

l'intégrale

TOUT-EN UN

MPSI | PTSI

Sous la direction de BERNARD **SALAMITO**

DAMIEN **JURINE**

STÉPHANE **CARDINI**

MARIE-NOËLLE **SANZ**

Physique tout-en-un

Avec la collaboration de :

EMMANUEL **ANGOT**

ANNE-EMMANUELLE **BADEL**

FRANÇOIS **CLAUSSET**

DUNOD

Conception et création de couverture : Atelier 3+

Les photos du chapitre 5 ont été réalisées par Pierre Canaguier.

Le pictogramme qui figure ci-contre mérite une explication. Son objet est d'alerter le lecteur sur la menace que représente pour l'avenir de l'écrit, particulièrement dans le domaine de l'édition technique et universitaire, le développement massif du photocopillage. Le Code de la propriété intellectuelle du 1^{er} juillet 1992 interdit en effet expressément la photocopie à usage collectif sans autorisation des ayants droit. Or, cette pratique s'est généralisée dans les établissements	d'enseignement supérieur, provoquant une baisse brutale des achats de livres et de revues, au point que la possibilité même pour les auteurs de créer des œuvres nouvelles et de les faire éditer correctement est aujourd'hui menacée. Nous rappelons donc que toute reproduction, partielle ou totale, de la présente publication est interdite sans autorisation de l'auteur, de son éditeur ou du Centre français d'exploitation du droit de copie (CFC, 20, rue des Grands-Augustins, 75006 Paris).
	

© Dunod, Paris, 2013
ISBN 978-2-10-070310-4

Le Code de la propriété intellectuelle n'autorisant, aux termes de l'article L. 122-5, 2° et 3° a), d'une part, que les « copies ou reproductions strictement réservées à l'usage privé du copiste et non destinées à une utilisation collective » et, d'autre part, que les analyses et les courtes citations dans un but d'exemple et d'illustration, « toute représentation ou reproduction intégrale ou partielle faite sans le consentement de l'auteur ou de ses ayants droit ou ayants cause est illicite » (art. L. 122-4).

Cette représentation ou reproduction, par quelque procédé que ce soit, constituerait donc une contrefaçon sanctionnée par les articles L. 335-2 et suivants du Code de la propriété intellectuelle.

Table des matières

I	Signaux physiques	21
1	Oscillateur harmonique	23
1	Un oscillateur harmonique mécanique	23
1.1	Système étudié	23
1.2	Obtention d’une équation différentielle	24
1.3	Définition d’un oscillateur harmonique	24
1.4	Résolution de l’équation différentielle	25
1.5	Conservation de l’énergie mécanique	27
1.6	Amplitude et période du mouvement	29
2	Signal sinusoïdal	30
2.1	Définition du signal sinusoïdal	30
2.2	Phase instantanée, phase initiale	30
2.3	Période, fréquence	31
2.4	Une interprétation géométrique	32
2.5	Représentation de Fresnel (MPSI)	33
2.6	Déphasage	34
2	Propagation d’un signal	47
1	Signaux physiques, spectre	47
1.1	Ondes et signaux physiques	47
1.2	Notion de spectre	48
1.3	Cas d’un signal périodique de forme quelconque	50
1.4	Cas d’un signal non périodique	53
1.5	Analyse harmonique expérimentale	53
1.6	Exemple : analyse de signaux sonores	55

TABLE DES MATIÈRES

2	Phénomène de propagation	57
2.1	Observations expérimentales	57
2.2	Onde progressive	58
2.3	Onde progressive sinusoïdale	62
3	Superpositions de deux signaux sinusoïdaux	79
1	Interférences entre deux ondes de même fréquence	79
1.1	Somme de deux signaux sinusoïdaux de même fréquence	79
1.2	Phénomène d’interférences	83
2	Ondes stationnaires et modes propres	88
2.1	Superposition de deux ondes progressives de même amplitude	88
2.2	Onde stationnaire	88
2.3	Expérience de la corde de Melde	90
2.4	Modes propres	94
4	Onde lumineuse	119
1	L’onde lumineuse	119
1.1	Existence et nature de l’onde lumineuse	119
1.2	Célérité de l’onde lumineuse	119
1.3	Longueurs d’onde et fréquences optiques	120
2	Récepteurs lumineux, éclairement	122
2.1	Comparaison avec les récepteurs d’onde sonore	122
2.2	Exemples de récepteurs d’onde lumineuse	123
2.3	Éclairement	124
2.4	Éclairement spectral	124
3	Les sources lumineuses	125
3.1	Les sources de lumière blanche	125
3.2	Les lampes spectrales	125
3.3	Faisceau laser	126
4	Rayon lumineux et source ponctuelle	126
4.1	Expérience	126
4.2	Définition d’un rayon lumineux	127
4.3	Propagation rectiligne	127
4.4	Modèle de la source ponctuelle et monochromatique	128
5	La diffraction de la lumière	129
5.1	Diffraction par une fente	129
5.2	Universalité du phénomène de diffraction	131

TABLE DES MATIÈRES

5	Optique géométrique	149
1	Approximation de l’optique géométrique	149
2	Lois de Descartes	150
2.1	Lois de Descartes pour la réflexion	150
2.2	Lois de Descartes pour la réfraction	151
3	Miroir plan	153
3.1	Image d’un point objet par un miroir plan	153
3.2	Image d’un objet par un miroir plan	155
4	Systèmes centrés et approximation de Gauss	155
4.1	Systèmes optiques centrés	155
4.2	Approximation de Gauss	155
4.3	Propriétés d’un système centré dans les conditions de Gauss	156
4.4	Foyers objet, foyers image	157
5	Lentilles minces	159
5.1	Présentation des lentilles	159
5.2	Constructions géométriques	161
5.3	Relations de conjugaison	167
5.4	Complément : démonstration des formules de conjugaison	168
6	Applications des lentilles	168
6.1	Projection d’une image	168
6.2	Le microscope	171
6.3	La lunette de Galilée	173
6.4	La lunette astronomique	174
7	L’œil	175
7.1	Description et modélisation	175
7.2	Caractéristiques optiques	176
7.3	Défauts de l’œil	176
8	Approche documentaire : influence des réglages sur l’image produite par un appareil photographique numérique	177
8.1	Modélisation	177
8.2	Influence du diaphragme d’ouverture	178
8.3	Influence de la distance focale	181
6	Introduction au monde quantique	207
1	La dualité onde-particule de la lumière	207
1.1	Introduction	207
1.2	Historique de la découverte du photon	208
1.3	Le photon	210

TABLE DES MATIÈRES

1.4	Une expérience avec des photons uniques	212
1.5	Franges d’interférences et photons	213
2	La dualité onde-particule de la matière	215
2.1	La longueur d’onde de de Broglie	215
2.2	Expériences d’interférences de particules	219
3	Fonction d’onde et probabilités	221
3.1	Analyse d’une expérience d’interférences quantiques	221
3.2	Notion de fonction d’onde et probabilité de détection	223
3.3	Interprétation de l’expérience des fentes de Young	223
3.4	Complémentarité	223
4	L’inégalité de Heisenberg (PTSI)	224
4.1	Indétermination quantique	224
4.2	Exemple : diffraction d’une particule par une fente	224
4.3	L’indétermination position-quantité de mouvement	225
5	Quantification de l’énergie d’une particule confinée	226
5.1	Notion de quantification	226
5.2	Particule dans un puits infini à 1 dimension	226
5.3	Généralisation : lien entre confinement spatial et quantification	228
7	Circuits électriques dans l’ARQS	243
1	Intensité du courant électrique	243
1.1	Charge électrique	243
1.2	Porteurs de charges	244
1.3	Le courant électrique	245
1.4	Intensité du courant	246
1.5	Mesure de l’intensité d’un courant	247
1.6	Loi des nœuds	248
1.7	Intensité constante, intensité variable	248
2	Tension électrique	248
2.1	Analogie hydraulique	248
2.2	Notion de tension	249
2.3	Mesure d’une tension	250
2.4	Tension constante, tension variable	250
2.5	Différence de potentiel et tension	250
2.6	Référence de potentiel	251
2.7	Additivité des tensions, loi des mailles	252
2.8	Notation simplifiée de la tension	253
3	Circuit électrique	253

TABLE DES MATIÈRES

3.1	Circuit fermé	253
3.2	Schéma électrique simplifié	254
3.3	La terre	254
3.4	Convention générateur, convention récepteur	254
3.5	L'approximation des régimes quasi-stationnaires	255
4	Dipôles	256
4.1	Définition	256
4.2	Dipôles passifs	256
4.3	Dipôles actifs	258
5	Associations de dipôles	259
5.1	Association de résistances	259
5.2	Associations de générateurs	261
5.3	Diviseur de tension	261
5.4	Diviseur de courant	262
6	Résistances de sortie et d'entrée	262
6.1	Résistance d'entrée	262
6.2	Résistance de sortie	263
7	Point de fonctionnement d'un circuit	263
7.1	Caractéristique d'un dipôle	263
7.2	Résolution graphique	264
7.3	Résolution algébrique	264
8	Puissance et énergie électriques	265
8.1	Les unités	265
8.2	Puissance électrique	265
8.3	Puissance Joule dans une résistance	266
8.4	Énergie stockée dans un condensateur ou une bobine	266
8	Circuit linéaire du premier ordre	285
1	Exemple expérimental	285
1.1	Montage	285
1.2	Régimes transitoire puis permanent	285
2	Modélisation	286
2.1	Simplification du montage	286
2.2	Équation différentielle sur $u_C(t)$	287
2.3	Unités et homogénéité de l'équation différentielle	287
2.4	Résolution de l'équation différentielle	288
2.5	Tracé et identification de la constante de temps	289
2.6	Portrait de phase	291

TABLE DES MATIÈRES

2.7	Bilan de puissance	291
2.8	Bilan d'énergie	292
3	Réalisation concrète de l'échelon	292
3.1	Réponse à un signal créneau	292
3.2	Modélisation du régime libre	293
4	Étude de la tension $u_R(t)$	294
4.1	Montage étudié et observations expérimentales	294
4.2	Équation différentielle sur $u_R(t)$	295
4.3	Portrait de phase	296
4.4	Solution de l'équation différentielle	296
4.5	Régime libre	297
5	Exemple de circuit inductif	298
5.1	Schéma du montage	298
5.2	Équation différentielle sur $i(t)$	299
5.3	Homogénéité de l'équation différentielle	299
5.4	Solution de l'équation différentielle	300
5.5	Bilan de puissance	301
6	Généralisation : systèmes du premier ordre	301
7	Régime permanent (PTSI)	301
9	Circuit linéaire du second ordre	313
1	Exemple expérimental	313
1.1	Montage	313
1.2	Régimes transitoire puis permanent	313
1.3	Portraits de phase	315
2	Équation différentielle sur la tension aux bornes du condensateur	316
2.1	Mise en équation	316
2.2	Forme canonique de l'équation différentielle	317
2.3	Conditions initiales	318
3	Solution de l'équation différentielle	319
3.1	Régime permanent ou établi	319
3.2	Solutions homogènes	319
3.3	Solution complète	321
4	Durée du régime transitoire	322
4.1	Définition du temps de réponse T_R	322
4.2	Complément : modélisation	323
5	Réponse à un signal créneaux	324
5.1	Observations expérimentales	324

TABLE DES MATIÈRES

5.2	Modélisation du régime libre	325
6	Bilan énergétique	326
7	Analogies	327
10	Régime sinusoïdal	337
1	Régimes transitoire et permanent sinusoïdal	337
2	Régime permanent ou établi	338
2.1	Observation à l’oscilloscope	338
2.2	Mesure d’un déphasage	339
3	Systèmes du premier ordre	340
3.1	Méthode des vecteurs de Fresnel (MPSI)	340
3.2	Méthode complexe	341
4	Impédance	344
4.1	Impédance d’un dipôle	344
4.2	Diviseur de tension	347
5	Résonance dans un système du deuxième ordre	348
5.1	Étude expérimentale de u_R	348
5.2	Interprétation	350
5.3	Complément : interprétation graphique du facteur de qualité	350
5.4	Remarque expérimentale	351
5.5	Étude de u_C	351
5.6	Interprétation	353
5.7	Mesures expérimentales	354
6	Fonction de transfert	355
6.1	Fonction de transfert harmonique	355
6.2	Lien entre équation différentielle et transmittance	355
6.3	Lien avec la transmittance de Laplace	356
7	Complément : déphasage et rapport des amplitudes dans un système du deuxième ordre	357
7.1	Position du problème	357
7.2	Méthode des vecteurs de Fresnel (MPSI)	358
7.3	Méthode complexe	359
11	Analyse fréquentielle d’un système linéaire	377
1	Représentation graphique de la fonction de transfert	377
1.1	Le gain en décibel	377
1.2	La phase	378
1.3	Relevés expérimentaux	378

TABLE DES MATIÈRES

1.4	Exemples de diagramme de Bode	379
1.5	Tracé à la calculatrice	381
2	Différents types de fonction de transfert	382
2.1	Filtres du premier ordre	382
2.2	Filtres du deuxième ordre	389
2.3	Récapitulatif	397
3	Gabarit (PTSI)	397
3.1	Amplitude	397
3.2	Phase	399
3.3	Exemple de réalisation d'un gabarit	400
3.4	Cahier des charges	401
12	Filtrage linéaire	409
1	Réponse d'un système linéaire en régime permanent	409
1.1	Légitimité de l'étude harmonique	409
1.2	Cadre de l'étude	409
1.3	Entrée sinusoïdale	410
1.4	Entrée combinaison linéaire de fonctions sinusoïdales	411
2	Contenu spectral d'un signal	412
2.1	Décomposition de Fourier	412
2.2	Complément : phénomène de Gibbs	416
3	Valeur efficace	417
3.1	Définition	417
3.2	Cas d'un signal sinusoïdal	417
3.3	Cas d'un signal créneau	417
3.4	Comparaison avec un signal constant	418
4	Filtrage linéaire d'un signal non sinusoïdal	419
4.1	Position du problème	419
4.2	Filtrage passe-bas	419
4.3	Réalisation d'un moyennneur	421
4.4	Filtrage passe-haut	421
4.5	Filtrage passe-bande	423
5	Réponse indicielle et contenu spectral	424
5.1	Possibilité de discontinuité, valeur moyenne	424
5.2	Complément : lien avec le théorème de la valeur initiale	427
6	Approche documentaire : accéléromètre	427

II Mécanique 1	453
13 Cinématique du point	455
1 Notion de point en physique	455
1.1 Définition d'un solide	455
1.2 Définition d'un point	455
1.3 Quand peut-on assimiler un système à un point ?	456
2 Repérage d'un point du plan	456
2.1 Intérêt d'avoir plusieurs systèmes de coordonnées	456
2.2 Repérage d'un point sur une droite.	458
2.3 Repérage d'un point dans le plan	459
3 Repérage d'un point dans l'espace	462
3.1 Repérage cartésien	462
3.2 Repérage cylindrique	463
3.3 Repérage sphérique	463
4 Cinématique du point	466
4.1 Notion de référentiel	466
4.2 Vecteurs position, déplacement, vitesse et accélération	468
5 Utilisation des différents systèmes de coordonnées	470
5.1 Coordonnées cartésiennes	470
5.2 Coordonnées cylindro-polaire	472
5.3 Coordonnées sphériques	477
6 Exemples de mouvements étudiés en coordonnées cartésiennes	479
6.1 Mouvements rectilignes	479
6.2 Mouvements à vecteur accélération constante	482
6.3 Mouvement rectiligne sinusoïdal : mouvement harmonique	484
7 Mouvements circulaires	485
7.1 Mouvement circulaire et uniforme	485
7.2 Généralisation : mouvement circulaire quelconque	486
8 Interprétation du vecteur accélération	487
8.1 Le vecteur vitesse et sa norme	487
8.2 Vecteur accélération et variation de la norme de la vitesse	488
8.3 Vecteur accélération et courbure de la trajectoire	489
9 Étude expérimentale de mouvements	490
9.1 Généralités	490
9.2 Étude expérimentale en coordonnées cartésiennes	491
9.3 Étude expérimentale en coordonnées polaires	495

TABLE DES MATIÈRES

14 Cinématique du solide	509
1 Repérage d'un solide	509
1.1 Définition d'un solide	509
1.2 Repérage d'un solide dans l'espace	509
2 Mouvement de translation	510
2.1 Définition	510
2.2 Mouvement d'un point d'un solide en translation	510
2.3 Conséquences	511
2.4 Deux mouvements de translations remarquables	511
3 Solides en rotation autour d'un axe fixe	512
3.1 Définition	512
3.2 Mouvement d'un point d'un solide en rotation	512
3.3 Conséquences	513
3.4 Quelques exemples de rotation autour d'un axe fixe	514
15 Principes de la dynamique newtonienne	521
1 Éléments cinétiques d'un point matériel	521
1.1 Masse	521
1.2 Quantité de mouvement	522
2 Les trois lois de Newton	523
2.1 Première loi de Newton : Principe d'inertie	523
2.2 Deuxième loi de Newton : Principe fondamental de la dynamique	524
2.3 Troisième loi de Newton : principe des actions réciproques	526
3 Limite de validité de la mécanique classique	527
3.1 Qu'est-ce qu'un principe ?	527
3.2 Les hypothèses de la mécanique classique	527
3.3 Les limites de la mécanique classique	527
4 Premières applications : détermination d'une loi de force	528
4.1 Détermination dynamique d'une force : mesure de g	528
4.2 Détermination statique d'une force	529
5 Classification des forces	529
5.1 Les quatre interactions fondamentales	530
5.2 Forces à distance	530
5.3 Forces de contact	534
6 Résolution d'un problème de mécanique du point	538
7 Chute libre dans le champ de pesanteur	539
7.1 Mise en équation	539
7.2 Chute libre dans le vide	539

TABLE DES MATIÈRES

7.3	Chute libre avec frottements proportionnels à la vitesse	541
7.4	Chute libre avec frottements proportionnels au carré de la vitesse . . .	543
7.5	Comparaison des deux modèles de frottements	545
8	Tir d'un projectile dans le champ de pesanteur	545
8.1	Mise en équation	545
8.2	Tir dans le vide	546
8.3	Tir en tenant compte de la résistance de l'air	549
9	Le pendule simple	551
9.1	Modélisation	551
9.2	Équation du mouvement	551
9.3	Résolution numérique	553
9.4	Cas des oscillations de faibles amplitudes	553
16	Aspects énergétiques de la dynamique du point	577
1	Travail et puissance d'une force	577
1.1	Introduction et notations	577
1.2	Puissance d'une force	578
1.3	Travail élémentaire d'une force	579
1.4	Travail d'une force au cours d'un déplacement	579
2	Premiers exemples de calculs de travaux	580
2.1	Travail d'une force constamment perpendiculaire au mouvement . . .	580
2.2	Travail d'une force constante	580
2.3	Travail d'une force de frottement de norme constante	580
3	Théorème de l'énergie cinétique	581
3.1	Définition de l'énergie cinétique	581
3.2	Théorème de l'énergie cinétique en référentiel galiléen	581
3.3	Utilisation du théorème de l'énergie cinétique	583
3.4	Intérêt d'une approche énergétique	583
3.5	Étude d'un problème à l'aide du théorème de l'énergie cinétique . . .	583
4	Énergie potentielle et forces conservatives	585
4.1	Définitions	585
4.2	Exemples de forces conservatives	586
4.3	Exemples de forces non conservatives	589
5	Énergie mécanique	589
5.1	Définition de l'énergie mécanique	589
5.2	Conservation de l'énergie mécanique	589
5.3	Cas général : non conservation de l'énergie mécanique	590
6	Étude qualitative des mouvements et des équilibres	591

TABLE DES MATIÈRES

6.1	Exemple introductif	591
6.2	Position du problème	591
6.3	Analyse du mouvement à l'aide d'un graphe énergétique	592
6.4	Analyse des équilibres à l'aide d'un graphe énergétique	593
7	Portraits de phase et lien avec le profil d'énergie potentielle	596
7.1	Définitions	596
7.2	Exemple introductif	597
7.3	Caractéristiques principales des portraits de phase	599
17	Mouvement dans un puits de potentiel	615
1	Mouvement conservatif dans un puits de potentiel	615
1.1	Mouvement dans un puits de potentiel harmonique	615
1.2	Mouvement dans un puits de potentiel quelconque	618
2	Mouvements dans un puits de potentiel : influence des frottements	622
2.1	Équation différentielle du mouvement	622
2.2	Équation caractéristique et observation	623
2.3	Résolution : les trois régimes	624
2.4	Évaluation rapide de la durée des différents régimes	628
2.5	Aspects énergétiques	629
18	Mouvement d'une particule chargée dans un champ électrique ou magnétique	651
1	Force de Lorentz	651
1.1	Rappel de l'expression	651
1.2	Différence fondamentale entre la composante électrique et la composante magnétique	652
1.3	Ordre de grandeur et conséquences	652
2	Mouvement dans un champ électrique uniforme	653
2.1	Équation du mouvement	653
2.2	Étude de la trajectoire	654
2.3	Accélération d'une particule chargée par un champ électrique	656
3	Mouvement dans un champ magnétique	659
3.1	Le mouvement est uniforme	659
3.2	Étude de la trajectoire	660
4	Quelques applications de ces mouvements	662
4.1	Expérience de Thomson	662
4.2	Spectromètre de masse	664
4.3	Cyclotron	666
5	Approche documentaire : limites relativistes en microscopie électronique	667

III Mécanique 2	683
19 Moment cinétique et solide en rotation	685
1 Observations préliminaires	685
1.1 Exemples introductifs	685
1.2 Notion intuitive de bras de levier	686
2 Moment cinétique d'un point matériel	686
2.1 Définition du moment cinétique	686
3 Moment cinétique d'un solide ou d'un système de points	689
3.1 Cas d'un système déformable	689
3.2 Cas d'un solide en rotation par rapport à un axe	690
4 Moment d'une force	692
4.1 Moment d'une force par rapport à un point O	692
4.2 Moment d'une force par rapport à un axe orienté Δ	693
5 Loi du moment cinétique pour un point matériel	695
5.1 Loi du moment cinétique par rapport à un point fixe	695
5.2 Cas de conservation du moment cinétique	696
5.3 Loi du moment cinétique par rapport à un axe fixe	696
6 Loi du moment cinétique pour un solide en rotation	697
6.1 Loi scalaire du moment cinétique pour un solide	697
6.2 Cas de conservation du moment cinétique	698
6.3 Couples	698
7 Application aux dispositifs rotatifs	699
7.1 Liaison pivot d'axe (Oz)	700
7.2 Notions sur les moteurs et les freins dans les dispositifs rotatifs (PTSI)	701
8 Pendule pesant	702
8.1 Position du problème et équation du mouvement	702
8.2 Oscillations de faible amplitude	703
8.3 Intégrale première du mouvement et étude qualitative	703
8.4 Portrait de phase	704
8.5 Résolution numérique	706
9 Énergie d'un solide en rotation autour d'un axe fixe	706
9.1 Énergie cinétique d'un solide en rotation	707
9.2 Puissance d'une force appliquée sur un solide en rotation	707
9.3 Loi de l'énergie cinétique pour un solide indéformable	708
20 Mouvement dans un champ de force centrale. Champs newtoniens.	717
1 Force centrale conservative	717

TABLE DES MATIÈRES

- 1.1 Qu'est-ce qu'une force centrale conservative ? 717
 - 1.2 Exemples de forces centrales conservatives 718
 - 1.3 Observations de mouvements à force centrale conservative 720
- 2 Généralités sur les forces centrales conservatives 721
 - 2.1 Conséquence du caractère central de la force 721
 - 2.2 Conséquence du caractère conservatif de la force 724
- 3 Cas particulier de l'attraction gravitationnelle 725
 - 3.1 Position du problème 725
 - 3.2 Étude qualitative du mouvement radial 726
- 4 Étude directe de la trajectoire circulaire 727
 - 4.1 Position du problème 727
 - 4.2 Étude à partir du principe fondamental de la dynamique 728
 - 4.3 Application aux satellites géostationnaires 730
 - 4.4 Vitesses cosmiques 733
 - 4.5 Complément : autres trajectoires envisageables (MPSI) 735

IV Thermodynamique 753

- 21 Système thermodynamique à l'équilibre 755**
 - 1 Descriptions microscopique et macroscopique de la matière 755
 - 1.1 Les phases solide, liquide et gaz 755
 - 1.2 L'agitation thermique 756
 - 1.3 Échelles microscopique, mésoscopique et macroscopique 757
 - 1.4 Le point de vue de la thermodynamique 757
 - 2 Système thermodynamique, variables d'état 758
 - 2.1 Système thermodynamique 758
 - 2.2 Variables d'état 759
 - 3 Équilibre thermodynamique 762
 - 3.1 Définition 762
 - 3.2 Équilibre thermodynamique local 762
 - 3.3 Conditions d'équilibre 762
 - 4 Équation d'état 764
 - 4.1 Définition 764
 - 4.2 Équation d'état d'un gaz parfait 764
 - 4.3 Équation d'état d'une phase condensée idéale 766
 - 5 Énergie interne, capacité thermique à volume constant 767
 - 5.1 L'énergie interne U 767

TABLE DES MATIÈRES

5.2	La capacité thermique à volume constant C_V	768
5.3	Cas d'un gaz parfait	768
5.4	Cas d'une phase condensée incompressible	769
6	Vitesse quadratique moyenne (PTSI)	770
6.1	Définition	770
6.2	Expression en fonction de la température	771
6.3	Ordre de grandeur de la vitesse quadratique moyenne	771
7	Corps pur diphasé en équilibre	772
7.1	Changements d'état physique	772
7.2	Diagramme de phases (P,T)	772
7.3	Variables d'état d'un système diphasé	775
8	Étude de l'équilibre liquide-gaz	777
8.1	Pression de vapeur saturante	777
8.2	Variation de P_{sat} avec T	778
8.3	Température d'ébullition	778
8.4	Diagramme de Clapeyron	779
8.5	Composition du mélange liquide-gaz	780
8.6	Point critique	781
8.7	Le stockage des fluides	782
22	Énergie échangée par un système au cours d'une transformation	795
1	Transformation thermodynamique	795
1.1	Transformation, état initial, état final	795
1.2	Différents types de transformations	796
1.3	Influence du choix du système	798
2	Travail des forces de pression	799
2.1	Expression générale du travail de la pression extérieure	799
2.2	Cas particulier d'un fluide en écoulement	802
2.3	Travail des forces de pression dans deux cas particuliers	803
2.4	Travail des forces de pression dans le cas d'une transformation mé- caniquement réversible	804
3	Transfert thermique	807
3.1	Définition	807
3.2	Les trois modes de transfert thermique (PTSI)	807
3.3	Transformation adiabatique	809
3.4	Notion de thermostat	810
3.5	Retour sur les transformations monotherme et isotherme	811
3.6	Choix d'un modèle : adiabatique ou isotherme ?	811

TABLE DES MATIÈRES

23	Premier principe. Bilans d'énergie.	819
1	Le premier principe de la thermodynamique	819
1.1	Énergie d'un système	819
1.2	Premier principe de la thermodynamique	820
1.3	Obtention de la valeur du transfert thermique	821
1.4	Transfert thermique dans une transformation isochore sans travail autre que celui de la pression et sans variation d'énergie cinétique . . .	822
1.5	Exemples d'application du premier principe	822
2	La fonction d'état enthalpie	826
2.1	Définitions	826
2.2	Premier principe pour une transformation monobare avec équilibre mécanique dans l'état initial et l'état final	827
2.3	Transfert thermique dans une transformation isobare sans travail autre que celui de la pression et sans variation d'énergie cinétique	827
2.4	Enthalpie d'un gaz parfait	828
2.5	Enthalpie d'une phase condensée indilatable et incompressible	830
2.6	Enthalpie d'un système diphasé	831
2.7	Variations d'enthalpie isobares	832
3	Mesures de grandeurs thermodynamiques	835
3.1	Le calorimètre	835
3.2	Détermination d'une capacité thermique massique	836
3.3	Détermination d'une enthalpie de changement d'état	837
3.4	Mesure de la valeur en eau du calorimètre	838
24	Deuxième principe. Bilans d'entropie.	849
1	Le deuxième principe de la thermodynamique	849
1.1	Transformations irréversibles et transformations réversibles	849
1.2	Le deuxième principe de la thermodynamique	852
2	Entropie d'un échantillon de corps pur	853
2.1	Entropie d'un gaz parfait	854
2.2	Entropie d'une phase condensée indilatable et incompressible	857
2.3	Entropie d'un système diphasé	858
3	Exemples de bilans d'entropie	860
3.1	Méthode générale	860
3.2	Exemple 1 : détente de Joule - Gay Lussac	860
3.3	Exemple 2 : mise en contact avec un thermostat	861
3.4	Exemple 3 : compression d'un gaz parfait	863
3.5	Exemple 4 : chauffage par effet Joule	866

TABLE DES MATIÈRES

3.6 Exemple 5 : solidification d’un liquide surfondu 867

25 Machines thermiques 885

1 Machine monotherme 885

2 Machines thermiques dithermes 886

2.1 Généralités sur les machines dithermes 886

2.2 Moteur thermique 887

2.3 Machine frigorifique 889

2.4 Pompe à chaleur 890

3 Étude de cycles théoriques réversibles 891

3.1 Cycle de Carnot pour un gaz parfait 891

3.2 Cycle de Carnot pour un système diphasé 893

4 Étude de machines thermiques réelles 894

4.1 Moteur à explosion 894

4.2 Machine frigorifique (MPSI) 897

V Induction et forces de Laplace 927

26 Le champ magnétique 929

1 Carte de champ 929

1.1 Les champs en physique 929

1.2 Un champ vectoriel permet de décrire une interaction à distance . . . 930

1.3 Unités et ordres de grandeur 930

1.4 Topographie du champ magnétique 931

1.5 Quelques cartes de champ magnétique 932

2 Moment magnétique 937

2.1 Vecteur surface 937

2.2 Définition du moment magnétique 937

2.3 Moment magnétique d’un aimant 937

2.4 Lignes de champ d’un moment magnétique 938

27 Actions d’un champ magnétique 943

1 Force de Laplace 943

1.1 Force de Laplace sur une tige en translation 943

1.2 Puissance de la force Laplace 944

2 Couple magnétique 944

2.1 Expression du Couple 944

2.2 Puissance de l’action de Laplace 944

TABLE DES MATIÈRES

2.3	Complément : établissement du couple	944
3	Action d'un champ magnétique sur un aimant	946
3.1	Orientation d'un aimant	946
3.2	Positions d'équilibre	947
3.3	Application : la boussole	947
3.4	Effet moteur d'un champ magnétique tournant	948
28	Lois de l'induction	957
1	Flux magnétique	957
1.1	Définition du flux magnétique	957
1.2	Orientation d'une surface	958
1.3	Unité de flux magnétique	959
2	Expériences d'induction électromagnétique	959
2.1	Expérience historique de Faraday	959
2.2	Expériences avec un aimant et une bobine	960
2.3	Le phénomène d'induction électromagnétique	961
2.4	Loi de Lenz	961
3	Loi de Faraday	962
3.1	Règle du flux	962
3.2	Convention d'algébrisation	962
3.3	Exceptions à la règle du flux	962
3.4	Loi de Faraday	963
29	Circuit fixe dans un champ magnétique variable	971
1	Auto-induction	971
1.1	Inductance propre	971
1.2	Calcul d'une inductance propre	973
1.3	Circuit électrique équivalent	974
1.4	Loi de Lenz	974
1.5	Mesure d'une inductance propre	975
1.6	Étude énergétique	975
2	Deux circuits en interaction	976
2.1	Inductance mutuelle	976
2.2	Circuits électriques équivalents	977
2.3	Étude harmonique	978
2.4	Étude énergétique	979
3	Transformateur de tension (PTSI)	980
3.1	Constitution	980

TABLE DES MATIÈRES

3.2	Principe de fonctionnement	980
3.3	Normalisation et orientation des courants	981
3.4	Courants de Foucault	983
3.5	Utilisation	983
30	Circuit mobile dans un champ magnétique stationnaire	999
1	Conversion de puissance mécanique en puissance électrique	999
1.1	Rails de Laplace générateurs	999
1.2	Freinage par induction	1004
1.3	Alternateur	1005
2	Conversion de puissance électrique en puissance mécanique	1008
2.1	Rails de Laplace moteurs	1008
2.2	Haut-parleur électrodynamique	1011
2.3	Machine à courant continu à entrefer plan (PTSI)	1014
VI	Appendices	1045
A	Mesures et incertitudes	1047
1	Mesure d’une grandeur physique	1047
1.1	Représentation d’une grandeur physique	1047
1.2	Mesure d’une grandeur physique	1048
2	Incertitudes et intervalle de confiance	1051
2.1	Notion d’intervalle de confiance	1051
2.2	Évaluation d’une incertitude-type	1052
2.3	Incertitude-type composée	1055
3	Présentation d’un résultat expérimental	1056
3.1	Notation d’un résultat	1056
3.2	Chiffres significatifs et arrondis	1057
4	Validité d’un résultat expérimental	1058
4.1	Comparaison entre une valeur mesurée et une valeur de référence	1058
4.2	Vérification d’une relation linéaire entre des données	1058
B	Outils mathématiques	1063
1	Équations algébriques	1063
1.1	Système linéaire de n équations à p inconnues	1063
1.2	Équation non linéaire	1064
2	Équations différentielles	1066
2.1	Équation différentielle linéaire d’ordre 1 à coefficients constants	1066

TABLE DES MATIÈRES

2.2	Équation différentielle linéaire homogène d'ordre 2 à coefficients constants	1070
2.3	Équation différentielle linéaire d'ordre 2 à coefficients constants, avec un second membre non nul	1074
2.4	Autres équations différentielles	1077
3	Fonctions	1080
3.1	Fonctions usuelles	1080
3.2	Dérivée	1082
3.3	Développements limités	1084
3.4	Primitive et intégrale	1085
3.5	Représentation graphique d'une fonction	1088
3.6	Développement en série de Fourier	1089
4	Géométrie	1091
4.1	Projection d'un vecteur, produit scalaire	1091
4.2	Produit vectoriel	1093
4.3	Transformations géométriques	1095
4.4	Courbes planes	1098
4.5	Courbes paramétrées	1102
4.6	Longueurs, aires et volumes classiques	1104
4.7	Barycentre d'un système de points	1105
5	Trigonométrie	1107
5.1	Angle orienté	1107
5.2	Fonctions trigonométriques	1108
5.3	Nombres complexes	1111
6	Gradient d'un champ scalaire	1113
6.1	Champ scalaire f et surfaces iso- f	1113
6.2	Dérivées partielles et différentielle	1113
6.3	Vecteur gradient	1114

Première partie

Signaux physiques

Trouver la bibliothèque des livres ici : <https://1024terabox.com/s/1mg0wxl22G8NjM8MA2RzgFw>

Oscillateur harmonique

1

On appelle **signal physique** une grandeur physique dépendant du temps. Dans cette partie du cours, on s'intéressera surtout à des **signaux périodiques**. Un signal périodique est un signal qui se reproduit identique à lui-même au cours du temps. Le plus fondamental des signaux périodiques est le **signal sinusoïdal**.

Dans ce chapitre, on introduit un modèle physique qui produit un signal sinusoïdal appelé l'**oscillateur harmonique**. Les adjectifs harmonique et sinusoïdal sont synonymes : on rencontre parfois les expressions « signal harmonique » ou « oscillateur sinusoïdal ».

L'exemple étudié dans ce chapitre est un oscillateur harmonique mécanique. Cependant on retrouve ce modèle dans bien d'autres domaines de la physique, notamment l'électricité.

1 Un oscillateur harmonique mécanique

1.1 Système étudié

Le système mécanique oscillant le plus simple est une masse accrochée à un ressort.

On considère dans ce paragraphe un mobile de masse m qui se déplace sans frottement le long d'une tige horizontale (figure 1.1). Sa position est repérée par l'abscisse x de son centre d'inertie G mesurée sur l'axe (Ox) matérialisé par la tige. On choisit de placer l'origine de l'axe (Ox) de manière que G coïncide avec O dans la position d'équilibre (voir figure 1.1).

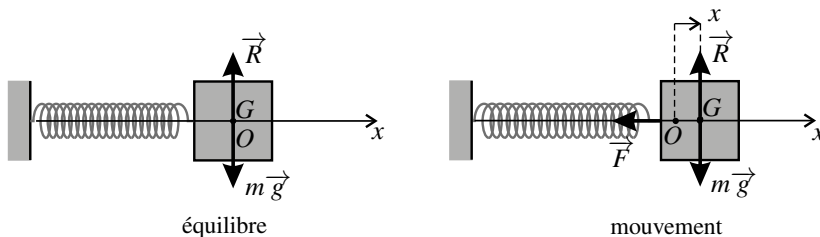


Figure 1.1 – Un exemple d'oscillateur harmonique mécanique.

CHAPITRE 1 – OSCILLATEUR HARMONIQUE

1.2 Obtention d’une équation différentielle

a) Forces s’exerçant sur le système

Le ressort exerce sur le mobile une force qui s’écrit :

$$\vec{F} = -kx\vec{u}_x,$$

où k est la **constante de raideur** du ressort. Cette force est une force de rappel : son sens est opposé au sens du déplacement du mobile par rapport à sa position d’équilibre (la position $x = 0$). Lorsque x est positif, la force est de sens opposé à $\vec{u}_x$ et inversement. On trouvera plus de renseignements sur la force d’un ressort dans le chapitre *Principes de la dynamique newtonienne*.

Le mobile est aussi soumis à son poids $m\vec{g}$ ainsi qu’à une réaction $\vec{R}$ de la tige. Ces deux forces sont verticales et n’ont pas d’influence sur le mouvement qui est horizontal.

b) Application du principe fondamental de la dynamique

Le mouvement du mobile est régi par le principe fondamental de la dynamique (ou troisième loi de Newton), qui s’écrit :

$$\frac{d\vec{p}}{dt} = \vec{F}$$

où $\vec{p}$ est la quantité de mouvement du mobile. En notant $\vec{v}_G$ la vitesse instantanée du centre d’inertie G du mobile, on a :

$$\vec{p} = m\vec{v}_G = m\frac{dx}{dt}\vec{u}_x.$$

Dans cette relation, $\frac{dx}{dt}$ est la dérivée de la fonction $x(t)$. On notera $\frac{d^2x}{dt^2}$ la dérivée seconde de cette fonction. Le principe fondamental de la dynamique conduit donc à la relation, vérifiée à chaque instant : $\frac{d}{dt}\left(m\frac{dx}{dt}\vec{u}_x\right) = -kx\vec{u}_x$, soit : $m\frac{d^2x}{dt^2}\vec{u}_x = -kx\vec{u}_x$, soit après simplification :

$$\frac{d^2x}{dt^2} = -\frac{k}{m}x. \tag{1.1}$$

La relation qui vient d’être obtenue est appelée **équation différentielle**. Une équation différentielle est une relation entre une fonction $x(t)$ et ses dérivées par rapport au temps $\frac{dx}{dt}$, $\frac{d^2x}{dt^2}$..., qui est vérifiée à chaque instant. Ici, on a trouvé une équation du deuxième ordre, car l’ordre de dérivation le plus grand qui apparaît dans l’équation (le seul ici, mais il pourrait y en avoir plusieurs) est égal à deux.

1.3 Définition d’un oscillateur harmonique

L’équation différentielle (1.1) est une équation différentielle d’oscillateur harmonique.

On appelle **oscillateur harmonique** un système physique décrit par une grandeur $x(t)$ dépendant du temps et vérifiant une équation différentielle de la forme :

$$\frac{d^2x}{dt^2} = -\omega_0^2 x(t) \tag{1.2}$$

où ω_0 est une constante réelle positive qui est appelée **pulsation propre** de l'oscillateur harmonique et qui s'exprime en rad.s^{-1} .

Remarque

L'unité de ω_0 se déduit de l'homogénéité de l'équation différentielle. En effet, la dérivée seconde par rapport au temps a pour dimension physique la dimension de x divisée par un temps au carré. Donc ω_0^2 a la dimension des s^{-2} et ω_0 a la dimension des s^{-1} . Le radian n'a pas de dimension physique.

La masse accrochée au ressort du paragraphe précédent est un oscillateur harmonique mécanique, de pulsation :

$$\omega_0 = \sqrt{\frac{k}{m}}.$$

On en verra d'autres exemples, en mécanique ou en électricité.

1.4 Résolution de l'équation différentielle

a) Position du problème

Résoudre une équation différentielle consiste à trouver l'expression de la fonction inconnue $x(t)$ qui vérifie cette relation. Mais l'équation ne détermine pas de manière unique $x(t)$. Parmi les fonctions qui la vérifient, on doit choisir celle qui respecte des **conditions initiales** qui sont connues *a priori*.

Pour une équation du deuxième ordre, comme l'équation d'un oscillateur harmonique, les conditions initiales consistent en la donnée :

- de la valeur de la fonction inconnue à l'instant initial $t = 0$: $x(0)$,
- de la valeur de la dérivée première de la fonction inconnue à l'instant initial : $\left(\frac{dx}{dt}\right)_{t=0}$.

On va résoudre cette équation dans le cas du mobile accroché au ressort. Les conditions initiales sont :

- la position initiale : $x(0) = x_0$,
- la vitesse initiale : $\left(\frac{dx}{dt}\right)_{t=0} = v_0$.

b) Solutions dans deux cas particuliers

Cas où $x_0 \neq 0$ et $v_0 = 0$. On considère la fonction :

$$x(t) = x_0 \cos(\omega_0 t) \tag{1.3}$$

CHAPITRE 1 – OSCILLATEUR HARMONIQUE

où x_0 est une constante. On a alors, en utilisant des formules classiques de dérivation (voir l'appendice mathématique) :

$$\frac{dx}{dt} = -\omega x_0 \sin(\omega_0 t), \quad \frac{d^2x}{dt^2} = \frac{d}{dt}(-\omega x_0 \sin(\omega_0 t)) = -\omega_0^2 x_0 \cos(\omega_0 t) = -\omega_0^2 x(t) = -\frac{k}{m}x(t).$$

Ainsi cette fonction $x(t)$ vérifie l'équation différentielle (1.1).

À quelle situation physique la solution trouvée correspond-elle ? Les conditions initiales vérifiées sont :

$$x(0) = x_0 \cos(0) = x_0 \quad \text{et} \quad \left(\frac{dx}{dt}\right)_{t=0} = \omega x_0 \sin(0) = 0.$$

Il s'agit du cas où on lâche le mobile dans la position $x = x_0$ sans lui communiquer de vitesse initiale. Dans ce cas, $x(t)$ oscille entre $-x_0$ et x_0 puisque la fonction cosinus oscille entre -1 et 1 (voir figure 1.2).

Cas où $x_0 = 0$ et $v_0 \neq 0$ On peut imaginer une autre manière de lancer le mobile : sans l'écarter de sa position d'équilibre, lui communiquer une vitesse $v_0 \vec{u}_x$. Quelle est alors la loi du mouvement $x(t)$? Les conditions initiales que doit vérifier cette solution sont :

$$x(0) = 0 \quad \text{et} \quad \left(\frac{dx}{dt}\right)_{t=0} = v_0.$$

La fonction : $x(t) = b \sin(\omega_0 t)$, où b est une constante est aussi une solution de l'équation différentielle du mouvement comme le montre le calcul suivant :

$$\frac{dx}{dt} = \omega_0 b \cos(\omega_0 t), \quad \frac{d^2x}{dt^2} = \frac{d}{dt}(\omega_0 b \cos(\omega_0 t)) = -\omega_0^2 b \sin(\omega_0 t) = -\omega_0^2 x(t) = -\frac{k}{m}x(t).$$

Elle vérifie les conditions initiales si : $x(0) = b \sin(0) = 0$, et si : $\left(\frac{dx}{dt}\right)_{t=0} = \omega_0 b \cos(0) = v_0$.

Il faut donc que : $b = \frac{v_0}{\omega}$. Finalement, la solution de l'équation différentielle correspondant à ces conditions initiales est :

$$x(t) = \frac{v_0}{\omega_0} \sin(\omega_0 t). \tag{1.4}$$

Cette solution $x(t)$ oscille entre $-\frac{v_0}{\omega_0}$ et $\frac{v_0}{\omega_0}$ puisque la fonction sinus oscille entre -1 et 1 (voir figure 1.2).

c) Solution pour des conditions initiales quelconques

On peut vérifier facilement que la somme des deux solutions précédemment trouvées,

$$x(t) = x_0 \cos(\omega_0 t) + \frac{v_0}{\omega_0} \sin(\omega_0 t), \tag{1.5}$$

est aussi une solution de l'équation différentielle et qu'elle vérifie les conditions initiales $x(0) = x_0$ et $\left(\frac{dx}{dt}\right)_{t=0} = v_0$. On a ainsi la solution pour un jeu de conditions initiales quelconque. Sur la figure 1.2 on voit que $x(t)$ oscille entre deux valeurs opposées $-A$ et A .

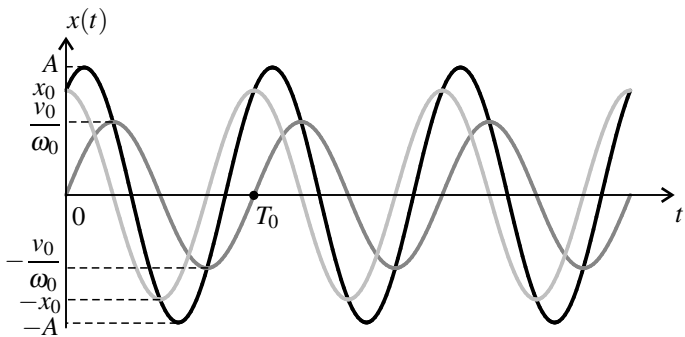


Figure 1.2 – Représentation graphique de $x(t)$ en fonction de t . En gris clair : cas $x_0 \neq 0$ et $v_0 = 0$; en gris foncé : cas $x_0 = 0$ et $v_0 \neq 0$; en noir : cas $x_0 \neq 0$ et $v_0 \neq 0$. La période des oscillations est $T_0 = \frac{2\pi}{\omega_0} = 2\pi\sqrt{\frac{m}{k}}$ (voir paragraphe 2).

d) Généralisation

La solution de l'équation de l'oscillateur harmonique (1.2) est de la forme :

$$x(t) = a \cos(\omega_0 t) + b \sin(\omega_0 t)$$

où a et b sont des constantes fixées par les conditions initiales.

1.5 Conservation de l'énergie mécanique

Lorsqu'il est en mouvement le mobile possède une **énergie cinétique** qui se calcule par la formule :

$$E_c = \frac{1}{2}mv^2 = \frac{1}{2}m\left(\frac{dx}{dt}\right)^2.$$

Le ressort quant à lui n'a pas de masse donc pas d'énergie cinétique, mais il possède une énergie appelée **énergie potentielle** liée à sa déformation et dont on admettra ici l'expression :

$$E_p = \frac{1}{2}kx^2.$$

La somme de ces deux énergies est l'**énergie mécanique** :

$$E_m = E_p + E_c = \frac{1}{2}kx^2 + \frac{1}{2}m\left(\frac{dx}{dt}\right)^2.$$

Que valent ces énergies au cours du temps ?

Si l'on injecte la solution $x(t) = x_0 \cos(\omega_0 t) + \frac{v_0}{\omega_0} \sin(\omega_0 t)$ dans les expressions précédentes

CHAPITRE 1 – OSCILLATEUR HARMONIQUE

on trouve :

$$\begin{aligned} E_c(t) &= \frac{1}{2}m(-\omega_0 x_0 \sin(\omega_0 t) + v_0 \cos(\omega_0 t))^2 \\ &= \frac{1}{2}m(\omega_0^2 x_0^2 \sin^2(\omega_0 t) + v_0^2 \cos^2(\omega_0 t) - 2\omega_0 x_0 v_0 \cos(\omega_0 t) \sin(\omega_0 t)), \end{aligned}$$

et :

$$\begin{aligned} E_p(t) &= \frac{1}{2}k \left(x_0 \cos(\omega_0 t) + \frac{v_0}{\omega_0} \sin(\omega_0 t) \right)^2 \\ &= \frac{1}{2}k \left(x_0^2 \cos^2(\omega_0 t) + \frac{v_0^2}{\omega_0^2} \sin^2(\omega_0 t) + 2x_0 \frac{v_0}{\omega_0} \cos(\omega_0 t) \sin(\omega_0 t) \right) \\ &= \frac{1}{2}m(\omega_0^2 x_0^2 \cos^2(\omega_0 t) + v_0^2 \sin^2(\omega_0 t) + 2\omega_0 x_0 v_0 \cos(\omega_0 t) \sin(\omega_0 t)), \end{aligned}$$

en utilisant la relation $k = m\omega_0^2$. L'énergie mécanique s'écrit :

$$\begin{aligned} E_m(t) &= E_p(t) + E_c(t) \\ &= \frac{1}{2}m(\omega_0^2 x_0^2 (\cos^2(\omega_0 t) + \sin^2(\omega_0 t)) + v_0^2 (\cos^2(\omega_0 t) + \sin^2(\omega_0 t))) \\ &= \frac{1}{2}m(\omega_0^2 x_0^2 + v_0^2) = \frac{1}{2}k \left(x_0^2 + \frac{v_0^2}{\omega_0^2} \right) = \frac{1}{2}kx_0^2 + \frac{1}{2}mv_0^2. \end{aligned} \quad (1.6)$$

en utilisant encore la relation $k = m\omega_0^2$. L'énergie mécanique est donc constante dans le temps. On reconnaît dans la dernière expression la valeur de l'énergie mécanique à l'instant initial. Ce résultat est cohérent avec le fait que l'on étudie un système idéalisé dont l'amortissement est négligé : on ne prend en compte aucun type de frottement. C'est pourquoi il y a conservation de l'énergie mécanique.

Remarque

Dans les deux cas particuliers du paragraphe b) on fournit son énergie au système :

- sous forme d'énergie potentielle uniquement quand $x_0 \neq 0$ et $v_0 = 0$;
- sous forme d'énergie cinétique uniquement quand $x_0 = 0$ et $v_0 \neq 0$.



L'équation différentielle (1.1) peut être établie à partir de la relation de conservation de l'énergie mécanique : $E_m = \frac{1}{2}m \left(\frac{dx}{dt} \right)^2 + \frac{1}{2}kx(t)^2 = \text{constante}$. En effet, il vient si l'on dérive par rapport au temps :

$$m \frac{dx}{dt} \frac{d^2x}{dt^2} + kx(t) \frac{dx}{dt} = 0,$$

ce qui redonne l'équation différentielle (1.1), après simplification par le terme $\frac{dx}{dt}$.

1.6 Amplitude et période du mouvement

a) Amplitude du mouvement

L'amplitude A du mouvement est la valeur maximale atteinte par $x(t)$. Dans le cas où $x_0 \neq 0$ et $v_0 = 0$, $A = x_0$; dans le cas où $x_0 = 0$ et $v_0 \neq 0$, $A = \frac{v_0}{\omega_0}$. Quelle est l'expression de A dans le cas général où $x_0 \neq 0$ et $v_0 \neq 0$?

On peut utiliser la conservation de l'énergie : lorsque $x(t)$ passe par la valeur maximale A , sa dérivée est nulle donc l'énergie cinétique est nulle. Par suite :

$$E_m = \frac{1}{2}kA^2 + 0.$$

L'expression (1.6) de l'énergie mécanique conduit alors à :

$$A = \sqrt{x_0^2 + \frac{v_0^2}{\omega_0^2}}.$$

b) Période

Comme on l'observe sur la figure 1.2, le mouvement du mobile est oscillatoire et périodique : les valeurs de $x(t)$ se répètent à intervalle régulier. Mathématiquement cela provient de la périodicité des fonctions cosinus et sinus : $\cos(\theta + 2\pi) = \cos \theta$ et $\sin(\theta + 2\pi) = \sin \theta$. Alors :

$$\cos(\omega_0 t) = \cos(\omega_0 t + 2\pi) = \cos\left(\omega_0 \left(t + \frac{2\pi}{\omega_0}\right)\right),$$

et de même : $\sin(\omega_0 t) = \sin\left(\omega_0 \left(t + \frac{2\pi}{\omega_0}\right)\right)$. Ainsi, on a :

$$x(t) = x(t + T_0) \quad \text{avec} \quad T_0 = \frac{2\pi}{\omega_0} = 2\pi\sqrt{\frac{m}{k}}. \quad (1.7)$$

Cette relation signifie que le mouvement est périodique de période T_0 . La période ne dépend pas de l'amplitude du mouvement, c'est la propriété d'**isochronisme** des oscillations de l'oscillateur harmonique. On peut alors parler de période de l'oscillateur harmonique. T_0 augmente avec la masse m du mobile et diminue avec la constante de raideur k du ressort, ce qui est intuitif.

2 Signal sinusoïdal

2.1 Définition du signal sinusoïdal

Un **signal sinusoïdal** est un signal de la forme :

$$s(t) = A \cos(\omega t + \varphi).$$

où A et ω sont des constantes positives et φ une constante.

ω est la **pulsation** du signal, A son **amplitude**, et φ sa **phase initiale**.



Un signal sinusoïdal peut aussi être écrit sous la forme :

$$s(t) = a \cos(\omega t) + b \sin(\omega t),$$

où a et b sont deux constantes. C'est sous cette forme qu'on l'obtient naturellement en résolvant l'équation différentielle de l'oscillateur harmonique.

Il faut savoir trouver la valeur de l'amplitude et de la phase initiale de ce signal. La relation de trigonométrie $\cos(\alpha + \beta) = \cos \alpha \cos \beta - \sin \alpha \sin \beta$ (voir appendice mathématique) permet d'écrire : $A \cos(\omega t + \varphi) = A \cos \varphi \cos(\omega t) - A \sin \varphi \sin(\omega t)$. On en déduit :

$$\cos \varphi = \frac{a}{A} \quad \text{et} \quad \sin \varphi = -\frac{b}{A} \quad \text{d'où} \quad A = \sqrt{a^2 + b^2}, \quad (1.8)$$

puisque $\cos^2 \varphi + \sin^2 \varphi = 1$.

Pour déterminer φ on peut utiliser la méthode suivante qui donne une valeur comprise entre $-\pi$ et π : il s'agit de $\arccos\left(\frac{a}{A}\right)$ si $\sin \varphi > 0$ et de $-\arccos\left(\frac{a}{A}\right)$ si $\sin \varphi < 0$.

2.2 Phase instantanée, phase initiale

L'argument de la fonction cosinus est appelée **phase instantanée**.

Le signal oscille entre $-A$ et A : il vaut A aux instants où la phase instantanée est égale à $2n\pi$ où n est un entier, il vaut $-A$ aux instants où elle est égale à $(2n+1)\pi$, il vaut 0 et a une pente négative aux instants où elle est égale à $\left(2n + \frac{1}{2}\right)\pi$, et il vaut 0 et a une pente positive aux instants où elle est égale à $\left(2n - \frac{1}{2}\right)\pi$. Ceci est illustré sur la figure 1.3.

La **phase initiale** φ donne la valeur de départ du signal à $t = 0$. Elle dépend de l'origine des temps choisie. De plus, le cosinus étant une fonction périodique de période 2π , la phase initiale n'est définie qu'à un multiple entier de 2π près. On peut ainsi toujours se ramener à une phase initiale comprise entre $-\pi$ et π .

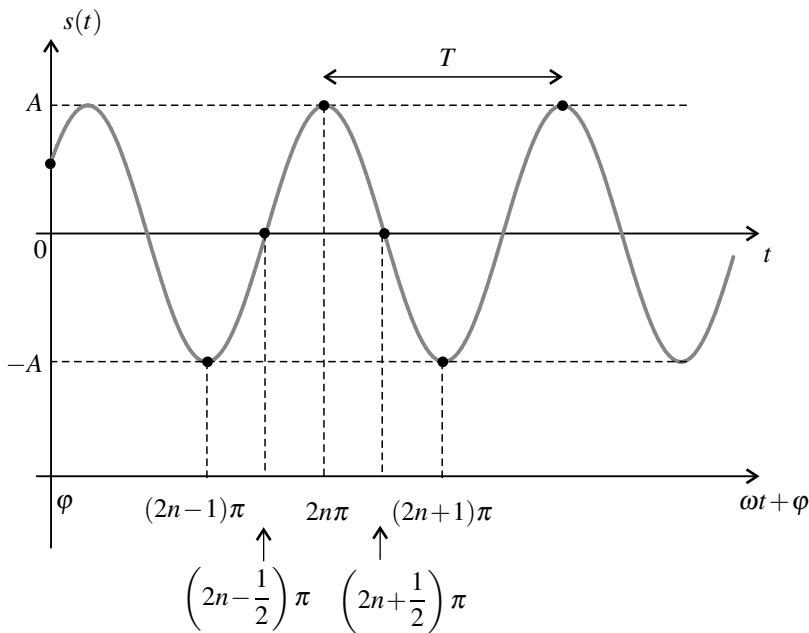


Figure 1.3 – Signal sinusoïdal. La valeur de la phase instantanée $\omega t + \varphi$ est indiquée en certains points, n est un entier.

2.3 Période, fréquence

a) Définitions

Un signal physique $s(t)$ est **périodique** s’il se répète dans le temps. Sa **période** T est la plus petite durée telle que :

$$s(t + T) = s(t).$$

La **fréquence** du signal :

$$f = \frac{1}{T}$$

est le nombre de répétitions du signal par unité de temps.

La période T se mesure en secondes (symbole s) ; la fréquence f se mesure en hertz (symbole Hz) : $1\text{Hz} = 1\text{s}^{-1}$.

b) Relations entre pulsation, période et fréquence

On a vu plus haut qu’un signal sinusoïdal est périodique de période :

$$T = \frac{2\pi}{\omega}.$$

CHAPITRE 1 – OSCILLATEUR HARMONIQUE

Sa fréquence est donc reliée à la pulsation par les formules :

$$f = \frac{\omega}{2\pi} \quad \text{ou} \quad \omega = 2\pi f.$$

c) Deux autres expressions du signal sinusoïdal

Le signal sinusoïdal $s(t) = A \cos(\omega t + \varphi)$ peut s'écrire aussi :

$$s(t) = A \cos\left(2\pi \frac{t}{T} + \varphi\right) \quad \text{ou} \quad s(t) = A \cos(2\pi f t + \varphi).$$

2.4 Une interprétation géométrique

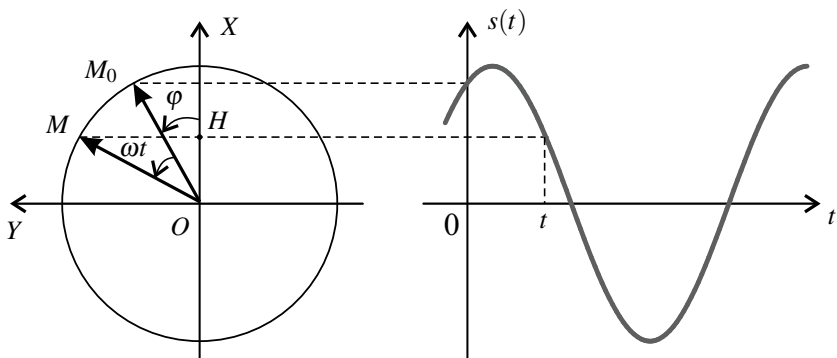


Figure 1.4 – Mouvement circulaire et signal sinusoïdal.

On peut donner du signal sinusoïdal $s(t) = A \cos(\omega t + \varphi)$ une image géométrique. On considère le cercle de rayon A centré à l'origine du repère orthonormé (OXY) (voir figure 1.4) et sur le cercle le point M_0 tel que l'angle entre le vecteur directeur $\vec{u}_X$ de l'axe (OX) et le vecteur $\vec{OM}_0$ vaut φ . Soit un point M se déplaçant sur le cercle avec la vitesse angulaire ω et passant par M_0 à l'instant initial. L'angle entre le vecteur $\vec{OM}$ et l'axe (OX) à l'instant t est $\theta(t) = \omega t + \varphi$ et l'abscisse de ce point est

$$X_M(t) = \overline{OH} = OM \cos \theta(t) = A \cos(\omega t + \varphi) = s(t).$$

Le signal sinusoïdal est ainsi l'abscisse d'un point tournant à la vitesse angulaire ω .

2.5 Représentation de Fresnel (MPSI)

a) Définition du vecteur de Fresnel

On associe au signal $s(t) = A \cos(\omega t + \varphi)$ un vecteur appelé **vecteur de Fresnel**, qui a une norme égale à A et qui fait, à l'instant t , l'angle $\omega t + \varphi$ avec l'axe des abscisses.

Ce vecteur tourne autour de l'origine à la vitesse angulaire ω (voir figure 1.5). Il s'agit du vecteur $\overrightarrow{OM}$ du paragraphe précédent. Dans cet ouvrage le vecteur de Fresnel associé au signal $s(t)$ est noté $\vec{S}$.

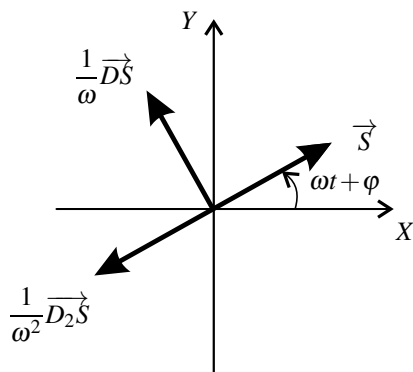


Figure 1.5 – Représentation de Fresnel d'un signal sinusoïdal et ses deux premières dérivées. Les vecteurs $\vec{S}$ et $\vec{D_2 S}$ ont été multipliés respectivement par $\frac{1}{\omega}$ et $\frac{1}{\omega^2}$ pour une question d'homogénéité.

b) Vecteur de Fresnel du signal dérivé

La dérivée par rapport au temps d'un signal sinusoïdal $s(t) = A \cos(\omega t + \varphi)$ est aussi un signal sinusoïdal, en effet :

$$\frac{ds}{dt} = -\omega A \sin(\omega t + \varphi) = \omega A \cos\left(\omega t + \varphi + \frac{\pi}{2}\right).$$

L'amplitude de $\frac{ds}{dt}$ est égale à l'amplitude de $s(t)$ multipliée par ω et sa phase initiale est égale à la phase initiale de $s(t)$ augmentée de $\frac{\pi}{2}$. Ainsi :

Le vecteur de Fresnel associé à $\frac{ds}{dt}$ s'obtient à partir du vecteur de Fresnel associé à $s(t)$ en effectuant les opérations suivantes :

- on tourne le vecteur d'un angle $\frac{\pi}{2}$ dans le sens trigonométrique,
- on multiplie la norme du vecteur par ω (voir figure 1.5).

CHAPITRE 1 – OSCILLATEUR HARMONIQUE

Le vecteur de Fresnel relatif à $\frac{ds}{dt}$ sera noté $\overrightarrow{D\dot{S}}$.

Si on dérive le signal deux fois, on tourne le vecteur d'un angle π et on multiplie sa norme par ω^2 . Ainsi le vecteur de Fresnel $\overrightarrow{D_2\dot{S}}$ associé à $\frac{d^2s}{dt^2}$ est :

$$\overrightarrow{D_2\dot{S}} = -\omega^2 \overrightarrow{S}.$$

Cette relation n'est autre que l'équation différentielle de l'oscillateur harmonique de pulsation propre ω dont le signal sinusoïdal est solution. Elle est visualisée sur la figure 1.5.

2.6 Déphasage

a) Déphasage entre deux signaux sinusoïdaux

On considère deux signaux sinusoïdaux : $s_1(t) = A_1 \cos(\omega_1 t + \varphi_1)$ et $s_2(t) = A_2 \cos(\omega_2 t + \varphi_2)$. On appelle **déphasage** du signal s_2 par rapport au signal s_1 la différence entre leurs phases instantanées :

$$\Delta\varphi(t) = (\omega_2 t + \varphi_2) - (\omega_1 t + \varphi_1) = (\omega_2 - \omega_1)t + \varphi_2 - \varphi_1.$$

Le déphasage est visible dans la représentation de Fresnel : il s'agit de l'angle allant du vecteur $\overrightarrow{S_1}$ au vecteur $\overrightarrow{S_2}$ (voir figure 1.6).

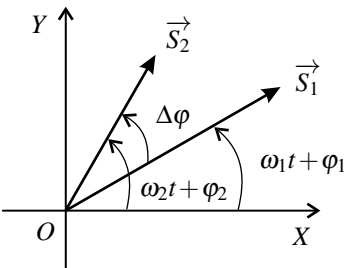


Figure 1.6 – Déphasage entre deux signaux sinusoïdaux.

Quand les deux signaux sinusoïdaux ont la même fréquence (soit $\omega_1 = \omega_2$) leur déphasage est constant dans le temps et égal à la différence de leurs phases initiales :

$$\Delta\varphi = \varphi_2 - \varphi_1.$$

Dans la suite on se place uniquement dans ce cas et on appelle ω la pulsation des deux signaux ($\omega_1 = \omega_2 = \omega$).

b) Valeurs remarquables du déphasage

Les signaux sont dits **en phase** si $\Delta\varphi$ est égal à 0 ou $2n\pi$ avec n entier. Dans ce cas :

$$\cos(\omega t + \varphi_2) = \cos(\omega t + \varphi_1 + 2n\pi) = \cos(\omega t + \varphi_1).$$

Les deux signaux passent par leurs valeurs maximales ou leur valeurs minimales en même temps, s’annulent en même temps. Les vecteurs de Fresnel ont à chaque instant même direction et même sens (voir figure 1.7).

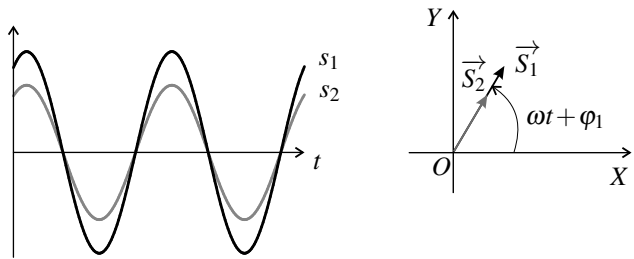


Figure 1.7 – Signaux sinusoidaux de même fréquence en phase.

Les signaux sont dits **en opposition phase** si $\Delta\varphi$ est égal à π ou $(2n + 1)\pi$ avec n entier. Dans ce cas :

$$\cos(\omega t + \varphi_2) = \cos(\omega t + \varphi_1 + (2n + 1)\pi) = \cos(\omega t + \varphi_1 + \pi) = -\cos(\omega t + \varphi_1).$$

À un instant où l’un des signaux passe par sa valeur maximale, l’autre passe par sa valeur minimale. Il s’annulent en même temps mais l’un en croissant et l’autre en décroissant. Les vecteurs de Fresnel ont à chaque instant la même direction mais sont de sens opposés (voir figure 1.8).

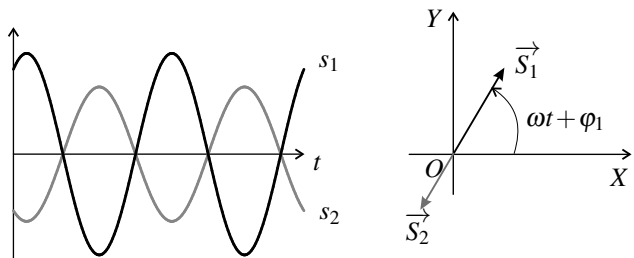


Figure 1.8 – Signaux sinusoidaux de même fréquence en opposition de phase.

Les signaux sont dits **en quadrature de phase** si $\Delta\varphi$ est égal à $\pm\frac{\pi}{2}$ ou $\left(2n \pm \frac{1}{2}\right)\pi$ avec n entier. Alors :

$$\cos(\omega t + \varphi_2) = \cos\left(\omega t + \varphi_1 + \left(2n \pm \frac{1}{2}\right)\pi\right) = \cos\left(\omega t + \varphi_1 \pm \frac{\pi}{2}\right) = \mp \sin(\omega t + \varphi_1).$$

À un instant où l’un des signaux passe par sa valeur maximale, l’autre passe par zéro et réciproquement. Les vecteurs de Fresnel sont orthogonaux à chaque instant. On parle de

CHAPITRE 1 – OSCILLATEUR HARMONIQUE

« quadrature avance » ou « quadrature retard » selon que s_2 est en avance ($\Delta\varphi = +\frac{\pi}{2} + 2n\pi$) ou en retard ($\Delta\varphi = -\frac{\pi}{2} + 2n\pi$) sur s_1 (voir figures 1.9 et 1.10).

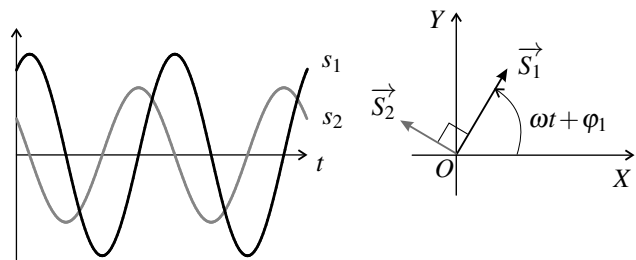


Figure 1.9 – Signaux sinusoïdaux de même fréquence en quadrature de phase. s_2 est en avance sur s_1 .

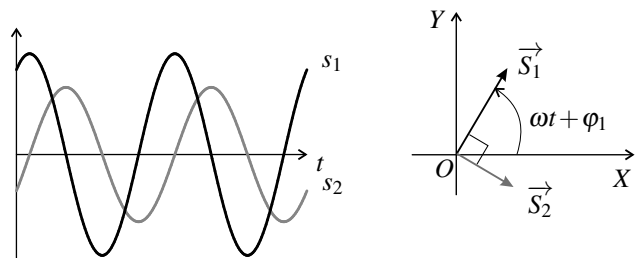


Figure 1.10 – Signaux sinusoïdaux de même fréquence en quadrature de phase. s_2 est en retard sur s_1 .

Remarque

La dérivée d’un signal sinusoïdal $s(t)$ est en quadrature avance sur $s(t)$. La dérivée seconde est en opposition de phase par rapport à $s(t)$. Un signal dont $s(t)$ est la dérivée (soit une primitive de $s(t)$), est en quadrature retard sur $s(t)$.

c) Mesure d’un déphasage

Expérimentalement on peut visualiser un signal en fonction du temps à l’aide d’un oscilloscope ou bien d’une carte d’acquisition reliée à un ordinateur. Pour cela, le signal doit être une tension électrique. Si le signal que l’on veut étudier n’est pas une tension, on utilise un capteur qui fournit une tension proportionnelle à ce signal.

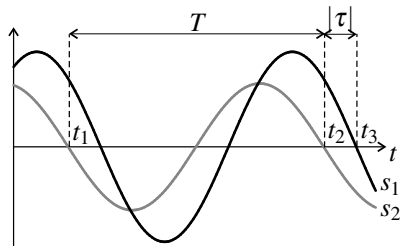


Figure 1.11 – Mesure du déphasage entre deux signaux sinusoïdaux.

La courbe donnant le signal au cours du temps est appelée **forme d’onde**. Dans le cas d’un signal sinusoïdal, on déduit de la forme d’onde :

- l’amplitude A , en mesurant la valeur maximale $s_{\max}$ et la valeur minimale $s_{\min}$, par la formule : $A = \frac{s_{\max} - s_{\min}}{2}$ (attention à ne pas oublier le facteur $\frac{1}{2}$) ;
- la période T , en mesurant les dates t_1 et $t_2 > t_1$ de deux annulations successives du signal avec la même pente (voir figure 1.11), par la formule : $T = t_2 - t_1$.

La phase initiale d’un seul signal sinusoïdal a peu d’intérêt pratique car elle dépend du choix de l’origine des temps. En revanche, le déphasage entre deux signaux sinusoïdaux de même fréquence est souvent une information importante.

Le déphasage est lié au décalage temporel entre les deux signaux. En effet :

$$\cos(\omega t + \varphi_2) = \cos(\omega t + \varphi_1 + \Delta\varphi) = \cos\left(\omega\left(t + \frac{\Delta\varphi}{\omega}\right) + \varphi_1\right) = \cos(\omega(t + \tau) + \varphi_1),$$

formule où apparaît le décalage temporel entre les deux signaux :

$$\tau = \frac{\Delta\varphi}{\omega}.$$

La mesure de τ conduit à la valeur de $\Delta\varphi$. Pour cela on repère : deux dates t_1 et t_2 consécutives en lesquelles la courbe s_2 s’annule avec la même pente, une date t_3 la plus proche possible de t_2 où le signal s_1 s’annule avec une pente du même signe (voir figure 1.11). On en déduit :

- la période du signal : $T = t_2 - t_1$;
- le décalage temporel : $\tau = t_3 - t_2$;
- le déphasage de s_2 par rapport à s_1 : $\Delta\varphi = \omega\tau = 2\pi\frac{t_3 - t_2}{t_2 - t_1}$.

La valeur du déphasage obtenue par cette méthode est comprise entre $-\pi$ et $+\pi$. Elle est positive si $s_2(t)$ est en avance sur $s_1(t)$ (cas de la figure 1.11) et négative si $s_2(t)$ est en retard.



Il faut contrôler le signe du résultat en observant le chronogramme : s_2 est en avance sur s_1 si, entre t_1 et t_2 , il atteint son maximum avant s_1 .

CHAPITRE 1 – OSCILLATEUR HARMONIQUE

d) Utilisation de l'oscilloscope en « mode XY »

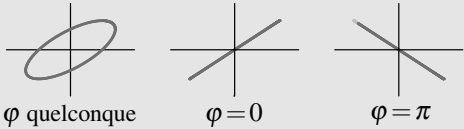
Lorsqu'on utilise une oscilloscope en « mode XY » on observe une courbe formée par les points donc l'abscisse est $X(t) = s_1(t)$ et l'ordonnée $Y(t) = s_2(t)$.

Dans le cas où $s_1(t)$ et $s_2(t)$ sont deux signaux sinusoïdaux, cette courbe est, pour un déphasage quelconque, une ellipse (voir annexe mathématique).

Si ces signaux sont en phase on a : $s_2(t) = \frac{A_2}{A_1}s_1(t)$ soit $Y(t) = \frac{A_2}{A_1}X(t)$. La courbe observée est donc une droite de pente positive.

Si ces signaux sont en opposition phase on a : $s_2(t) = -\frac{A_2}{A_1}s_1(t)$ soit $Y(t) = -\frac{A_2}{A_1}X(t)$. La courbe observée est donc une droite de pente négative.

L'utilisation de l'oscilloscope « en mode XY » permet de repérer facilement des signaux en phase ou en en opposition de phase :



SYNTHÈSE

SAVOIRS

- équation différentielle de l'oscillateur harmonique
- expression de la pulsation de l'oscillateur constitué par une masse accrochée à un ressort
- définitions d'un signal sinusoïdal, de son amplitude, sa pulsation, sa phase initiale
- relations entre la pulsation, la fréquence et la période
- représentation de Fresnel

SAVOIR-FAIRE

- établir l'équation différentielle d'une masse accrochée à un ressort
- résoudre l'équation différentielle d'un oscillateur harmonique avec des conditions initiales données
- trouver l'amplitude et la phase initiale de la solution
- vérifier la conservation de l'énergie mécanique
- reconnaître l'amplitude, la phase initiale, la période, la fréquence, la pulsation d'un signal sinusoïdal donné
- trouver la phase instantanée à un instant où le signal est maximal, minimal, nul
- dessiner une représentation de Fresnel (MPSI)
- déterminer expérimentalement un déphasage

MOTS-CLÉS

- | | | |
|--------------------------|-------------|-----------------------------|
| • oscillateur harmonique | • période | • vecteur de Fresnel |
| • signal sinusoïdal | • fréquence | • conservation de l'énergie |
| • amplitude | • phase | |
| • pulsation | • déphasage | |

CHAPITRE 1 – OSCILLATEUR HARMONIQUE

S'ENTRAÎNER

1.1 Reconnaître un oscillateur harmonique (★)

1. La tension électrique $v(t)$ aux bornes d'un oscillateur à quartz (tel qu'on en trouve dans les montres) vérifie l'équation différentielle : $\frac{d^2v}{dt^2} + Av(t) = 0$ avec $A = 4,239 \cdot 10^{10}$ USI. Quelle est l'unité de A ? Quelle est la fréquence de cet oscillateur ?

2. Un électron de masse $m_e = 9,11 \cdot 10^{-31}$ kg et de charge $q = -1,61 \cdot 10^{-19}$ C est piégé à l'intérieur d'un dispositif tel que son énergie potentielle est $E_p = \frac{1}{2} \frac{qV_0}{d^2} z^2$ où $V_0 = -5,0$ V et $d = 6,0$ mm. On s'intéresse à un mouvement de l'électron selon l'axe (Oz) .

- Exprimer l'énergie mécanique de l'électron en fonction des données et de $z(t)$ de $\frac{dz}{dt}$.
- On suppose que l'énergie mécanique est constante dans le temps. Calculer la fréquence des oscillations de l'électron selon (Oz) dans le piège.

1.2 Mouvement sinusoïdal (★)

On filme avec une webcam le mobile étudié dans le premier paragraphe. A l'aide d'un logiciel permettant de relever image après image la position de G , on trouve que G passe par sa position d'équilibre, avec une vitesse dans le sens de (Ox) , à l'instant $t_1 = 1,44 \pm 0,03$ s et qu'il passe ensuite pour la première fois au milieu entre la position d'équilibre et la position d'élongation maximale à l'instant $t_2 = 2,52 \pm 0,1$ s.

- Calculer la période des oscillations.
- Pourquoi n'a-t-on pas choisi de mesurer l'instant où le mobile passe par la position d'élongation maximale ?

1.3 Vibration d'un diapason (★)

Un diapason vibre à la fréquence du La4 soit $f = 440$ Hz. On mesure sur une photo l'amplitude du mouvement de l'extrémité des branches $A = 0,5$ mm. Quelle est la vitesse maximale de l'extrémité du diapason ? Quelle est l'accélération maximale de ce point ?

1.4 Énergie de l'oscillateur harmonique (★★)

L'énergie mécanique d'un oscillateur harmonique s'écrit : $E_m(t) = \frac{1}{2} m \omega_0^2 x^2 + \frac{1}{2} m \left(\frac{dx}{dt} \right)^2$.

On suppose qu'il n'y a aucun phénomène dissipatif : l'énergie mécanique est donc constante.

1. En utilisant la conservation de l'énergie, retrouver l'équation différentielle de l'oscillateur harmonique.

2. On suppose que $x(t) = A \cos(\omega_0 t + \varphi)$. Exprimer l'énergie cinétique et l'énergie potentielle en fonction de m , ω_0 , A et $\cos(2\omega t + 2\varphi)$. On utilisera les formules : $\cos^2 \alpha = \frac{1}{2}(1 + \cos(2\alpha))$ et $\sin^2 \alpha = \frac{1}{2}(1 - \cos(2\alpha))$. Vérifier que l'énergie mécanique est bien constante.

ÉTAPES D'INSCRIPTION



pour voir comment créer un compte

CLICK TO WATCH

1XBET



STREAMING EN DIRECT SUR MOBILE

Profitez du bonus!



UTILISATEUR VIP: D. VOLKOV - COMPTE PRO

12:36

SOLDE PRO: 147 500 ₺

1XBET PRO

Populaires Sports Esports Casino 1xGames

Tout Football Tennis Basket-ball Hockey sur glace Volleyball Tennis de table

World Win 26 World Flight 26 Brésil - Norvège Mexique Angleterre

Tous les Jeux Crash Crystal Burning Hot Fruit Cocktail Western slot

TOP-EVENTS

WC2026

COMPTE AFFILIÉ CERTIFIÉ - ACCÈS ÉLITE

Populaire EN DIRECT

Sport Tout

Wimbledon, 1/8 de finale

1er set

	Jeu	1	Score
Roman Safiullin (Q)	0	0	0
Novak Djokovic (7)	0	0	0

1X2

Populaire Favoris Coupon Historique Menu

SM-G965F

12:16

1XBET

Se connecter Inscription

Populaires Sports Esports Casino 1xGames

Tout Football Tennis Basket-ball Hockey sur glace Volleyball Tennis de table

World Win 26 World Flight 26 Brésil - Norvège Mexique Angleterre

Tous les Jeux Crash Crystal Burning Hot Fruit Cocktail Western slot

TOP-EVENTS

WC2026

Populaire EN DIRECT

Sport Tout

Corée du Sud. K Ligue 1

Gwangju 1 : 1 Ulsan Hyundai

Populaire Favoris Coupon Historique Menu

SM-G965F

12:17

Inscription

En un clic

Par téléphone

Par e-mail

Par les réseaux sociaux

Vous avez déjà un compte ? Se connecter

SM-G965F

12:31

Inscription

Par e-mail

Nom * Volkov

Prénom * Dmitry

Pays * Russie

Région Novosibirskaya obl.

Ville/Commune Novosibirsk

Date de naissance * 24.07.1996

Devise * Rouble russe (RUB)

E-mail * dmitriy-volkov@gmail.com

Code +7 Téléphone 912 345-25-26

S'inscrire

SM-G965F

12:32

Inscription

Par e-mail

Code +7 Téléphone 912 345-25-26

Mot de passe * Azerty_25_26

Répéter mot de passe * Azerty_25_26

Exigences du mot de passe

Code promo (facultatif) 1SYSCASHBACK

Bonus Bonus pour le sport

Conditions générales

Politique de confidentialité

Recevoir des informations sur les événements par e-mail

Recevoir les notifications concernant les résultats des paris par e-mail

S'inscrire

3. Tracer sur un même graphe les courbes donnant l'énergie cinétique et l'énergie potentielle en fonction du temps. Quelle est la fréquence de variation de ces énergies ?

1.5 Caractéristiques de signaux sinusoïdaux (★)

1. Donner l'amplitude, la période, la fréquence et la phase initiale des signaux suivants :

a. $x(t) = 15 \cos(100\pi t + 0,5)$;

b. $x(t) = 5 \sin(7,854 \cdot 10^6 t)$;

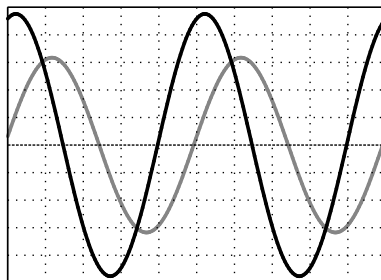
c. $x(t) = 2 \sin(120\pi t - \frac{\pi}{4})$;

d. $x(t) = 15 \cos(2,0 \cdot 10^3 \pi t) - 5 \sin(2,0 \cdot 10^3 \pi t)$ (indication : utiliser la représentation de Fresnel).

2. Quelle est la phase initiale d'un signal sinusoïdal qui vaut la moitié de sa valeur maximale et croît à l'instant $t = \frac{T}{4}$ où T est la période ?

1.6 Détermination d'un déphasage (★)

La figure représente un écran d'oscilloscope avec deux signaux sinusoïdaux de même fréquence $s_1(t)$ (en noir) et $s_2(t)$ (en gris). La ligne en tireté représente le niveau zéro pour les deux signaux. Une division de l'axe des temps correspond à 20 ms.



1. Déterminer la fréquence des signaux.
2. Caculer le déphasage de s_2 par rapport à s_1 .
3. Quelle est la phase de s_1 au point le plus à gauche de l'écran ?

APPROFONDIR

1.7 Vibration d'une molécule (★★)

La fréquence de vibration de la molécule de chlorure d'hydrogène HCl est $f = 8,5 \cdot 10^{13}$ Hz. On donne les masses atomiques molaires : $M_H = 1 \text{ g} \cdot \text{mol}^{-1}$ et $M_{Cl} = 35,5 \text{ g} \cdot \text{mol}^{-1}$, ainsi que le nombre d'Avogadro : $N_A = 6,02 \cdot 10^{23} \text{ mol}^{-1}$.

On modélise la molécule par un atome d'hydrogène mobile relié à un atome de chlore fixe par un « ressort » de raideur k .

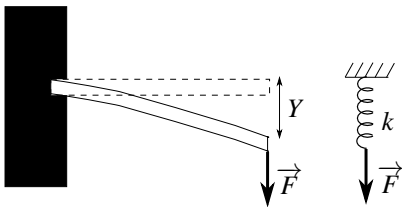
1. Justifier l'hypothèse d'un atome de chlore fixe.
2. Calculer k .
3. On admet que l'énergie mécanique de la molécule est égale à $\frac{1}{2}hf$ où $h = 6,63 \cdot 10^{-34} \text{ J} \cdot \text{s}$ est la constante de Planck. Calculer l'amplitude du mouvement de l'atome d'hydrogène.
4. Calculer sa vitesse maximale.

CHAPITRE 1 – OSCILLATEUR HARMONIQUE

1.8 Un modèle d'élasticité d'une fibre de verre (★★)

Le verre est un matériau très dur. On peut toutefois le déformer légèrement sans le casser : on parle d'élasticité. Récemment, des expériences de biophysique ont été menées pour étudier l'ADN. Le capteur utilisé était simplement une fibre optique en silice amincie à l'extrémité de laquelle on accroche un brin d'ADN. L'expérience consistait à suivre la déformation de flexion de la fibre. La masse volumique du verre est $\rho = 2500 \text{ kg.m}^{-1}$.

La fibre de verre de longueur ℓ et de diamètre d est encastrée horizontalement dans une paroi immobile. Au repos, la fibre est horizontale (on néglige son poids). Quand on applique une force verticale F (on supposera que la force F reste verticale tout au long de l'expérience) à l'extrémité libre de la fibre, celle-ci est déformée. L'extrémité est déplacée verticalement d'une distance Y que l'on appelle la flèche (voir figure).



La flèche Y est donnée par la relation suivante (on notera la présence du facteur numérique 7, sans dimension, qui est en fait une valeur approchée pour plus de simplicité) : $\frac{7\ell^3 F}{Ed^4}$, où E est appelé module d'Young du verre. Pour les applications numériques on prendra pour le module d'Young $E = 7.10^{10} \text{ S.I.}$.

1. Quelle est l'unité S.I. du module d'Young E ?
2. En considérant uniquement la force F , montrer que l'on peut modéliser la fibre de verre par un ressort de longueur à vide nulle et de constante de raideur k dont on donnera l'expression analytique en fonction de E , d et ℓ .
3. Calculer numériquement k pour une fibre de longueur $\ell = 7 \text{ mm}$ et de diamètre $d = 10 \mu\text{m}$.

On a tous fait l'expérience suivante : faire vibrer une règle ou une tige lorsqu'une de ses extrémités est bloquée. On cherche ici à trouver les grandeurs pertinentes qui fixent la fréquence des vibrations. L'extrémité de la tige vaut $Y(t)$ à l'instant t . On admet que lors des vibrations de la fibre, l'énergie cinétique de la fibre de verre est donnée par l'expression $E_c = \rho \ell d^2 \left(\frac{dY}{dt} \right)^2$. Son énergie potentielle élastique lorsque la flèche vaut Y est :

$$E_p = \frac{1}{2} \frac{Ed^4}{7\ell^3} Y^2.$$

4. Écrire l'expression de l'énergie mécanique de la fibre en négligeant l'énergie potentielle de pesanteur.
5. Justifier que l'énergie mécanique se conserve au cours du temps. En déduire l'équation différentielle qui régit les vibrations de la fibre.
6. Quelle est l'expression de la fréquence propre de vibration d'une tige de verre de module d'Young E , de longueur ℓ et de diamètre d ?
7. Calculer numériquement la fréquence des vibrations d'une fibre de verre de longueur 7 mm et de diamètre 0,01 mm.

CORRIGÉS

1.1 Reconnaître un oscillateur harmonique

1. A se mesure en s^{-2} . L'équation différentielle est celle d'un oscillateur harmonique de pulsation $\omega = \sqrt{A}$ donc de fréquence $f = \frac{\omega}{2\pi} = \frac{\sqrt{A}}{2\pi} = 32,77 \text{ kHz}$.

2. a. L'énergie mécanique est la somme de l'énergie cinétique $E_c = \frac{1}{2}m_e \left(\frac{dz}{dt}\right)^2$ et de l'énergie potentielle donnée, soit : $E_m = \frac{1}{2}m_e \left(\frac{dz}{dt}\right)^2 + \frac{1}{2}\frac{|qV_0|}{d^2}z^2$, car $qV_0 = (-q)(-V_0) = |qV_0|$.

b. L'énergie mécanique étant conservée, $\frac{dE_m}{dt} = m_e \frac{d^2z}{dt^2} \frac{dz}{dt} + \frac{|qV_0|}{d^2} \frac{dz}{dt} z(t) = 0$, soit après simplification par $\frac{dz}{dt}$ et division par m : $\frac{d^2z}{dt^2} + \frac{|qV_0|}{m_e d^2} z(t) = 0$.

On reconnaît l'équation d'un oscillateur harmonique de pulsation $\omega = \sqrt{\frac{|qV_0|}{m_e d^2}}$ et donc de fréquence $f = \frac{1}{2\pi} \sqrt{\frac{|qV_0|}{m_e d^2}} = 25.10^6 \text{ Hz} = 25 \text{ MHz}$.

1.2 Mouvement sinusoïdal

1. Le mobile passe par sa position d'équilibre avec une vitesse dans le sens positif à t_1 . D'après la figure 1.3 du cours la phase du signal $x(t)$ à $t = t_1$ est $\varphi = -\frac{\pi}{2}$. La loi horaire du mouvement est donc : $x(t) = A \cos\left(\omega(t - t_1) - \frac{\pi}{2}\right) = A \sin(\omega(t - t_1))$.

2. A l'instant t_2 , $x(t_2) = \frac{A}{2}$, donc $\sin(\omega(t_2 - t_1)) = \frac{1}{2}$, soit $\omega(t_2 - t_1) = \arcsin\left(\frac{1}{2}\right) = \frac{\pi}{3}$.

Ainsi : $\omega = \frac{\pi}{3(t_2 - t_1)}$ d'où $T = \frac{2\pi}{\omega} = 6(t_2 - t_1) = (6,48 \pm 24) \text{ s}$.



Pour encadrer le résultat il faut le calculer avec les valeurs extrêmes des données. Par exemple la valeur maximale de T est : $6 \times (2,52 + 0,1 - (1,44 - 0,03)) = 6,72 \text{ s}$.

3. La mesure de l'instant de passage à la position d'élongation maximale est imprécise car la vitesse du mobile s'annule à cet instant.

1.3 Vibration d'un diapason

La loi horaire du mouvement de l'extrémité du diapason est : $x(t) = A \cos(2\pi f t + \varphi)$, et la vitesse est : $\frac{dx}{dt} = -2\pi f A \sin(2\pi f t + \varphi)$. La valeur maximale de la vitesse est : $v_{\max} = 2\pi f A = 2\pi \times 440 \times 0,5.10^{-3} = 1,4 \text{ m}\cdot\text{s}^{-1}$.

CHAPITRE 1 – OSCILLATEUR HARMONIQUE

L'accélération instantanée est $\frac{d^2x}{dt^2} = -4\pi^2 f^2 A \cos(2\pi ft + \varphi)$. La valeur maximale de l'accélération est : $a_{\max} = 4\pi^2 f^2 A = 3,8 \cdot 10^3 \text{ m} \cdot \text{s}^{-2}$.

1.4 Énergie de l'oscillateur harmonique

1. Il suffit de dériver par rapport au temps la relation : $\frac{1}{2}m\omega_0^2 x(t)^2 + \frac{1}{2}m\left(\frac{dx}{dt}\right)^2 = E_m =$ constante. Il vient : $m\omega_0^2 x(t)\frac{dx}{dt} + m\frac{dx}{dt}\frac{d^2x}{dt^2} = 0$, puis après simplification par $\frac{dx}{dt}$ et division par m on obtient : $\frac{d^2x}{dt^2} + \omega_0^2 x(t) = 0$.

2. $x(t) = A \cos(\omega_0 t + \varphi)$ donc $\frac{dx}{dt} = -\omega_0 A \sin(\omega_0 t + \varphi)$. Il vient donc :

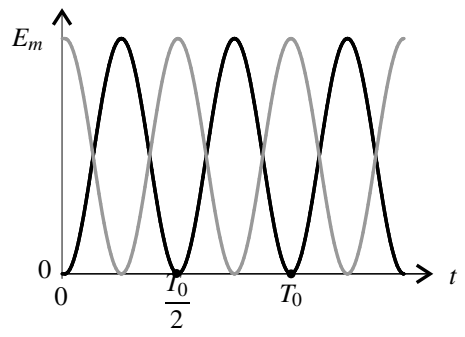
$$E_c(t) = \frac{1}{2}m\omega_0^2 A^2 \sin^2(\omega_0 t + \varphi) = \frac{1}{4}m\omega_0^2 A^2 (1 - \cos(2\omega_0 t + 2\varphi)),$$

et

$$E_p(t) = \frac{1}{2}m\omega_0^2 A^2 \cos^2(\omega_0 t + \varphi) = \frac{1}{4}m\omega_0^2 A^2 (1 + \cos(2\omega_0 t + 2\varphi)).$$

Il apparaît clairement que $E_c(t) + E_p(t) = \frac{1}{2}m\omega_0^2 A^2$.

3. L'évolution dans le temps des énergies cinétique et potentielle est représentée sur la figure ci-contre dans le cas où $v_0 = 0$: $E_c(t)$ en noir et $E_p(t)$ en gris. L'énergie cinétique et l'énergie potentielle oscillent au cours du temps, avec une pulsation $2\omega_0$ donc une période deux fois plus petite que la période $T_0 = \frac{2\pi}{\omega_0}$ de l'oscillateur. Chacune des deux vaut en moyenne $\frac{E_m}{2}$. Il y a un échange continu entre ces deux formes d'énergie.



1.5 Caractéristiques de signaux sinusoïdaux

1. a. Amplitude : $A = 15$, période : $T = 2,00 \cdot 10^{-3} \text{ s}$, fréquence : $f = 50 \text{ Hz}$, phase initiale : $\varphi = 0,5 \text{ rad}$.

b. $A = 5$, $T = 8,000 \cdot 10^{-7} \text{ s}$, $f = 1,250 \text{ MHz}$, $\varphi = -\frac{\pi}{2}$.

c. $A = 2$, $T = 1,67 \cdot 10^{-3} \text{ s}$, $f = 60,0 \text{ Hz}$, $\varphi = -\frac{3\pi}{4} \text{ rad}$.

d. $A = \sqrt{15^2 + 5^2} = 15,8$, $T = 1,0 \cdot 10^{-3} \text{ s}$, $f = 1,0 \cdot 10^3 \text{ Hz}$, $\varphi = \arctan \frac{5}{15} = 0,32 \text{ rad}$.

2. Le signal s'écrit $s(t) = A \cos(\omega t + \varphi)$. On sait que $s(\frac{T}{4}) = \frac{A}{2}$ soit $\cos(\frac{\pi}{2} + \varphi) = \frac{1}{2}$, donc $\frac{\pi}{2} + \varphi = \pm \arccos(\frac{1}{2}) = \pm \frac{\pi}{3}$, soit $\varphi = -\frac{\pi}{6}$ ou $-\frac{5\pi}{6}$.

De plus le signal est croissant à $t = \frac{T}{4}$, donc $-A\omega \sin(\frac{\pi}{2} + \varphi) > 0$ soit $\cos(\varphi) > 0$. La bonne valeur est donc : $\varphi = -\frac{\pi}{6}$.

1.6 Détermination d'un déphasage

- 1. On mesure la période qui vaut 5 carreaux sur le graphe, soit $T = 5 \times 20.10^{-3} = 0,1$ s. La fréquence des signaux est donc : $f = \frac{1}{T} = 10$ Hz.
- 2. s_2 est en retard sur s_1 de la durée correspondant à 1 carreau soit de $\tau = 20$ ms. Le déphasage de s_2 par rapport à s_1 est donc : $\Delta\varphi = -\frac{2\pi}{T} \tau = -\frac{2\pi}{5} = -1,25$ rad.
- 3. A l'instant le plus à gauche de l'écran, s_2 passe par zéro avec une pente positive donc sa phase est $-\frac{\pi}{2}$; on en déduit que la phase de s_1 à cet instant est $-\frac{\pi}{2} + \frac{2\pi}{5} = -\frac{\pi}{10}$.

1.7 Vibrations d'une molécule

- 1. L'atome de chlore étant beaucoup plus lourd que l'atome d'hydrogène on peut considérer que c'est l'atome d'hydrogène qui bouge.
- 2. Ce système est analogue à l'oscillateur harmonique étudié dans le cours. La fréquence d'oscillation est $f = \frac{1}{2\pi} \sqrt{\frac{k}{m_H}}$ où $m_H = \frac{M_H}{\mathcal{N}_A}$ est la masse d'un atome d'hydrogène. Il vient donc : $k = 4\pi^2 f^2 \frac{M_H}{\mathcal{N}_A} = 4\pi^2 \times (8,5.10^{13})^2 \times \frac{1.10^{-3}}{6,02.10^{23}} = 4,7.10^2 \text{ N}\cdot\text{m}^{-1}$.
- 3. L'énergie mécanique est : $E_m = \frac{1}{2}hf = \frac{1}{2}m(2\pi f)^2 m_H A^2$ où A est l'amplitude du mouvement de l'atome d'hydrogène. Il vient :

$$A = \sqrt{\frac{h\mathcal{N}_A}{4\pi^2 f M_H}} = \sqrt{\frac{6,63.10^{-34} \times 6,02.10^{23}}{4\pi^2 \times 8,5.10^{13} \times 1.10^{-3}}} = 1,1.10^{-11} \text{ m}.$$

- 4. La vitesse maximale de l'atome d'hydrogène est :

$$v_{\max} = 2\pi f A = 2\pi \times 8,5.10^{13} \times 1,1.10^{-11} = 5,8.10^3 \text{ m}\cdot\text{s}^{-1}.$$

1.8 Un modèle d'élasticité d'une fibre de verre

- 1. La quantité $\frac{7\ell^3 F}{Ed^4}$ est égale à Y donc homogène à une longueur. On effectue une analyse dimensionnelle sachant qu'une force est homogène à $\frac{ML}{T^2}$. Ainsi :

$$\frac{7\ell^3 F}{Ed^4} \sim L \Rightarrow E \sim \frac{L^3}{L^4} \frac{ML}{T^2} \frac{1}{L} \sim \frac{M}{LT^2}.$$

E est homogène à une force surfacique (ou à une pression).

CHAPITRE 1 – OSCILLATEUR HARMONIQUE

2. On considère l'extrémité de la fibre comme un point matériel. Les deux forces s'exerçant sur ce point sont F et la tension T du ressort équivalent cherché qui est vers le haut.

L'équilibre de l'extrémité donne la relation entre les forces : $T = -F$ or $Y = \frac{7\ell^3 F}{Ed^4}$ soit

$T = -\frac{Ed^4}{7\ell^3}Y = -kY$. La fibre est équivalente à un ressort de longueur à vide nulle et constante de raideur $k = \frac{Ed^4}{7\ell^3}$.

3. L'application numérique donne $k = 2,9 \cdot 10^{-4} \text{ N.m}^{-1}$.

4. L'énergie mécanique est la somme de l'énergie potentielle donnée et de l'énergie cinétique soit :

$$E_m = \frac{Ed^4}{14\ell^3}Y^2 + \rho\ell d^2 \left(\frac{dY}{dt} \right)^2.$$

5. Il n'y pas d'autres forces à prendre en compte que la tension du ressort équivalent, donc si on applique le théorème de l'énergie cinétique, il vient : $dE_c = \delta W$ or $\delta W = -dE_p$, d'où $dE_c = -dE_p$ soit $dE_m = dE_c + dE_p = 0$. L'énergie mécanique est donc constante.

Pour obtenir l'équation du mouvement, on dérive E_m par rapport au temps :

$$\frac{dE_m}{dt} = 0 \Rightarrow \dot{Y} \left(\frac{Ed^4}{7\ell^3}Y + 2\rho\ell d^2\dot{Y} \right) = 0.$$

On simplifie par $\dot{Y}$ qui n'est pas identiquement pour obtenir l'équation :

$$\ddot{Y} + \frac{Ed^2}{14\rho\ell^4}Y = 0$$

qui est l'équation d'un oscillateur harmonique.

6. La pulsation propre de l'oscillateur précédent est $\omega_0 = \sqrt{\frac{Ed^2}{14\rho\ell^4}}$ et la fréquence propre $f_0 = \omega_0/2\pi$.

7. L'application numérique donne : $\omega_0 = 288 \text{ rad.s}^{-1}$ et $f_0 = 45,9 \text{ Hz}$.

Propagation d'un signal

2

Les signaux envisagés dans ce chapitre sont des **ondes**. Ces signaux ne sont pas sinusoïdaux dans la majorité des cas, mais ils s'interprètent comme la superposition de signaux sinusoïdaux et peuvent être représentés par un **spectre**.

Une onde est modélisée par une fonction du temps et d'une variable spatiale. La **propagation** de l'onde se traduit par une forme mathématique particulière de cette fonction. On étudiera le cas des **ondes sinusoïdales** qui sont périodiques dans le temps et dans l'espace.

1 Signaux physiques, spectre

1.1 Ondes et signaux physiques

On appelle **onde** un phénomène physique dans lequel une perturbation locale se déplace dans l'espace sans qu'il y ait de déplacement de matière en moyenne. Toute grandeur physique, nulle dans l'état de repos et apparaissant avec la perturbation, est appelée **signal physique** transporté par l'onde.

On peut citer quelques exemples d'ondes :

1. Tout le monde connaît les « ronds dans l'eau » générés à la surface de l'eau par une action locale comme l'impact d'un projectile : le signal physique est le déplacement vertical de la surface de l'eau. Les particules d'eau ont un mouvement vertical « sur place » et elles retrouvent leur position initiale après le passage de l'onde de sorte que le phénomène n'entraîne aucun déplacement global de l'eau. Le mot onde vient du latin *unda* qui signifie eau.
2. Il existe des **ondes élastiques** se propageant dans les milieux solides. Le signal transporté par ces ondes est une déformation locale et réversible de la matière. L'onde est qualifiée de :
 - **transversale** lorsque le mouvement local de la matière est perpendiculaire à la direction dans laquelle l'onde se propage ;
 - **longitudinale** lorsque le mouvement local de la matière est parallèle à la direction dans laquelle l'onde se propage.

CHAPITRE 2 – PROPAGATION D’UN SIGNAL

On peut propager des ondes transversales le long d’une corde, des ondes longitudinales ou transversales le long d’un ressort à boudin. Parmi les **ondes sismiques** se propageant dans la croûte terrestre, il y a des ondes longitudinales (ondes *P* qui, lors d’un tremblement de terre, arrivent en premier) et des ondes transversales (ondes *S* qui arrivent en second).

- 3. Les **ondes acoustiques** se propagent dans tous les milieux matériels, solides ou fluides. Dans un solide il s’agit d’ondes élastiques longitudinales. Dans un fluide, la perturbation est une modification très faible de la pression due à un mouvement longitudinal imperceptible des couches de fluide. Les signaux transportés par l’onde acoustique sont la variation de pression par rapport à l’état de repos, appelée surpression acoustique, et la vitesse de vibration. Les microphones sont sensibles à l’un ou l’autre de ces signaux.
- 4. À la différence des exemples précédents, les **ondes électromagnétiques** n’ont pas besoin de milieu matériel pour se propager : elles traversent par exemple l’espace vide entre une galaxie lointaine et la Terre. Les deux grandeurs physiques associées à ces ondes sont un champ électrique et un champ magnétique. La lumière est une onde électromagnétique.

Ces ondes peuvent être guidées, par exemple le long d’un câble de transmission constitué par deux conducteurs. Dans ce cas, on peut associer à l’onde deux signaux électriques : la tension entre les deux conducteurs et l’intensité passant à travers leurs sections (dans des sens opposés). On peut parler alors d’« onde de courant » le long du câble.

- 5. La théorie de la relativité générale prédit l’existence d’**ondes gravitationnelles**. Ces ondes se déplacent dans l’espace sans support matériel. Le signal est une déformation de l’espace induisant des variations des longueurs. Ces ondes n’ont encore jamais été détectées directement.

Type d’onde	Milieu de propagation	Signaux physiques
Ondes élastiques	solide	déplacement transversal ou longitudinal
Ondes sonores	fluide	surpression acoustique, vitesse (longitudinale)
Ondes électromagnétiques	vide	champ électrique, champ magnétique
Ondes de courant	câble de transmission	tension électrique, intensité électrique
Ondes gravitationnelles	vide	déformation de l’espace

Tableau 2.1 – Quelques ondes et les signaux associés.

1.2 Notion de spectre

Le signal d’une onde n’est pas, en général, sinusoïdal. Cependant, une théorie mathématique due à Joseph Fourier, mathématicien et physicien du début du *XIX*^{ème} siècle, montre que tout signal réalisable en pratique peut être décomposé en **somme de signaux sinusoïdaux**.

a) Analyse spectrale

L'opération qui consiste à déterminer les signaux sinusoïdaux composant un signal donné est appelée **analyse spectrale**. Le résultat de l'analyse spectrale est :

- la liste des fréquences f_i des composantes sinusoïdales contenues dans le signal,
- l'amplitude A_i de chaque composante sinusoïdale de fréquence f_i ,
- la phase initiale φ_i de chaque composante sinusoïdale de fréquence f_i .

Elle permet de savoir que le signal s'écrit :

$$s(t) = \sum_i A_i \cos(2\pi f_i t + \varphi_i).$$

(2.1)

Chaque signal sinusoïdal $s_i(t) = A_i \cos(2\pi f_i t + \varphi_i)$ est une **composante sinusoïdale** du signal $s(t)$. On dit que le signal $s(t)$ « contient les fréquences f_i ». Le **spectre du signal** est l'ensemble des fréquences contenues dans le signal.

b) Représentation graphique du spectre

La représentation graphique des A_i en fonction des f_i constitue le **spectre d'amplitude**. Elle permet de visualiser le contenu fréquentiel du signal.

On préfère parfois représenter les carrés des amplitudes A_i^2 en fonction des fréquences f_i pour visualiser la contribution de chaque composante à l'énergie du signal, c'est le **spectre d'énergie**.

La représentation des phases initiales φ_i en fonction des f_i est le **spectre de phase**. Ce spectre est plus difficilement interprétable car les phases initiales, à la différence des amplitudes, dépendent du choix de l'origine des temps qui est arbitraire.

Exemple

Le signal $s(t) = 3 \cos(30\pi t + \frac{\pi}{3}) + 4 \cos(100\pi t)$ contient les fréquences $f_1 = 15$ Hz et $f_2 = 50$ Hz. Son spectre, d'amplitude et de phase, est représenté sur la figure 2.1.

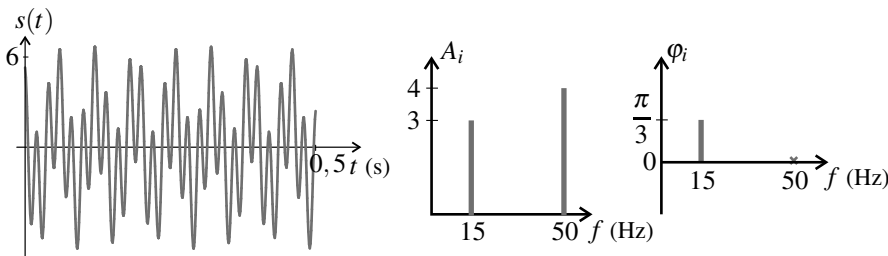


Figure 2.1 – Signal $s_1(t)$ et son spectre : amplitudes et phases initiales.

On remarque que le chronogramme (courbe $s(t)$ en fonction de t) est difficile à interpréter. Le spectre donne une information bien plus lisible.

CHAPITRE 2 – PROPAGATION D’UN SIGNAL

c) Synthèse de Fourier

Toute l’information que l’on peut avoir sur un signal est contenue dans les spectres d’amplitude (ou énergie) et de phase. L’opération consistant à reconstituer un signal $s(t)$ dont on connaît le spectre (fréquences f_i , amplitudes A_i et phases initiales φ_i) par la formule 2.1 est appelée **synthèse de Fourier**.

1.3 Cas d’un signal périodique de forme quelconque

Dans ce paragraphe on s’intéresse à un signal périodique dont on note T_S la période et $f_S = \frac{1}{T_S}$ la fréquence. Il s’agit d’un signal de forme quelconque, *a priori* non sinusoïdal.

a) Analyse spectrale d’un signal périodique

Un théorème mathématique important découvert au début du $XIX^{\text{ème}}$ siècle par le mathématicien et physicien français Joseph Fourier indique que :

Tout signal périodique de fréquence f_S et de forme quelconque peut se reconstituer par la superposition de signaux sinusoïdaux de fréquences multiples de f_S : $0, f_S, 2f_S, \dots, nf_S, \dots$. Il peut donc s’écrire sous la forme :

$$s(t) = A_0 + \sum_{n=1}^{\infty} A_n \cos(2\pi n f_S t + \varphi_n), \tag{2.2}$$

où les A_n sont des constantes positives et les φ_n des constantes.

La somme infinie de la formule (2.2) est le **développement en série de Fourier** du signal $s(t)$. Les formules permettant de trouver les amplitudes A_n et les phases initiales φ_n des composantes sinusoïdales du signal ne sont pas au programme.

L’usage a consacré des noms pour les différents termes qui apparaissent dans la série de Fourier :

- A_0 est appelée **composante continue**, c’est la moyenne du signal qui est la plupart du temps nulle dans le cas d’une onde ;
- la composante sinusoïdale $A_1 \cos(2\pi f_S t + \varphi_1)$ qui a la même fréquence f_S que le signal $s(t)$ est appelée **fondamental** ;
- la composante sinusoïdale $A_n \cos(2\pi n f_S t + \varphi_n)$, qui a une fréquence n fois plus élevée que la fréquence du signal, avec $n \geq 2$, est appelée **harmonique de rang n** .

Ainsi le spectre d’un signal périodique a l’allure représentée sur la figure 2.2.

Remarque

Le spectre d’un signal sinusoïdal ne comporte qu’un fondamental.

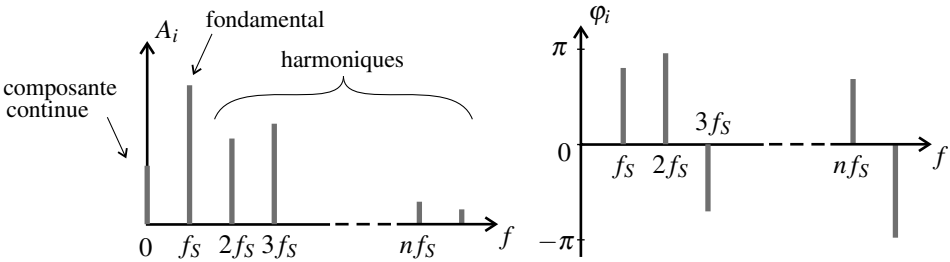


Figure 2.2 – Spectre d'un signal périodique $s(t)$: amplitude et phase initiale.

b) Exemple de synthèse de Fourier d'un signal périodique (MPSI)

On considère le signal représenté sur la figure 2.3, de période T_S . En définissant un temps adimensionnel $t^* = \frac{t}{T_S}$, on peut écrire ce signal : $s(t) = 0,625(t^* - 0,4)$ si $0 \leq t^* \leq 0,8$ et $s(t) = 2,5(t^* - 0,9)$ si $0,8 \leq t^* \leq 1$.

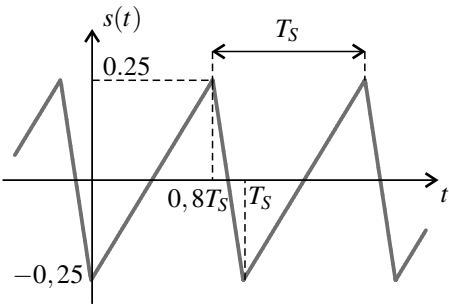


Figure 2.3 – Exemple de signal périodique non sinusoïdal de période T_S .

Les amplitudes et phases initiales des composantes sinusoïdales de ce signal sont facilement calculables et on trouve :

$$A_n = \frac{1}{\pi n^2} |\sin(0,8\pi n)| \quad \text{et} \quad \varphi_n = (-0,5 - 0,2r)\pi,$$

où r est le reste de la division de n par 5. Ce spectre est représenté sur la figure 2.4.

On peut tester la synthèse de Fourier en calculant et représentant les sommes partielles :

$$S_n(t) = \sum_{i=1}^n A_i \cos(2\pi t^* - \varphi_i).$$

Les résultats sont donnés sur la figure 2.5. On voit le signal se construire petit à petit à partir de son fondamental et de ses harmoniques de rang de plus en plus élevé. La somme S_3 limitée

CHAPITRE 2 – PROPAGATION D’UN SIGNAL

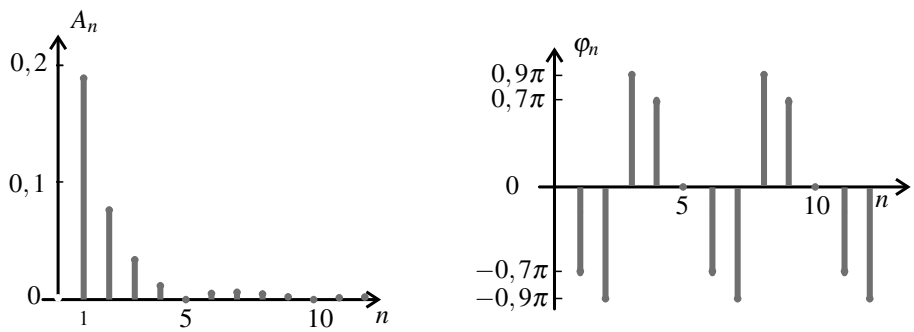


Figure 2.4 – Spectre du signal $s(t)$ défini dans le texte : amplitudes et phases initiales (ramenées entre $-\pi$ et π).

aux harmoniques d’ordre inférieur ou égal à 3 a déjà l’allure globale du signal ; la somme S_{10} est pratiquement égale au signal et on ne peut pas distinguer à l’œil la somme S_{100} du signal. Ce procédé peut s’appliquer, après détermination du spectre, à un signal physique réel.

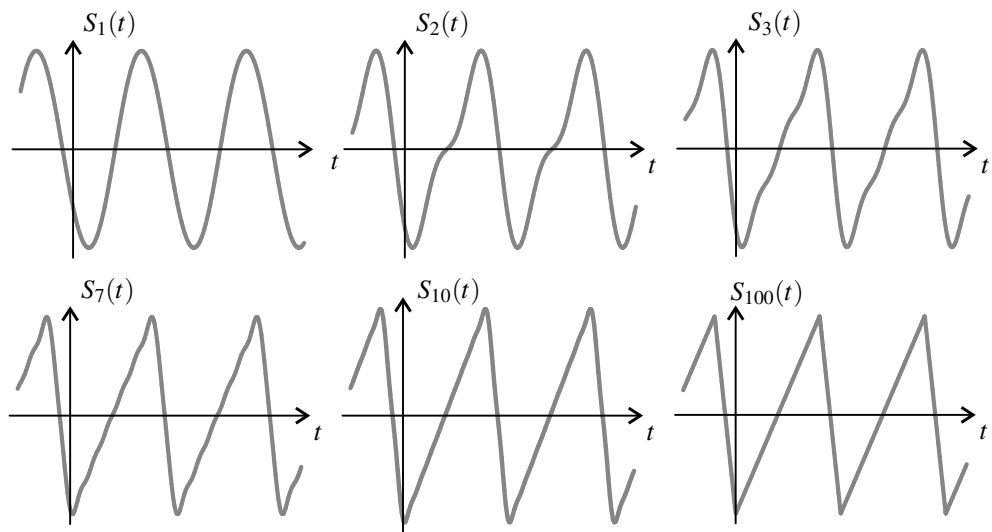


Figure 2.5 – Sommes partielles de Fourier $S_n(t)$ pour $n = 1, 2, 3, 7, 10$ et 100 .

1.4 Cas d'un signal non périodique

a) Spectre continu

La théorie de la **transformée de Fourier** montre qu'il est possible de reconstituer tout signal physique, même non périodique, par superposition de signaux sinusoïdaux. La différence avec le cas des signaux périodiques est que les fréquences f (et pulsation $\omega = 2\pi f$) des composante sinusoïdales prennent *a priori* toutes les valeurs de 0 à l'infini.

La décomposition d'un signal non périodique s'écrit sous la forme :

$$s(t) = \int_0^\infty A(\omega) \cos(\omega t + \varphi(\omega)) d\omega.$$

$A(\omega)$ est la densité spectrale d'amplitude qui a la dimension du signal $s(t)$ divisée par l'unité de pulsation, c'est-à-dire multipliée par des secondes.

Le spectre du signal est l'ensemble des fréquences pour lesquelles la densité spectrale $A(\omega)$ est non nulle. Il comporte un continuum de fréquences entre une fréquence minimale et une fréquence maximale, d'où le nom de **spectre continu**.

L'information complète concernant le signal est donnée par deux courbes :

- le spectre d'amplitude : courbe de $A(\omega)$ en fonction de ω ;
- le spectre de phase initiale : courbe $\varphi(\omega)$ en fonction de ω .

b) Étude d'un exemple

Dans le cas du signal $s(t) = \exp(-at^2) \cos(2\pi t)$ un calcul en dehors du programme montre que $A(\omega) = 2\sqrt{\frac{\pi}{a}} \cosh\left(\frac{\pi\omega}{a}\right) \exp\left(-\frac{\pi^2 + \omega^2}{4a}\right)$ et $\varphi(\omega) = 0$. La figure 2.6 montre ce signal et son spectre d'amplitude pour deux valeurs différentes de a .

Dans les deux cas la courbe d'amplitude est maximale pour $\omega_0 = 2\pi \text{ rad}\cdot\text{s}^{-1}$ qui est la pulsation du cosinus apparaissant dans l'expression du signal. Lorsque a diminue, le signal occupe une plage de temps plus grande et le spectre se resserre autour de ω_0 . Si on faisait tendre a vers 0, le signal deviendrait un signal purement sinusoïdal et son spectre un pic de largeur nulle en ω_0 . Ceci est général : plus un signal est court, plus son spectre est large.

Des exemples de spectres continus réels sont donnés dans la suite de ce chapitre (voir figures 2.8 et 2.9), ainsi que dans le chapitre 4.

1.5 Analyse harmonique expérimentale

La plupart des oscilloscopes numériques et des logiciels de traitement des données expérimentales sont pourvus d'un module d'analyse de Fourier permettant d'obtenir le spectre d'un signal physique.

CHAPITRE 2 – PROPAGATION D’UN SIGNAL

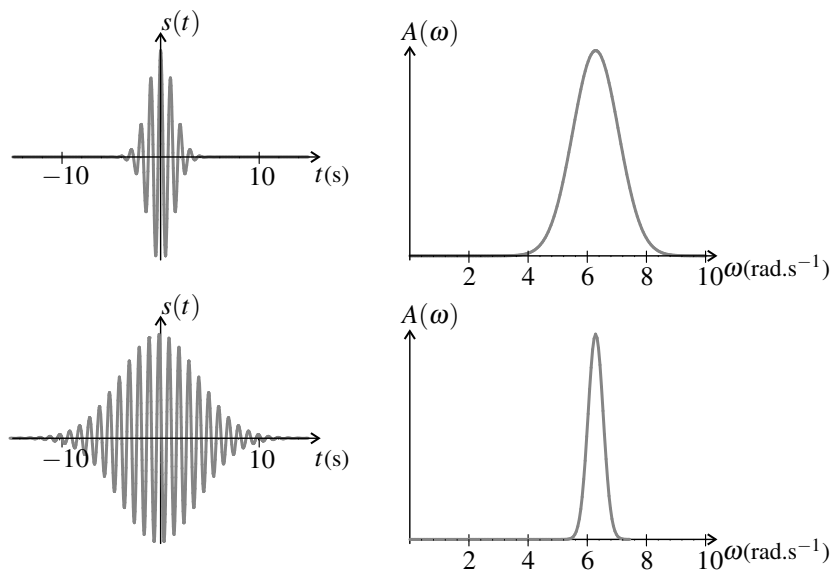


Figure 2.6 – Signal $s(t) = \exp(-at^2) \cos(2\pi t)$ et son spectre d'amplitude pour deux valeurs de a : en haut $a = 0,3$, en bas $a = 0,033$.

a) Obtention du spectre d’un signal périodique

Dans le cas d’un signal périodique, pour obtenir un spectre exact, il est nécessaire de faire les calculs sur un **échantillon de signal correspondant à un nombre entier de périodes**.

Les logiciels sont munis d’un système de « sélection de périodes » pouvant fonctionner soit de manière manuelle soit de manière automatique (si le signal est propre).

Sur certains oscilloscopes, le calcul du spectre porte sur la portion de signal visible à l’écran. Si l’appareil est réglé de manière à ce qu’on voie la forme du signal sur un petit nombre de périodes le spectre a une mauvaise résolution. Pour avoir une bonne résolution en fréquence il faut augmenter la base de temps et voir des oscillations très serrées.

b) Obtention d’un spectre continu

Dans le cas d’un signal non périodique, on cherche à déterminer un spectre continu. Pour avoir la meilleure résolution en fréquence, il faut utiliser un échantillon aussi long que possible de signal, tout en préservant un pas d’échantillonnage δt suffisamment petit pour avoir une large gamme de fréquences.

c) Utilisation d’une fenêtre

Le calcul de spectre génère des erreurs lorsque le signal n’a pas la même valeur à l’instant t_{initial} et à l’instant t_{final} . Pour éviter cela on peut le multiplier par une fonction appelée *fenêtre* nulle (ou presque nulle) aux deux bornes de l’intervalle de temps.

Les logiciels proposent de nombreuses fenêtres parmi lesquelles :

- la fenêtre rectangulaire qui vaut 1 sur l'intervalle d'échantillonnage (ce qui revient à ne pas utiliser de fenêtre) ;
- la fenêtre de *Hamming* qui donne une bonne résolution sur la position des pics mais des valeurs des amplitudes peu fiables ;
- la fenêtre *Flattop* qui préserve la valeur des amplitudes des pics mais ne permet pas une lecture précise de leurs positions sur l'axe des fréquences.

1.6 Exemple : analyse de signaux sonores

Dans ce paragraphe on s'intéresse à des signaux sonores. Le mot « harmonique » désignant une composante sinusoïdale d'un signal périodique provient du vocabulaire musical.

a) Ordre de grandeur des fréquences sonores

Le **domaine audible**, intervalle des fréquences f_{son} perçues par l'oreille humaine, s'étend de 20 Hz à 20 kHz.

C'est l'intervalle des fréquences perçues par un individu « moyen ». En fait, l'intervalle des fréquences perçues est très variable suivant l'individu et il se rétrécit avec l'âge à partir de 20 ans. Une exposition fréquence à des sons très intenses tend à accélérer ce processus.

b) Exemples de spectres de sons

Expérience

On peut facilement analyser un son en utilisant un micro de bonne qualité et un ordinateur muni d'une carte son et d'un logiciel de traitement du son.

La figure 2.7 montre le signal (tension électrique) fourni par une guitare électrique et destiné à être envoyé sur un amplificateur. Le signal est périodique et le spectre a été calculé sur 20 périodes. On observe le fondamental à $f_s = 178 \text{ Hz}$ (il s'agit d'un Fa3), une harmonique 2 très prononcée et des harmoniques 3, 6 et 7.

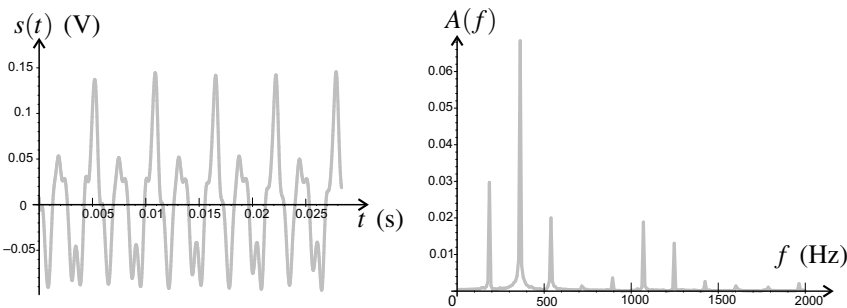


Figure 2.7 – Le son d'une guitare électrique jouant un Fa3 et son spectre en amplitude obtenu avec un calcul sur 20 périodes.

CHAPITRE 2 – PROPAGATION D’UN SIGNAL

La figure 2.8 montre le signal délivré par un micro captant le son d’un battement unique de tambour. Le spectre été calculé sur la totalité du son qui dure un peu plus d’un dixième de seconde. On observe un spectre continu. Le maximum d’amplitude est à 176 Hz. On repère aussi des pics à 546 Hz et 899 Hz. S’agit-il d’harmoniques ? On a $\frac{546}{176} = 3,1$ et $\frac{899}{176} = 5,1$. Ces fréquences, proches de multiples entiers de la fréquence la plus importante dans le signal sont appelés partiels.

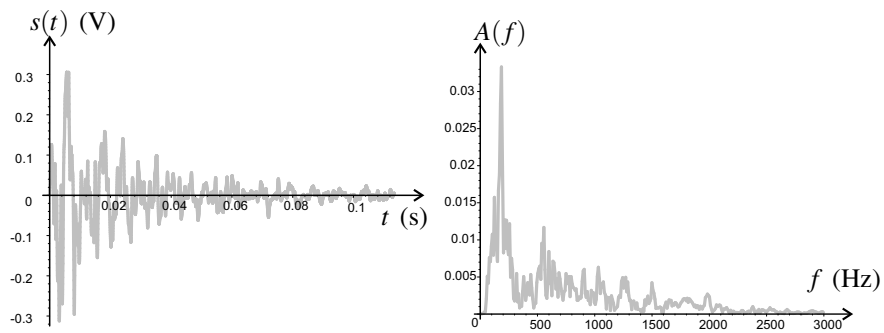


Figure 2.8 – Le son d’un tambour et son spectre en amplitude.

La figure 2.9 représente le signal (tension électrique) fournit par un micro captant le son d’une feuille de papier que l’on froisse. Le spectre a été calculé sur un échantillon durant 1,5 s. C’est un spectre continu qui contient toutes les fréquences entre 400 Hz environ et 7 kHz.

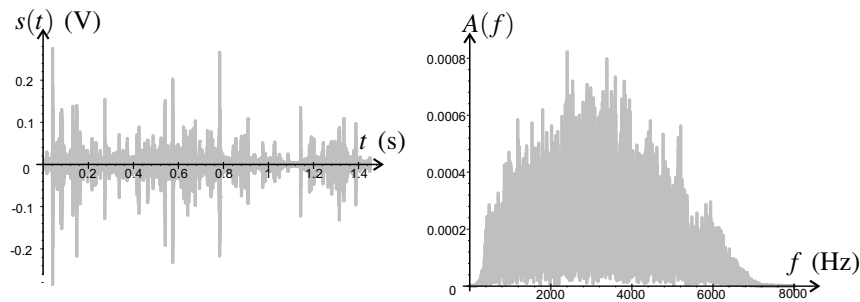


Figure 2.9 – Le son d’une feuille de papier qu’on froisse et son spectre en amplitude.

2 Phénomène de propagation

La deuxième partie de ce chapitre est consacré à la propagation des ondes. Après avoir introduit le modèle de l'onde progressive, on étudiera le modèle de l'onde progressive sinusoïdale.

2.1 Observations expérimentales

a) Onde à la surface de l'eau

Lorsqu'on jette une pierre dans un étang, on observe à la surface de l'eau des rides circulaires, les « ronds dans l'eau ». La perturbation engendrée par la pierre se transmet de proche en proche à des points de plus en plus éloignés du lieu de l'impact. Si l'on filme le phénomène, en se plaçant à la verticale du point d'impact, on peut mesurer, à différents instants $t_1, t_2, \dots, t_n$, le rayon d'un « rond » donné, les valeurs successives $R_1, R_2, \dots, R_n$ sont bien représentées par une loi de la forme : $R_i = R_0 + vt_i$, où v est une constante homogène à une vitesse. v est la vitesse de propagation de l'ébranlement à la surface de l'eau qui se déplace entre les instants t_i et t_{i+1} sur une distance $R_{i+1} - R_i = v(t_{i+1} - t_i)$.

b) Onde sonore

Expérience

L'expérience suivante met en évidence le phénomène de propagation du son.

On dispose deux micros à distance d l'un de l'autre de manière à ce qu'une onde sonore puisse parvenir à chacun d'entre eux sans que l'autre ne lui fasse obstacle. Les signaux délivrés par les micros sont envoyés sur les deux voies d'un oscilloscope numérique réglé en « monocoup » ou une carte d'acquisition pour microordinateur.

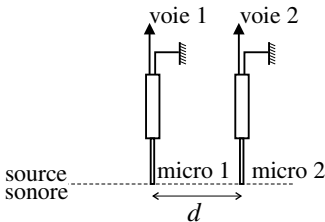


Figure 2.10 – Mesure de la vitesse de propagation du son dans l'air.

La source sonore doit produire un son sec et bref. Le document ci-dessous correspond à un claquement de doigts.

Le déclenchement de l'acquisition est provoqué par le signal du micro 1 utilisé comme « source de déclenchement » : elle commence à l'instant où ce signal atteint une valeur que l'on choisit. Cette valeur doit être ni trop basse (pour qu'un signal parasite ne déclenche pas l'acquisition), ni trop élevée (pour que le signal soit suffisamment fort pour la déclencher).

CHAPITRE 2 – PROPAGATION D’UN SIGNAL

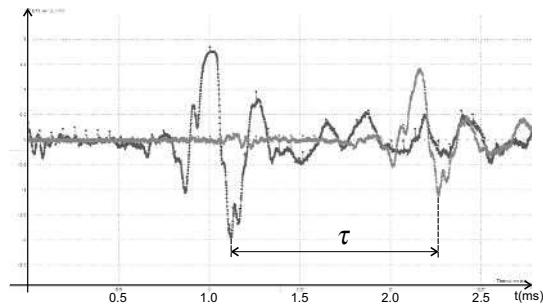


Figure 2.11 – Signaux détectés par les deux micros : le micro 1 reçoit en premier le signal, le micro 2 le reçoit avec un retard τ et une certaine atténuation.

L’enregistrement de la figure 2.11 a été fait avec une distance $d = 40 \text{ cm} \pm 1 \text{ mm}$ entre les micros. On observe que le signal donné par le micro 2 reproduit celui donné par le micro 1, avec une atténuation et un retard $\tau = 1,152 \text{ ms} \pm 0,004 \text{ ms}$.

On interprète cette expérience en considérant que le signal sonore se déplace à la vitesse c_{son} . Il atteint donc le micro 2 après le micro 1, avec un retard : $\tau = \frac{d}{c_{\text{son}}}$.

On en déduit :

$$c_{\text{son}} = \frac{d}{\tau} = \frac{0,40}{1,125 \cdot 10^{-3}} = 347 \pm 2 \text{ m} \cdot \text{s}^{-1}.$$

2.2 Onde progressive

a) Vitesse de propagation ou célérité

Les deux situations précédentes illustrent le phénomène de **propagation**. Dans les deux, un signal physique (ébranlement de l’eau, surpression de l’onde sonore) créé en un point de l’espace, se transmet de proche en proche dans la matière (l’eau, l’air) et peut être ainsi observé à distance de l’endroit où il a été produit, et ceci sans qu’il y ait de déplacement de matière entre le point où le signal est produit et celui où il est mesuré.

La vitesse de déplacement du signal est appelée **vitesse de propagation** ou encore **célérité**. Dans la suite on la notera c .

Pour modéliser le phénomène de propagation, on introduit dans ce paragraphe l’**onde progressive à une dimension**, qui se propage sans atténuation ni déformation à la vitesse constante c dans la direction d’un axe (Ox) . Cette onde est représentée mathématiquement par une fonction $s(x,t)$ de deux variables : la coordonnée x selon l’axe (Ox) et le temps t . $s(x,t)$ est la valeur du signal, mesurée à l’abscisse x , à l’instant t .

On considérera le cas où l’onde progresse dans le sens positif de (Ox) et le cas où elle progresse dans le sens négatif de (Ox) .

b) Première expression de l'onde progressive

On considère une onde progressive, se propageant avec la célérité c dans la direction de l'axe (Ox) et dans le sens positif de cet axe, c'est-à-dire vers les x croissants. La figure 2.12 représente le signal mesuré à une abscisse x_0 , en fonction du temps, ainsi que le signal mesuré à une abscisse $x_1 > x_0$.

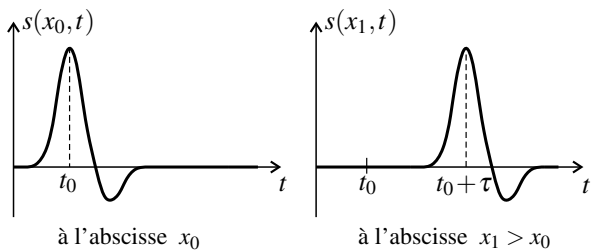


Figure 2.12 – Onde se propageant sans atténuation ni déformation dans le sens positif de (Ox) , en deux points différents.

Les valeurs observées en x_0 au cours du temps sont observées aussi en x_1 , mais avec un retard τ . Ceci s'écrit :

$$s(x_1, t) = s(x_0, t - \tau).$$

La durée τ est celle qu'il faut à l'onde pour se propager de x_0 à x_1 . La vitesse de propagation étant c on a :

$$\tau = \frac{x_1 - x_0}{c}$$

Il vient donc :

$$s(x_1, t) = s\left(x_0, t - \frac{x_1 - x_0}{c}\right). \tag{2.3}$$

On peut remarquer que la formule (2.3) est valable aussi quand $x_1 < x_0$. Elle s'écrit en effet : $s(x_1, t) = s\left(x_0, t + \frac{x_0 - x_1}{c}\right)$ et signifie alors que l'on trouve en x_1 le même signal qu'en x_0 avec une *avance* $\frac{x_0 - x_1}{c}$.

Si on écrit la formule (2.3) en posant $x_0 = 0$ et $x_1 = x$, valeur quelconque, on trouve :

$$s(x, t) = s\left(0, t - \frac{x}{c}\right).$$

Le membre de droite de cette équation est simplement une fonction d'une seule variable, $t - \frac{x}{c}$. Pour simplifier l'écriture on le note $f\left(t - \frac{x}{c}\right)$.

CHAPITRE 2 – PROPAGATION D’UN SIGNAL

Une **onde progressive** se propageant à la vitesse c dans la direction de l’axe (Ox) , *dans le sens positif* de cet axe, sans atténuation ni déformation, est de la forme mathématique suivante :

$$s(x,t) = f\left(t - \frac{x}{c}\right),$$

où f est une fonction quelconque dont l’argument a la dimension d’un temps.

On peut s’intéresser aussi à une onde se propageant dans la direction de l’axe (Ox) mais cette fois dans le sens négatif de cet axe, c’est-à-dire vers les x décroissants. Ce cas se ramène au cas précédent si l’on définit un axe (Ox') tel que $x' = -x$. Il suffit donc de changer x en $-x$ dans le résultat. Ainsi :

Une **onde progressive** se propageant à la vitesse c dans la direction de l’axe (Ox) , *dans le sens négatif* de cet axe, sans atténuation ni déformation, est de la forme mathématique suivante :

$$s(x,t) = g\left(t + \frac{x}{c}\right),$$

où g est une fonction quelconque dont l’argument a la dimension d’un temps.

c) Deuxième expression de l’onde progressive

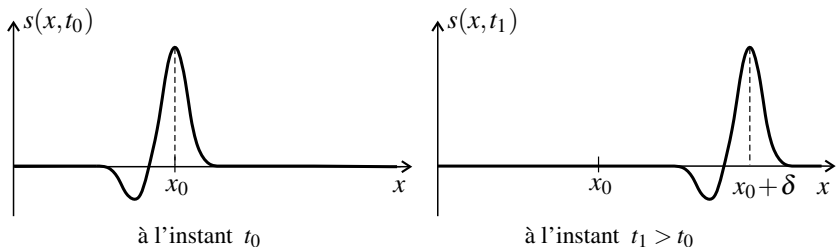


Figure 2.13 – Onde se propageant sans atténuation ni déformation dans le sens positif de (Ox) , à deux instants différents.

La figure 2.13 représente les valeurs du signal à deux instants différents t_0 et $t_1 > t_0$ pour la même onde que sur la figure 2.12.

Entre les instants t_0 et t_1 , l’onde qui progresse dans le sens positif de l’axe (Ox) se déplace d’une distance δ . Ainsi la valeur observée à l’instant t_0 en x est observée à l’instant t_1 en $x + \delta$:

$$s(x,t_0) = s(x + \delta, t_1).$$

L’onde se propage à la vitesse c donc on a :

$$\delta = c(t_1 - t_0).$$

Il vient ainsi :

$$s(x,t_0) = s(x + c(t_1 - t_0), t_1). \tag{2.4}$$

De même que plus haut, on peut se convaincre facilement que la formule (2.4) est valable aussi si $t_1 < t_0$. Si on l'écrit en posant $t_0 = t$, instant quelconque, et $t_1 = 0$ elle devient alors :

$$s(x, t) = s(x - ct, 0).$$

Le membre de droite de cette équation est simplement une fonction d'une seule variable, $x - ct$. Pour simplifier l'écriture on le note $F(x - ct)$. Ainsi :

Une **onde progressive** se propageant à la vitesse c dans la direction de l'axe (Ox) , *dans le sens positif* de cet axe, sans atténuation ni déformation, est de la forme mathématique suivante :

$$s(x, t) = F(x - ct),$$

où F est une fonction quelconque dont l'argument a la dimension d'une longueur.

Si l'on s'intéresse aussi à une onde se propageant le long de l'axe (Ox) dans le sens négatif de cet axe, c'est-à-dire le sens des x décroissants, le même raisonnement s'applique si l'on prend un axe (Ox') tel que $x' = -x$. Il suffit donc de changer x en $-x$ dans le résultat. On obtient une fonction de la variable $-x - ct$ que l'on peut changer, pour avoir une notation plus simple, pour son opposé $x + ct$. Ainsi :

Une **onde progressive** se propageant à la vitesse c dans la direction de l'axe (Ox) , *dans le sens négatif* de cet axe, sans atténuation ni déformation, est de la forme mathématique suivante :

$$s(x, t) = G(x + ct),$$

où G est une fonction quelconque dont l'argument a la dimension d'une longueur.

d) Équivalence des deux expressions

Comme on vient de le montrer, une onde se propageant dans le sens positif de l'axe (Ox) peut s'écrire de deux manières :

$$s(x, t) = f\left(t - \frac{x}{c}\right) = F(x - ct).$$

Si l'on connaît la fonction f on trouve la fonction F en posant $t = 0$ dans l'équation ci-dessus :

$$F(x) = f\left(-\frac{x}{c}\right),$$

et si l'on connaît la fonction F , on trouve la fonction f en posant $x = 0$ dans cette équation :

$$f(t) = F(-ct).$$

Dans le cas d'une onde se propageant dans le sens négatif de l'axe (Ox) , on obtient la même onde avec les deux formules $g\left(t + \frac{x}{c}\right)$ et $G(x + ct)$ si :

$$G(x) = g\left(\frac{x}{c}\right) \quad \text{soit} \quad g(t) = G(ct).$$

CHAPITRE 2 – PROPAGATION D’UN SIGNAL

2.3 Onde progressive sinusoïdale

Dans ce paragraphe on considère une onde progressive se propageant *dans le sens positif de l’axe (Ox)*, sans déformation ni atténuation. Cette onde aura une forme particulière : elle sera sinusoïdale.

a) Définition

On parle d’**onde sinusoïdale**, ou encore d’**onde harmonique**, lorsque le signal mesuré en tout point est une fonction sinusoïdale du temps de pulsation ω , indépendante du point. Une telle onde a la forme mathématique suivante :

$$s(x,t) = A(x) \cos(\omega t + \varphi(x)).$$

$A(x)$ est l’amplitude de l’onde au point d’abscisse x et $\varphi(x)$ la phase initiale de l’onde en ce même point.

On note : $A_0 = A(0)$ et $\varphi(0) = \varphi_0$, l’amplitude et la phase initiale à l’origine O de l’axe Ox .

b) Double périodicité spatio-temporelle

L’onde se propageant sans atténuation ni déformation dans le sens positif de (Ox) , la vibration observée à toute abscisse $x > 0$ reproduit la vibration observée en $x = 0$ avec le retard de propagation $\tau = \frac{x}{c}$. Ceci s’exprime ainsi :

$$s(x,t) = s(0,t - \tau) = A_0 \cos\left(\omega\left(t - \frac{x}{c}\right) + \varphi_0\right).$$

On retiendra ce résultat sous la forme suivante :

Une **onde progressive sinusoïdale** de pulsation ω se propageant dans le sens positif de l’axe (Ox) avec la vitesse c a pour expression :

$$s(x,t) = A_0 \cos(\omega t - kx + \varphi_0) \quad \text{avec} \quad k = \frac{\omega}{c}. \tag{2.5}$$

k est appelé **vecteur d’onde**. A_0 est l’**amplitude** de l’onde et φ_0 sa **phase initiale à l’origine**.

Ainsi l’onde progressive sinusoïdale est une fonction sinusoïdale à la fois :

- du temps (pour x fixé) avec une pulsation temporelle ω ,
- de la variable spatiale x (pour t fixé), avec une pulsation spatiale k .

On parle de la **double périodicité spatio-temporelle** de l’onde.

De même que la période temporelle est : $T = \frac{2\pi}{\omega}$, la période spatiale est la **longueur d’onde** :

$$\lambda = \frac{2\pi}{k} = \frac{2\pi c}{\omega} = cT.$$

La fréquence spatiale est le **nombre d’onde** :

$$\sigma = \frac{1}{\lambda}.$$

Une onde progressive sinusoïdale se propageant dans le sens positif de l'axe (Ox) peut s'écrire :

$$s(x,t) = A_0 \cos \left(2\pi \left(\frac{t}{T} - \frac{x}{\lambda} \right) + \varphi_0 \right) \tag{2.6}$$

où T est la période temporelle et λ la période spatiale appelée **longueur d'onde**.
La longueur d'onde est égale à la distance sur laquelle l'onde se propage pendant une durée égale à la période temporelle T :

$$\lambda = cT. \tag{2.7}$$

On peut trouver toutes les grandeurs relatives à la périodicité spatio-temporelle quand on connaît l'une d'entre elles et la vitesse de propagation c . Si on connaît la période temporelle T alors :

$$\omega = \frac{2\pi}{T}, \quad f = \frac{1}{T}, \quad \lambda = cT \quad \text{et} \quad \sigma = \frac{1}{cT}.$$

Si on connaît la fréquence f :

$$\omega = 2\pi f, \quad T = \frac{1}{f}, \quad \lambda = \frac{c}{f} \quad \text{et} \quad \sigma = \frac{f}{c}.$$

Si on connaît la longueur d'onde λ :

$$\sigma = \frac{1}{\lambda}, \quad T = \frac{\lambda}{c}, \quad f = \frac{c}{\lambda}, \quad \text{et} \quad \omega = \frac{2\pi c}{\lambda}.$$

Les grandeurs relatives à la double périodicité spatio-temporelle de l'onde progressive sinusoïdale sont rassemblées dans le tableau suivant :

	Période	Fréquence	Pulsation
Temps	T	f	ω
Espace	λ	σ	k

Tableau 2.2 – Grandeurs caractérisant la double périodicité de l'onde progressive sinusoïdale.

c) Interprétation physique

Le lien entre les périodes temporelle T et spatiale λ exprimé dans la formule (2.7) se retrouve par le raisonnement suivant.

Comme on le voit sur la figure 2.14, la courbe représentant l'onde en fonction de x à l'instant t_1 est la courbe à l'instant $t_0 < t_1$ décalée vers la droite de $\delta = c(t_1 - t_0)$. Ceci fait qu'en chaque point la valeur de l'onde change entre t_0 et t_1 . Si on suppose maintenant que $t_1 = t_0 + T$ l'onde a les mêmes valeurs à t_0 et à t_1 . Il faut pour cela que la courbe soit décalée d'une période spatiale soit $\delta = \lambda$. Ainsi : $\lambda = c(t_1 - t_0) = cT$.

CHAPITRE 2 – PROPAGATION D’UN SIGNAL

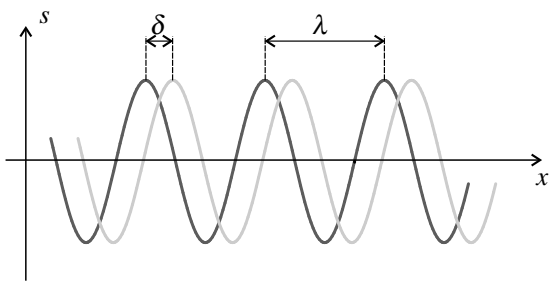


Figure 2.14 – Onde sinusoïdale se propageant dans le sens positif de (Ox) à deux instants t_0 (en gris foncé) et $t_1 > t_0$ (en gris clair). Le décalage des courbes est $\delta = c(t_1 - t_0)$.

d) Déphasage entre les vibrations en deux points

La phase initiale de la vibration à l’abscisse x est, d’après l’équation (2.5) :

$$\varphi(x) = \varphi_0 - kx = \varphi_0 - \frac{2\pi x}{\lambda}.$$

Les signaux d’une onde sinusoïdale se propageant dans le sens positif de (Ox) en deux points d’abscisses x_1 et x_0 sont déphasés de :

$$\varphi(x_1) - \varphi(x_0) = -\frac{2\pi}{\lambda}(x_1 - x_0).$$

Remarque

Ce résultat se retrouve très facilement. Si $x_1 > x_0$, la vibration en x_1 est en retard sur celle en x_0 de $\tau = \frac{x_1 - x_0}{c}$. Le retard temporel se traduit par un retard de phase :

$$\varphi(x_1) - \varphi(x_0) = -2\pi \frac{\tau}{T} = -\frac{2\pi}{cT}(x_1 - x_0) = -\frac{2\pi}{\lambda}(x_1 - x_0).$$

A quelle condition les vibrations en deux points d’abscisses x_0 et x_1 sont-elles en phase ? Cela signifie que $\varphi(x_0) - \varphi(x_1) = m \times 2\pi$ où m est un entier relatif, soit $\frac{2\pi}{\lambda}(x_1 - x_0) = m \times 2\pi$, soit :

$$x_1 - x_0 = m \times \lambda.$$

De même, les vibrations en ces deux points sont en opposition de phase si $\varphi(x_0) - \varphi(x_1) = \pi + m \times 2\pi$ où m est un entier relatif, soit si :

$$x_1 - x_0 = \left(m + \frac{1}{2}\right) \times \lambda.$$

Deux points en lesquels l'onde est **en phase** sont séparés le long de la direction de propagation (Ox) d'un **nombre entier de fois la longueur d'onde**.

Deux points en lesquels l'onde est **en opposition de phase** sont séparés le long de la direction de propagation (Ox) d'un **nombre entier de fois la longueur d'onde plus une demi-longueur d'onde**.

Ces résultats sont mis à profit dans l'expérience suivante pour mesurer la célérité d'ondes ultra-sonores dans l'air.

Expérience

On place un émetteur d'ultrasons E , produisant une onde sinusoïdale de fréquence $f = 44,0 \text{ kHz}$, devant deux récepteurs R_1 et R_2 fournissant des signaux (tensions électriques) proportionnels au signal de l'onde qu'ils reçoivent et que l'on envoie sur les deux voies d'un oscilloscope.

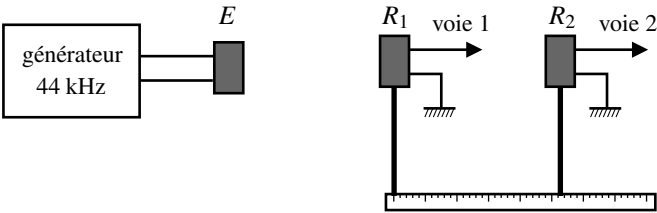


Figure 2.15 – Mesure de la célérité des ondes ultrasonores dans l'air.

En « mode normal », on voit à l'écran de l'oscilloscope deux sinusoïdes d'amplitudes à peu près égales si les récepteurs sont à peu près à la même distance de l'émetteur. Si l'on passe l'oscilloscope en « mode XY », on voit (en général) une ellipse. Lorsque le déphasage des deux signaux varie, cette ellipse se transforme en un segment de droite de pente positive si les deux signaux sont en phase, et en un segment de droite de pente négative si les deux signaux sont en opposition de phase (voir chapitre précédent).

R_1 est fixe. R_2 peut être translaté le long d'une règle graduée, parallèle à la direction entre l'émetteur et ce récepteur. On commence par ajuster sa position pour voir deux signaux en phase, c'est-à-dire, sur l'oscilloscope, un segment de droite de pente positive. Puis on le recule progressivement en observant l'écran de l'oscilloscope : on voit que les signaux sont alternativement en opposition de phase et en phase. Quand les signaux sont pour la 5^{ème} fois de suite en phase on note la position du récepteur mobile sur la règle, puis on continue. Lors d'un essai, on a noté les positions successives, en centimètres : 42,0 ; 37,9 ; 33,9 ; 29,9 ; 26,0. La règle étant graduée en millimètre l'incertitude sur chacune de ces valeurs est de 0,3 mm.

La distance moyenne entre deux positions successives de la liste précédente est : $\frac{42,0 - 26,0}{4} = 4,0 \text{ cm} \pm 0,15 \text{ mm}$. On en déduit une évaluation de la longueur d'onde :

CHAPITRE 2 – PROPAGATION D’UN SIGNAL

$\lambda = \frac{4,0}{5} = 0,80 \text{ cm} = 8,0 \text{ mm} \pm 0,03 \text{ mm}$, puis de la célérité des ondes ultra-sonores dans l’air :

$$c = \lambda f = 352 \pm 2 \text{ m.s}^{-1}.$$

SYNTHÈSE

SAVOIRS

- connaître les grands types d’onde et la nature du signal propagé
- notion de spectre
- composition du spectre d’un signal périodique non sinusoïdal
- domaine fréquentiel des ondes sonores
- formes mathématiques d’une onde progressive
- formes mathématiques d’une onde progressive sinusoïdale (ou harmonique)
- relations entre ω et k , entre λ et T
- évolution de la phase d’une onde progressive sinusoïdale dans l’espace

SAVOIR-FAIRE

- obtenir expérimentalement le spectre d’un signal
- écrire une onde progressive de forme quelconque
- écrire une onde progressive sinusoïdale quelconque
- calculer les grandeurs relatives à la périodicité de l’onde à partir de l’une d’entre elles
- calculer le déphasage d’une onde entre deux points
- mesurer expérimentalement la célérité d’une onde
- mesurer expérimentalement une longueur d’onde

MOTS-CLÉS

- | | | |
|--------------------------|--------------------|--------------------|
| • onde | • harmoniques | • onde sinusoïdale |
| • signal physique | • son | • vecteur d’onde |
| • spectre | • propagation | • longueur d’onde |
| • composante sinusoïdale | • célérité | • déphasage |
| • fondamental | • onde progressive | |

S'ENTRAÎNER

2.1 Signal sans fondamental (★)

On considère le signal $s(t) = 10 \sin(80\pi t) + 5 \sin(120\pi t + 0,6\pi)$ où le temps est exprimé en secondes. Quelle est la fréquence de $s(t)$?

2.2 Spectre d'un produit de fonctions sinusoïdales (★)

On considère le signal : $s(t) = A \cos(2\pi f_1 t) \cos(2\pi f_2 t + \varphi)$ où A et φ sont des constantes.

- 1. En utilisant la formule de trigonométrie $\cos a \cos b = \frac{1}{2} (\cos(a + b) + \cos(a - b))$ déterminer les fréquences contenues dans $s(t)$. Représenter son spectre d'amplitude et de phase.
- 2. Examiner le cas où $f_1 = f_2$.

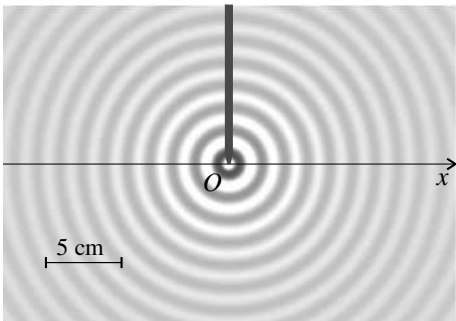
2.3 Relation entre fréquence et longueur d'onde (★)

- 1. Calculer la longueur d'onde de l'onde électromagnétique qui existe dans un four micro-onde sachant que sa fréquence est $f = 2,45 \text{ GHz}$ et que la célérité des ondes électromagnétiques est $c = 3,00 \cdot 10^8 \text{ m} \cdot \text{s}^{-1}$. Est-elle de l'ordre du micromètre ?
- 2. La vitesse du son dans l'air dépend de la température T selon la formule $c = \sqrt{\gamma \frac{RT}{M_{\text{air}}}}$ où $\gamma = 1,4$, $R = 8,314 \text{ J} \cdot \text{K}^{-1} \cdot \text{mol}^{-1}$ et $M_{\text{air}} = 29 \cdot 10^{-3} \text{ kg} \cdot \text{mol}^{-1}$. Calculer la fréquence d'un son de longueur d'onde $\lambda = 78 \text{ cm}$ lorsque la température vaut $T_1 = 290 \text{ K}$, puis $T_2 = 300 \text{ K}$. Le changement de hauteur du son dû au changement de température est-il de plus d'un demi-ton ? Un demi-ton correspond à une variation relative de fréquence égale à $2^{\frac{1}{12}} - 1$.

2.4 Cuve à ondes (★)

La figure représente la surface d'une cuve à onde éclairée en éclairage stroboscopique. L'onde est engendrée par un vibreur de fréquence $f = 18 \text{ Hz}$. L'image est claire là où la surface de l'eau est convexe, foncée là où elle est concave.

- 1. En mesurant sur la figure, déterminer la longueur d'onde.
- 2. En déduire la célérité de l'onde.
- 3. On suppose l'onde sinusoïdale, d'amplitude A constante et de phase initiale nulle en O . Écrire le signal $s(x, t)$ pour $x > 0$ et pour $x < 0$.
- 4. Expliquer pourquoi A n'est pas, en fait, constante.



2.5 Ondes progressives sinusoïdales (★)

- 1. Donner la période, la fréquence, la pulsation, la longueur d'onde, le nombre d'onde et le vecteur d'onde, de l'onde : $s(x, t) = 5 \sin(2,4 \cdot 10^3 \pi t - 7,0 \pi x + 0,7 \pi)$ où x et t sont exprimés respectivement en mètres et en secondes. Quelle est sa vitesse de propagation ?

CHAPITRE 2 – PROPAGATION D'UN SIGNAL

2. Une onde sinusoïdale se propage dans la direction de l'axe (Ox) dans le sens positif avec la célérité c . L'expression du signal de l'onde au point d'abscisse x_1 est $s_1(x_1, t) = A \cos(\omega t)$. Déterminer l'expression de $s_1(x, t)$. Représenter $s_1(x, 0)$ en fonction de x .

3. Une onde sinusoïdale se propage dans la direction de l'axe (Ox) dans le sens négatif avec la célérité c . On donne : $s_2(0, t) = A \sin(\omega t)$. Déterminer l'expression de $s_2(x, t)$. Représenter graphiquement $s_2\left(\frac{\lambda}{4}, t\right)$ et $s_2\left(\frac{\lambda}{2}, t\right)$ en fonction de t .

2.6 Trains d'ondes (★)

Une onde se propage dans la direction de l'axe (Ox), dans le sens positif avec la célérité c . La source, située en $x = 0$, émet un train d'ondes, c'est-à-dire une oscillation de durée limitée τ :

$$s(0, t) = \begin{cases} 0 & \text{si } t < 0 \\ \sin\left(2\pi \frac{t}{T}\right) & \text{si } 0 \leq t < \tau \\ 0 & \text{si } t \geq \tau \end{cases}$$

1. Exprimer $s(x, t)$ pour x positif quelconque.
2. Représenter $s\left(x, \frac{\tau}{2}\right)$ et $s\left(x, \frac{3\tau}{2}\right)$ en fonction de x pour $x > 0$ (prendre $\tau = 4T$ pour le dessin). Quelle est la longueur du train d'ondes dans l'espace ?
3. On suppose à présent que :

$$s(0, t) = \begin{cases} 0 & \text{si } t < 0 \\ \exp\left(-\frac{t}{\tau}\right) \sin\left(2\pi \frac{t}{T}\right) & \text{si } t \geq 0 \end{cases}$$

avec $\tau = T$. En s'aidant d'une calculatrice graphique représenter $s(0, t)$ en fonction de t , puis $s(x, 6\tau)$ en fonction de x . On considère habituellement que ce train d'ondes dure 6τ . Justifier et donner la longueur du train d'ondes.

2.7 Propagation d'un signal périodique (★)

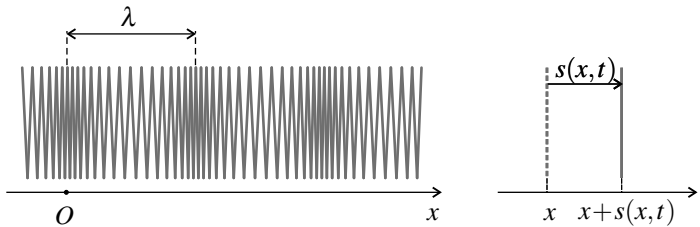
Une onde se propage dans le sens positif de (Ox) à la célérité c . En $x = 0$ son signal est périodique de fréquence f_S et s'écrit : $s(0, t) = f(t) = \sum_{n=1}^{\infty} A_n \cos(2\pi n f_S t + \varphi_n)$.

1. Quelle est la longueur d'onde associée au fondamental du signal $f(t)$? À son harmonique de rang n ? Quelle est la période spatiale de l'onde ?
2. Montrer que le signal reçu en x_0 s'écrit : $s(x_0, t) = \sum_{n=1}^{\infty} A'_n \cos(2\pi n f_S t + \varphi'_n)$ et exprimer les A'_n et φ'_n en fonction de x_0 .

2.8 Onde longitudinale sur un ressort (★★)

L'onde de compression-dilatation le long d'un ressort est une onde longitudinale analogue à une onde sonore. Lors du passage de cette onde chaque spire bouge dans la direction de l'axe (Ox), axe parallèle au ressort. Un signal associé à l'onde est le déplacement $s(x, t)$ de la

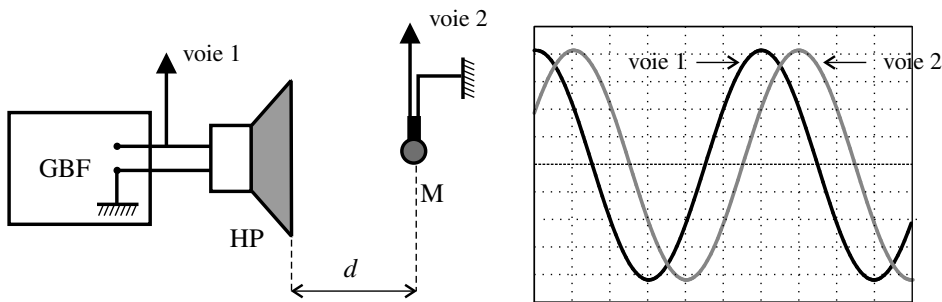
spire qui est située à l'abscisse x en l'absence d'onde : cette spire passe ainsi de la position x qu'elle occupe au repos à la position $x + s(x, t)$ (voir figure).



- 1. On appelle a l'espacement entre deux spires consécutives dans l'état de repos. Lors du passage de l'onde, la distance entre les spires situées au repos en $x_i = ia$ et $x_{i+1} = (i + 1)a$ devient d_i . Exprimer d_i en fonction de a , $s(x_i, t)$ et $s(x_{i+1}, t)$.
- 2. On suppose que $s(x, t) = A \cos(\omega t - kx + \varphi)$. On pose $\Phi = \omega t + \varphi$. En utilisant une formule de trigonométrie, montrer que : $d_i = a + 2A \sin\left(\frac{ka}{2}\right) \sin\left(\Phi - kx_i - \frac{ka}{2}\right)$.
- 3. On suppose que $ka \ll 1$. Cette hypothèse correspond-elle bien à la figure ci-dessus ? Montrer que, lors du passage de l'onde, la distance entre deux spires consécutives situées au repos au voisinage de x devient : $d(x, t) \simeq a(1 + kA \sin(\omega t - kx + \varphi))$.
- 4. La figure donne l'allure l'allure du ressort à $t = 0$. En déduire φ . Où se trouvent, sur la figure, les spires dont le déplacement est nul ?

APPROFONDIR

2.9 Étude expérimentale d'une onde progressive sinusoïdale (★)



Un haut-parleur HP est mis en vibration à l'aide d'un générateur de basses fréquences GBF réglé sur la fréquence $f = 1500 \text{ Hz}$. L'onde sonore ainsi créée se propage dans l'air à la célérité $c = 342 \text{ m}\cdot\text{s}^{-1}$. Un microphone M placé à distance d du haut-parleur reçoit le signal

CHAPITRE 2 – PROPAGATION D’UN SIGNAL

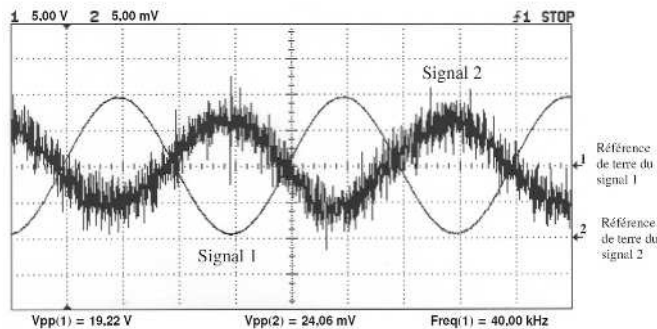
sonore et le transforme en un signal électrique. Les signaux du GBF et du micro sont envoyés respectivement sur les voies 1 et 2 d’un oscilloscope.

1. Pour une certaine position de M et un réglage adéquat de l’oscilloscope, l’écran a l’aspect représenté sur la figure. Quel est le déphasage des signaux visualisés ?
2. L’oscilloscope étant synchronisé sur la voie 1, comment évolue la courbe de la voie 2 lorsqu’on éloigne M de HP ?
3. De combien doit-on augmenter d pour voir deux signaux en phase ? Quel est le meilleur moyen pour savoir si les signaux sont en phase ?

2.10 Principe de la télémétrie (CAPES externe 2010) (★)

On place un émetteur et un récepteur à ultrasons côte à côte. Ce bloc est appelé le télémètre. À la distance D , on place un obstacle réfléchissant les ondes sonores, que nous appellerons la cible. Une onde sinusoïdale, de période T , est émise par l’émetteur du télémètre, elle se réfléchit sur la cible et est détectée par le récepteur du télémètre. Sur l’écran d’un oscilloscope, on visualise simultanément deux signaux ; celui capté (par un dispositif non décrit) en sortie de l’émetteur et celui du récepteur.

1. On appelle temps de vol, noté t_v , la durée du trajet aller-retour de l’onde entre le télémètre et la cible. Exprimer t_v en fonction de la distance D séparant le télémètre de la cible et de la célérité c de l’onde.
2. Pour illustrer le principe de la mesure, on colle la cible au télémètre, puis on l’éloigne lentement, en comptant le nombre de coïncidences, c’est-à-dire le nombre de fois où les signaux sont en phase. Pour simplifier, on suppose que lorsque $D = 0$, les signaux sont en phase. On se place dans le cas où l’on a compté exactement un nombre n de coïncidences. Exprimer D en fonction de n et de la longueur d’onde des ondes ultrasonores.



3. Lors du recul de la cible, 50 coïncidences ont été comptées avant d’observer les signaux suivants sur l’écran de l’oscilloscope (voir figure). Dans les conditions de l’expérience, la longueur d’onde des ondes sonores valait 8,5 mm. En exploitant les données de l’enregistrement, calculer la distance séparant le télémètre de la cible.
4. Pourquoi les deux signaux de la figure sont-ils si différents ? Identifier quel est, selon toute vraisemblance, le signal capté en sortie de l’émetteur et celui reçu par le récepteur.

5. Le comptage des coïncidences a été réalisé en plaçant l'oscilloscope en mode XY. Dans le cas des signaux de la figure, représenter la figure que l'on obtiendrait en se plaçant dans ce mode.

2.11 Effet Doppler (★)

Une onde sinusoïdale de fréquence f se propage dans la direction de (Ox) dans le sens positif de (Ox) avec la célérité c . Un observateur se déplace avec une vitesse $\vec{v} = v\vec{u}_x$ parallèle à (Ox) .

1. Écrire le signal $s(x, t)$ de l'onde en définissant les notations nécessaires.
2. Pour l'observateur en mouvement, le point d'abscisse x est repéré par une abscisse le long d'un axe (Ox') qui lui est lié telle que $x' = x - vt$. Exprimer $s(x', t)$.
3. En déduire l'expression de la fréquence f' pour l'observateur en mouvement. Comparer f' et f suivant le signe de v .
4. Vous marchez dans la rue et un camion de pompier, sirène en marche, arrive de derrière et vous dépasse. Qu'entendez-vous ?

2.12 Position et date d'un séisme (★)

Un séisme produit deux types d'ondes sismiques : les ondes P, longitudinales, qui se propagent avec la célérité c_P et les ondes S, transversales, qui se propagent avec la célérité $c_S < c_P$.

1. Lors d'un séisme, on commence à détecter les premières à l'instant de date t_P et les secondes à l'instant de date t_S . Montrer qu'on peut en déduire, connaissant c_P et c_S , la distance Δ entre le foyer du séisme et l'appareil ainsi que la date du début du séisme.
2. Pour un séisme, on mesure les distances Δ_1 , Δ_2 et Δ_3 entre le foyer du séisme et trois stations de mesures. Sans faire de calcul, montrer que cette information permet de localiser le foyer du séisme à l'intérieur de la Terre. Quel système fonctionne sur ce même principe ?

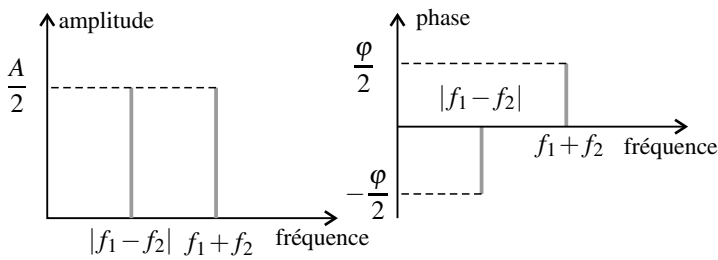
CORRIGÉS

2.1 Signal sans fondamental

$s(t)$ est la somme de deux signaux de fréquences $f_1 = 40 \text{ Hz} = 2 \times 20 \text{ Hz}$ et $f_2 = 60 \text{ Hz} = 3 \times 20 \text{ Hz}$. Il s’agit d’un signal de fréquence $f = 20 \text{ Hz}$ qui n’a pas de fondamental mais seulement deux harmoniques de rang 2 et 3.

2.2 Spectre d’un produit de fonctions sinusoïdales

1. $s(t) = \frac{A}{2} \cos(2\pi(f_1 + f_2)t + \varphi) + \frac{A}{2} \cos(2\pi(f_1 - f_2)t - \varphi)$. Ainsi le spectre de $s(t)$ comporte les deux fréquences $f_1 + f_2$ et $|f_1 - f_2|$ (attention à ne pas oublier la valeur absolue : une fréquence est obligatoirement positive).



2. Dans le cas où $f_1 = f_2$, $s(t) = \frac{A}{2} \cos \varphi + \frac{A}{2} \cos(2f_1 t + \varphi)$. Le spectre de $s(t)$ comporte les fréquences 0, correspondant à un signal constant, et $2f_1$.

2.3 Relation entre fréquence et longueur d’onde

1. $\lambda = \frac{c}{f} = \frac{3,00.10^8}{2,45.10^9} = 1,22.10^{-1} \text{ m}$. Les « micro-ondes » n’ont pas une longueur d’onde de l’ordre du micromètre.

2. La fréquence d’un son de longueur d’onde λ est : $f = \frac{c}{\lambda} = \sqrt{\gamma \frac{RT}{M_{\text{air}}}} \frac{1}{\lambda}$. Pour $\lambda = 78 \text{ cm}$, on trouve $f_1 = 437 \text{ Hz}$ à $T_1 = 290 \text{ K}$ et $f_2 = 445 \text{ Hz}$ à $T_2 = 300 \text{ K}$. La variation relative de fréquence entre les deux températures est $\frac{f_2 - f_1}{f_1} = 1,7\%$. Pour une note montant d’un demi-ton la variation relative de fréquence est $2^{\frac{1}{12}} - 1 = 5,9\%$. La variation de température entraîne un changement de hauteur inférieur au demi-ton.

2.4 Cuve à ondes

1. La longueur d'onde est la distance entre deux crêtes consécutives le long de l'axe (Ox) , donc entre les centres de deux zones claires. On mesure sur la figure 2, 1 cm pour 7λ , l'échelle étant de $\frac{1}{5}$ la longueur d'onde est $\lambda = 5 \times \frac{2,1}{7} = 1,5$ cm.
2. La célérité de l'onde est $c = \lambda f = 1,5 \cdot 10^{-2} \times 18 = 0,27$ m·s⁻¹.
3. Pour $x > 0$ l'onde se propage dans le sens de (Ox) donc $s(x,t) = A \cos\left(2\pi f\left(t - \frac{x}{c}\right)\right)$; pour $x < 0$ elle se propage dans le sens inverse de l'axe (Ox) : $s(x,t) = A \cos\left(2\pi f\left(t + \frac{x}{c}\right)\right)$. On peut résumer les deux expressions en :

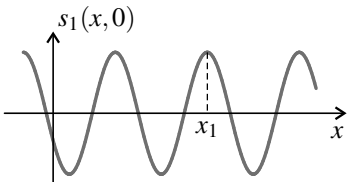
$$s(x,t) = A \cos\left(2\pi f\left(t - \frac{|x|}{c}\right)\right).$$

4. L'amplitude n'est pas constante parce que l'énergie de l'onde se répartit sur ces cercles de rayon de plus en plus grand. Des considérations en dehors du programme permettent d'établir que $A = \frac{\text{constante}}{\sqrt{x}}$.

2.5 Ondes progressives sinusoïdales

1. Période : $T = 8,3 \cdot 10^{-4}$ s, pulsation : $\omega = 7,5 \cdot 10^3$ rad·s⁻¹, fréquence : $f = 1,2 \cdot 10^3$ Hz. Longueur d'onde : $\lambda = 0,29$ m, vecteur d'onde : $k = 22$ rad·m⁻¹, nombre d'onde : $\sigma = 3,5$ m⁻¹.
- La vitesse de propagation est : $c = \lambda f = 348$ m·s⁻¹.
2. L'onde se propageant avec la célérité c dans le sens positif de (Ox) , on a :

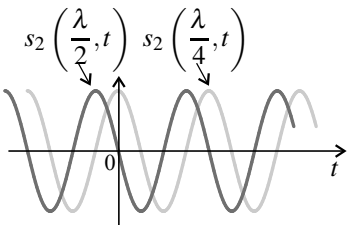
$$s_1(x,t) = s_1\left(x_1, t - \frac{x-x_1}{c}\right) = A \cos(\omega t - kx + kx_1),$$



en posant $k = \frac{\omega}{c}$. Pour tracer $s(x, t = 0)$ on remarque que le signal est maximal en $x = x_1$ à $t = 0$.

3. L'onde se propageant avec la célérité c dans le sens négatif de (Ox) , on a :

$$s_2(x,t) = s_2\left(0, t + \frac{x}{c}\right) = A \sin(\omega t + kx).$$



$s_2\left(\frac{\lambda}{4}, t\right)$ est en quadrature avance sur $s_2(0, t)$ et $s_2\left(\frac{\lambda}{2}, t\right)$ est en quadrature avance sur $s_2\left(\frac{\lambda}{4}, t\right)$ et en opposition de phase par rapport à $s_2(0, t)$.

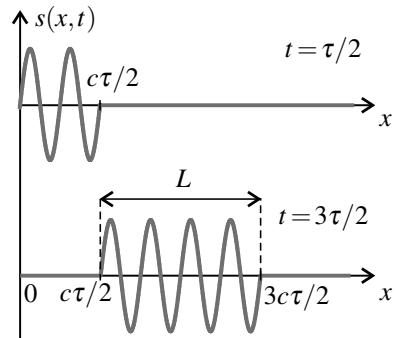
CHAPITRE 2 – PROPAGATION D'UN SIGNAL

2.6 Trains d'ondes

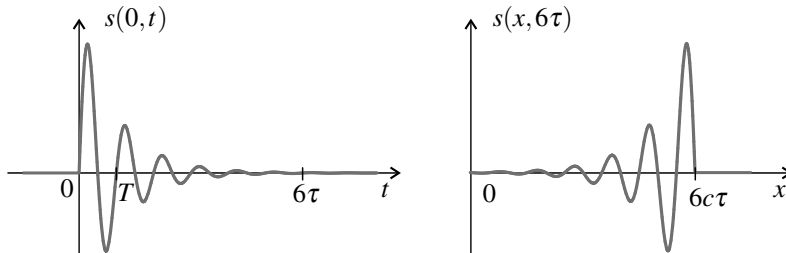
1. $s(x, t) = s(0, t - x/c)$, donc :

$$s(x, t) = \begin{cases} 0 & \text{si } x > ct \\ \sin\left(\frac{2\pi t}{T} - \frac{2\pi x}{cT}\right) & \text{si } c(t - \tau) < x \leq ct \\ 0 & \text{si } x \leq c(t - \tau) \end{cases}$$

La longueur du train d'onde dans l'espace est : $L = c\tau$.



2. Sur le graphe de $s(0, t)$ on constate que les oscillations ne sont plus visibles en gros à partir de $t = 6\tau$. Plus précisément, pour $t > 6\tau$, $\exp\left(-\frac{t}{\tau}\right) < \exp(-6) = 2,5 \cdot 10^{-3}$ donc l'amplitude d'oscillation est inférieure à $\frac{1}{400}$ de sa valeur initiale. Ceci justifie le fait de considérer que la durée du train d'onde est 6τ . Dans ce cas la longueur du train d'onde dans l'espace est $6c\tau$, comme on peut le voir sur la partie droite de la figure ci-dessous.



2.7 Propagation d'un signal périodique

1. La longueur d'onde associée au fondamental, composante sinusoïdale de fréquence f_s , est $\lambda_1 = \frac{c}{f_s}$. La longueur d'onde associée à l'harmonique de rang n est $\lambda_n = \frac{c}{nf_s} = \frac{\lambda_1}{n}$.

La période spatiale de l'onde est λ_1 .

2. Le phénomène de propagation à la célérité c dans le sens positif de l'axe (Ox) se traduit par la relation :

$$s(x_0, t) = s\left(0, t - \frac{x_0}{c}\right) = \sum_{n=1}^{\infty} A_n \cos\left(2\pi n f_s t - 2\pi n \frac{f_s}{c} x_0 + \varphi_n\right).$$

On trouve bien la forme attendue, avec : $A'_n = A_n$ et $\varphi'_n = \varphi_n - 2\pi n \frac{f_s}{c} x_0 = \varphi_n - 2\pi \frac{x_0}{\lambda_n}$. On remarque que le retard de phase dû à la propagation dépend de l'ordre n de l'harmonique.

2.8 Onde longitudinale le long d'un ressort

1. Lors du passage de l'onde l'abscisse de la $i^{\text{ème}}$ spire devient $x_i + s(x_i, t)$, et l'abscisse de la $(i + 1)^{\text{ème}}$ spire devient $x_{i+1} + s(x_{i+1}, t)$. La distance entre ces deux spires est alors :

$$d_i = ((x_{i+1} + s(x_{i+1}, t)) - (x_i + s(x_i, t))) = a + s(x_{i+1}, t) - s(x_i, t).$$

2. En utilisant la formule $\cos p - \cos q = -2 \sin \frac{p+q}{2} \sin \frac{p-q}{2}$, on trouve :

$$\begin{aligned} d_i &= a + A \cos(\Phi - kx_{i+1}) - A \cos(\Phi - kx_i) \\ &= a + 2A \sin\left(\frac{1}{2}k(x_{i+1} - x_i)\right) \sin\left(\Phi - \frac{1}{2}k(x_i + x_{i+1})\right) \\ &= a + 2A \sin\left(\frac{ka}{2}\right) \sin\left(\Phi - k\left(i + \frac{1}{2}\right)a\right) = a + 2A \sin\left(\frac{ka}{2}\right) \sin\left(\Phi - kx_i - \frac{ka}{2}\right). \end{aligned}$$

3. Si $ka \ll 1$ on peut écrire : $\sin\left(\frac{ka}{2}\right) \simeq \frac{ka}{2}$ et $\sin\left(\Phi - kx_i - \frac{ka}{2}\right) \simeq \sin(\Phi - kx_i)$. Il vient alors : $d_i \simeq a + kAa \sin(\Phi - kx_i)$.

Ainsi la distance entre des spires situées à l'abscisse x devient, lors du passage de l'onde :

$$d(x, t) = a(1 + kA \sin(\omega t - kx + \varphi)).$$

Ce signal est en quadrature de phase avec le signal $s(x, t)$ de déplacement des spires.

4. Sur la figure on observe qu'à $t = 0$ les spires sont serrées au maximum en $x = 0$. D'après la formule précédente, cela signifie que : $\sin \varphi = -1$, soit que : $\varphi = -\frac{\pi}{2}$.

Le déplacement des spires est donc $s(x, t) = A \cos(\omega t - kx - \frac{\pi}{2}) = A \sin(\omega t - kx)$.

Sur la figure, à $t = 0$, $s(x, 0) = -A \sin kx$, donc le déplacement est nul en tous les points d'abscisses $x = n\frac{\lambda}{2}$ où n est un entier et $\lambda = \frac{2\pi}{k}$ est la longueur d'onde. Ces points correspondent sur la figure aux endroits où les spires sont soit écartées au maximum, soit serrées au maximum.

2.9 Étude expérimentale d'une onde progressive sinusoïdale

1. La période T des signaux correspond à 6 carreaux sur l'axe horizontal. Leur décalage temporel τ correspond à 1 carreau et le signal de la voie 2 (micro) est en retard sur le signal de la voie 1 (GBF). Ainsi, le déphasage du signal capté par le microphone par rapport au signal du GBF est : $\Delta\varphi = 2\pi \frac{\tau}{T} = -2\pi \times \frac{1}{6} = -\frac{\pi}{3} = -1,05 \text{ rad}$.

2. Au fur et à mesure que l'on éloigne le micro, le signal correspondant (voie 2) est de plus en plus en retard sur le signal du GBF (voie 1). La courbe 2 se décale progressivement vers la droite sur l'écran.

CHAPITRE 2 – PROPAGATION D'UN SIGNAL

3. Pour voir deux signaux en phase, il faut augmenter le retard de phase du signal de la voie 2 de $2\pi - 1,05 = 5,24$ rad, c'est-à-dire augmenter le retard temporel de $\frac{5,24}{2\pi}T$ où $T = \frac{1}{f}$ est

la période. Pour cela il faut reculer le micro de $\frac{5,24}{2\pi} \times \frac{c}{f} = \frac{5,24}{2\pi} \times \frac{342}{1500} = 0,19$ m.

Pour constater expérimentalement que deux signaux sont en phase, la meilleure méthode consiste à passer l'oscilloscope en mode XY ; on voit un segment de droite de pente positive lorsque les signaux sont en phase.

2.10 Principe de la télémétrie

1. Le temps de vol est la durée de la propagation de l'onde sur une distance $2D$ correspondant à l'aller-retour jusqu'à la cible. Ainsi : $t_v = \frac{2D}{c}$.

2. Au fur et à mesure que l'on éloigne la cible, le décalage temporel entre les deux signaux augmente. Chaque fois que ce décalage est un multiple entier de la période T , les signaux sont en phases. Si on éloigne d'une distance D , on augmente de décalage temporel de $\frac{2D}{c}$. La

première fois que les signaux sont en phase à nouveau, on a $\frac{2D}{c} = T$ soit $D = \frac{cT}{2} = \frac{\lambda}{2}$ où λ est la longueur d'onde des ondes ultrasonores. La deuxième fois, $\frac{2D}{c} = 2T$ soit $D = 2 \times \frac{\lambda}{2}$.

À la $n^{\text{ième}}$ coïncidence, $D = n \times \frac{\lambda}{2}$.

3. Dans l'expérience on a reculé la cible d'une distance comprise entre $50 \times \frac{\lambda}{2}$ et $51 \times \frac{\lambda}{2}$. Les deux signaux sont en opposition de phase, ce qui veut dire qu'après la 50^e coïncidence on a reculé la cible d'une distance ΔD telle que $\frac{2\Delta D}{c} = \frac{T}{2}$ soit $\Delta D = \frac{\lambda}{4}$. La distance de la cible est donc : $D = 50 \frac{\lambda}{2} + \frac{\lambda}{4} = 10,8$ cm.

4. Le signal reçu par le récepteur en provenant de la cible est faible en raison de l'éloignement de celle-ci. Pour l'observer sur l'oscilloscope il faut augmenter la sensibilité, ce qui a pour conséquence d'amplifier le bruit électronique. C'est pourquoi ce signal a une allure irrégulière que l'on qualifie de « bruitée ».

5. Les deux signaux étant en opposition de phase on observerait en mode XY un segment de droite de pente négative (de contour assez flou à cause du bruit).

2.11 Effet Doppler

1. En notant A l'amplitude du signal et φ sa phase initiale à l'origine :

$$s(x, t) = A \cos \left(2\pi f \left(t - \frac{x}{c} \right) + \varphi \right).$$

2. $x = x' + vt$ donc

$$s(x', t) = A \cos \left(2\pi f \left(t - \frac{x' + vt}{c} \right) + \varphi \right) = A \cos \left(2\pi f \left(1 - \frac{v}{c} \right) t - 2\pi f \frac{x'}{c} + \varphi \right).$$

Du point de vue de l'observateur en mouvement l'onde a une fréquence

$$f' = f \left(1 - \frac{v}{c} \right).$$

La fréquence perçue par l'observateur est modifiée par son mouvement. C'est l'effet Doppler. D'après la formule, si $v > 0$ et $f' < f$: si l'observateur se déplace dans le sens de propagation de l'onde il perçoit une fréquence plus petite. Inversement, si $v < 0$, $f' > f$: si l'observateur se déplace en sens inverse de la propagation il perçoit une fréquence plus grande.

Remarque

L'expression de l'onde montre aussi que la longueur d'onde mesurée par l'observateur est $\lambda' = \frac{c}{f'} = \lambda$. On peut en déduire la vitesse de propagation de l'onde du point de vue de l'observateur : $c' = \lambda' f' = c - v$. Cette formule correspond à la cinématique classique.

3. Quand le camion des pompiers est derrière, il se rapproche. Le sens de la propagation est opposé au sens du déplacement de l'observateur par rapport au camion donc il perçoit une fréquence plus aigüe. Quand le camion est passé devant c'est le contraire.

2.12 Position et date d'un séisme

1. Si t_0 est la date à laquelle le séisme commence, on commence à détecter des ondes P à l'instant $t_P = t_0 + \frac{\Delta}{c_P}$ et des ondes S à l'instant $t_S = t_0 + \frac{\Delta}{c_S}$.

Ces deux relations constituent un système de 2 équations pour les 2 inconnues t_0 et Δ . La solution est :

$$\Delta = \frac{c_P c_S}{c_P - c_S} (t_S - t_P) \quad \text{et} \quad t_0 = \frac{c_P t_P - c_S t_S}{c_P - c_S}.$$

2. La connaissance de Δ_1 permet de savoir que le séisme s'est produit en un point situé sur la sphère $\mathcal{S}_1$, centrée sur la position de la première station et de rayon Δ_1 . Connaissant de plus Δ_2 on sait qu'il s'est produit en un point de l'intersection de $\mathcal{S}_1$ avec la sphère $\mathcal{S}_2$ centrée sur la position de la deuxième station et de rayon Δ_2 . L'intersection de ces deux sphères (nécessairement non vide) est un cercle $\mathcal{C}_{12}$. Connaissant de plus Δ_3 on sait que le séisme est localisé sur l'intersection de $\mathcal{C}_{12}$ avec la sphère $\mathcal{S}_3$ de rayon Δ_3 centrée sur la position de la troisième station. Cette intersection (nécessairement non vide) contient en général deux points dont un seul se trouve à l'intérieur de la Terre. On peut ainsi localiser le séisme.

Le système GPS (initiales de *Global Positioning System*), ainsi que le futur réseau européen Galileo, fonctionnent sur ce principe.

CHAPITRE 2 – PROPAGATION D’UN SIGNAL

Superpositions de deux signaux sinusoïdaux

3

Dans ce chapitre on va envisager différentes situations où deux signaux sinusoïdaux se superposent. Il s'agira de deux signaux sinusoïdaux de même fréquence, donnant le phénomène d'**interférences**. Puis on considérera deux signaux de fréquences voisines donnant le phénomène de **battements**. Enfin on étudiera un nouveau type d'onde, obtenu par superposition de deux ondes progressant en sens inverses, l'**onde stationnaire**.

1 Interférences entre deux ondes de même fréquence

1.1 Somme de deux signaux sinusoïdaux de même fréquence

Dans ce paragraphe on va s'intéresser au signal obtenu en faisant la somme de deux signaux sinusoïdaux de même fréquence f , donc de même pulsation $\omega = 2\pi f$.

a) Expérimentation mathématique

On peut visualiser, à l'aide d'une calculatrice graphique, le signal $s(t)$, somme des signaux $s_1(t) = 4\cos(2\pi t)$ et $s_2(t) = 3\cos(2\pi t + \varphi)$ de même fréquence $f = 1$ Hz, pour différentes valeurs de φ . On observe que $s(t)$ est un signal sinusoïdal de fréquence 1 Hz, dont l'amplitude et la phase initiale dépendent de φ .

En particulier, on peut voir sur la figure 3.1 que :

- pour $\varphi = 0$ l'amplitude est $4 + 3 = 7$,
- pour $\varphi = \pi/2$ l'amplitude est proche de 5,
- pour $\varphi = \pi$ l'amplitude est $4 - 3 = 1$,
- pour $\varphi = 3\pi/2$ l'amplitude est proche de 5.

De plus, si l'on change le deuxième signal en $s_2(t) = 4\cos(2\pi t + \varphi)$, signal de même amplitude que $s_1(t)$, on obtient, pour $\varphi = \pi$, un signal résultant $s(t)$ nul.

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

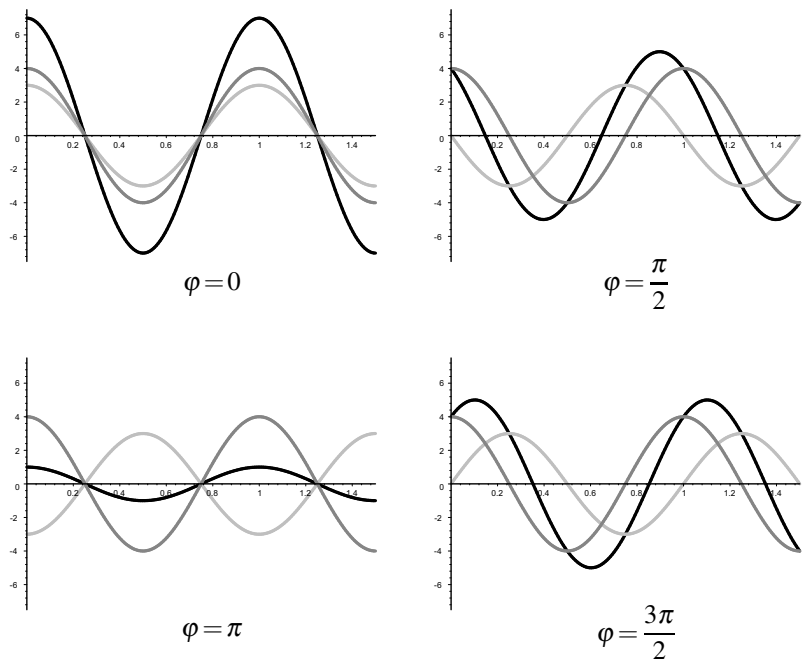


Figure 3.1 – Signaux $s_1(t)$ (en gris foncé), $s_2(t)$ (en gris clair) et $s_1(t) + s_2(t)$ (en noir) pour différentes valeurs de φ .

b) Calcul de l'amplitude du signal résultant

Soit les deux signaux : $s_1(t) = A_1 \cos(\omega t + \varphi_1)$ et $s_2(t) = A_2 \cos(\omega t + \varphi_2)$. Quelle est l'amplitude du signal $s(t) = s_1(t) + s_2(t)$?

Méthode analytique On peut écrire le signal résultant sous la forme :

$$\begin{aligned} s(t) &= s_1(t) + s_2(t) = A_1 \cos(\omega t + \varphi_1) + A_2 \cos(\omega t + \varphi_2) \\ &= (A_1 \cos \varphi_1 \cos(\omega t) - A_1 \sin \varphi_1 \sin(\omega t)) + (A_2 \cos \varphi_2 \cos(\omega t) - A_2 \sin \varphi_2 \sin(\omega t)) \\ &= (A_1 \cos \varphi_1 + A_2 \cos \varphi_2) \cos(\omega t) - (A_1 \sin \varphi_1 + A_2 \sin \varphi_2) \sin(\omega t), \end{aligned}$$

en utilisant la formule de trigonométrie $\cos(\alpha + \beta) = \cos \alpha \cos \beta - \sin \alpha \sin \beta$. Il apparaît ainsi clairement qu'il s'agit d'un signal sinusoïdal de même pulsation (donc fréquence) que les deux signaux s_1 et s_2 . Son amplitude A est donnée par la formule 1.8 du chapitre 1 :

$$\begin{aligned} A^2 &= (A_1 \cos \varphi_1 + A_2 \cos \varphi_2)^2 + (A_1 \sin \varphi_1 + A_2 \sin \varphi_2)^2 \\ &= A_1^2 + A_2^2 + 2A_1 A_2 (\cos \varphi_1 \cos \varphi_2 + \sin \varphi_1 \sin \varphi_2) \\ &= A_1^2 + A_2^2 + 2A_1 A_2 \cos(\varphi_1 - \varphi_2), \end{aligned}$$

où on a utilisé la formule $\cos^2 \alpha + \sin^2 \alpha = 1$, ainsi que la même formule de trigonométrie que plus haut.

Le signal résultant de la superposition de deux signaux sinusoïdaux de même fréquence, d'amplitudes A_1 et A_2 et de phases initiales φ_1 et φ_2 est un signal sinusoïdal de même fréquence et d'amplitude donnée par la **formule des interférences** :

$$A = \sqrt{A_1^2 + A_2^2 + 2A_1A_2 \cos(\varphi_2 - \varphi_1)}. \tag{3.1}$$

Cette amplitude n'est pas égale à la somme des amplitudes des deux signaux.

Méthode géométrique (MPSI) La formule (3.1) peut s'établir de manière élégante en utilisant la représentation des signaux sinusoïdaux par des vecteurs de Fresnel.

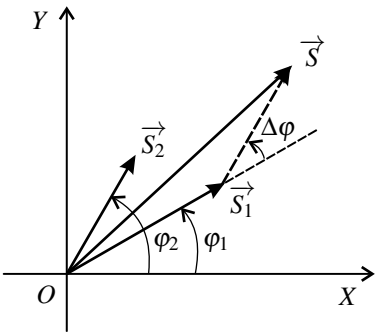


Figure 3.2 – Addition des vecteurs de Fresnel correspondant à deux signaux sinusoïdaux de même fréquence. Les vecteurs sont représentés à l'instant $t = 0$.

La figure 3.2 représente les vecteurs $\vec{S}_1$ et $\vec{S}_2$ correspondant respectivement aux signaux $s_1(t)$ et $s_2(t)$ (sur cette figure tous les vecteurs de Fresnel sont représentés à l'instant $t = 0$). On se rappelle que les normes de ces vecteurs sont égales aux amplitudes des signaux : $\|\vec{S}_1\| = A_1$ et $\|\vec{S}_2\| = A_2$.

On a représenté aussi le vecteur de Fresnel du signal somme $s(t)$: $\vec{S} = \vec{S}_1 + \vec{S}_2$.

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX



Pour dessiner $\vec{S} = \vec{S}_1 + \vec{S}_2$, il suffit de reproduire à partir de l'extrémité de $\vec{S}_1$ la flèche représentant $\vec{S}_2$ (même longueur et même direction).

L'amplitude A de $s(t)$ est donnée par :

$$\begin{aligned} A^2 &= \|\vec{S}\|^2 = \vec{S}^2 = (\vec{S}_1 + \vec{S}_2)^2 = \vec{S}_1^2 + \vec{S}_2^2 + 2\vec{S}_1 \cdot \vec{S}_2 \\ &= \|\vec{S}_1\|^2 + \|\vec{S}_2\|^2 + 2\|\vec{S}_1\|\|\vec{S}_2\|\cos(\widehat{\vec{S}_1, \vec{S}_2}) \\ &= A_1^2 + A_2^2 + 2A_1A_2\cos\Delta\varphi, \end{aligned}$$

où $\Delta\varphi = \varphi_2 - \varphi_1$ est l'angle entre les vecteurs $\vec{S}_1$ et $\vec{S}_2$. On retrouve bien la formule (3.1).



Le calcul ci-dessus utilise différentes propriétés mathématiques du produit scalaire : bilinéarité, expression en fonction des normes et du cosinus de l'angle des vecteurs, égalité du carré scalaire et du carré de la norme. On trouvera plus de détails sur ces propriétés dans l'annexe mathématique.

c) Maximum et minimum d'amplitude

Les amplitudes A_1 et A_2 étant fixées, l'amplitude A est maximale lorsque $\cos\Delta\varphi = 1$ soit lorsque $\Delta\varphi = m \times 2\pi$ où m est un entier relatif.

L'amplitude du signal somme de deux signaux sinusoïdaux de même pulsation est **maximale** lorsque les signaux sont **en phase**.

La valeur maximale de A est :

$$A_{\max} = \sqrt{A_1^2 + A_2^2 + 2A_1A_2} = \sqrt{(A_1 + A_2)^2} \quad \text{soit} \quad A_{\max} = A_1 + A_2.$$

Ce résultat se voit bien sur la figure 3.1 lorsque les deux sinusoïdes sont en phase, ainsi que sur la figure 3.2 puisque dans ce cas les deux vecteurs de Fresnel sont colinéaires et de même sens.

A est minimale lorsque $\cos\Delta\varphi = -1$ soit lorsque $\Delta\varphi = \pi + m \times 2\pi$ où m est un entier relatif.

L'amplitude du signal somme de deux signaux sinusoïdaux de même pulsation est **minimale** lorsque les signaux sont **en opposition de phase**.

La valeur minimale de A est :

$$A_{\min} = \sqrt{A_1^2 + A_2^2 - 2A_1A_2} = \sqrt{(A_1 - A_2)^2} \quad \text{soit} \quad A_{\min} = |A_1 - A_2|.$$

Ce résultat se voit bien sur la figure 3.1 lorsque les deux sinusoïdes sont en opposition de phase, ainsi que sur la figure 3.2 puisque dans ce cas les deux vecteurs de Fresnel sont colinéaires et de sens opposés.

d) Cas de deux ondes de même amplitude

Dans le cas où les deux ondes ont des amplitudes égales, la formule (3.1) devient :

$$A = \sqrt{A_1^2 + A_1^2 + 2A_1^2 \cos \Delta\phi} = A_1 \sqrt{2 + 2 \cos \Delta\phi} = 2A_1 \left| \cos \frac{\Delta\phi}{2} \right|,$$

en utilisant la formule : $1 + \cos \alpha = 2 \cos^2 \frac{\alpha}{2}$ (voir appendice mathématique).

Les valeurs maximale et minimale de A sont dans ce cas :

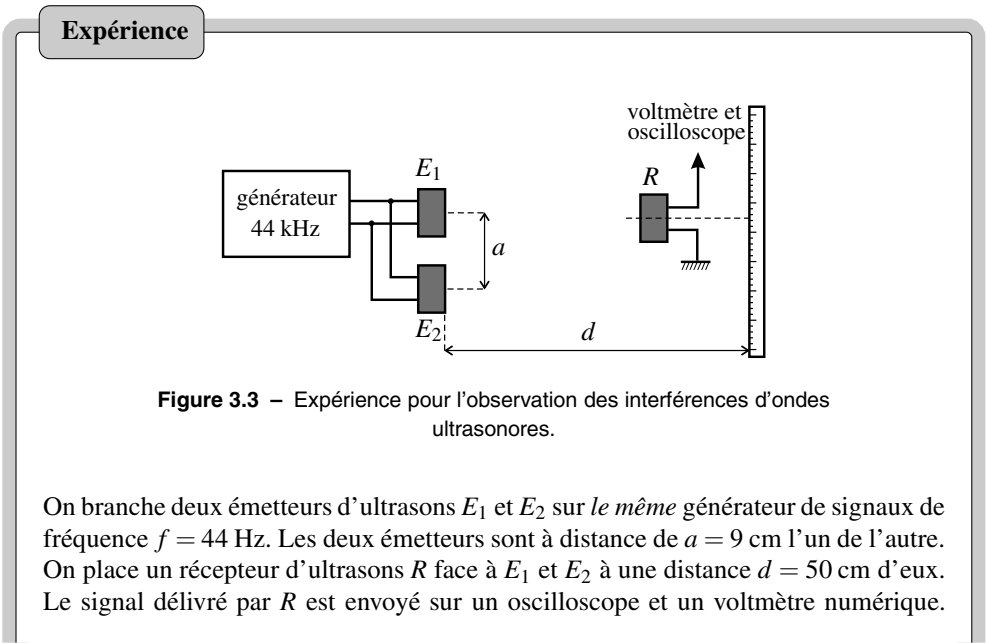
$$A_{\max} = 2A_1 \quad \text{et} \quad A_{\min} = 0.$$

Ainsi la superposition de deux ondes en opposition de phase peut donner un signal nul. Ceci est utilisé pour les isolations phoniques actives de certains casques. Un dispositif capte le bruit ambiant et envoie dans l'oreille un signal exactement en opposition de phase qui annule ce bruit.

1.2 Phénomène d'interférences

Lorsque deux ondes de même nature et de même fréquence parviennent en un point elles se superposent : leurs signaux s'additionnent. Cependant leur amplitudes ne s'additionnent pas. En effet, le calcul du paragraphe précédent a montré que l'amplitude du signal résultant dépend du déphasage $\Delta\phi = \phi_2 - \phi_1$ entre les deux signaux. Or ce déphasage varie d'un point à un autre. L'onde résultante a alors une amplitude modulée dans l'espace. C'est le **phénomène d'interférences** que l'on va étudier dans ce paragraphe.

a) Observation expérimentale.



CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

Lorsqu'on déplace R parallèlement à la droite reliant les deux émetteurs, l'amplitude du signal sinusoïdal observé sur l'oscilloscope varie, de même que l'indication du multimètre qui est égale à sa valeur efficace V_{eff} (c'est-à-dire son amplitude divisée par $\sqrt{2}$, voir le chapitre sur le régime harmonique).

Sur la figure 3.4 on voit la courbe donnant V_{eff} en fonction de la position x du récepteur mesurée le long d'une règle.

On observe une variation alternée de l'amplitude du signal qui est caractéristique du phénomène d'interférences.

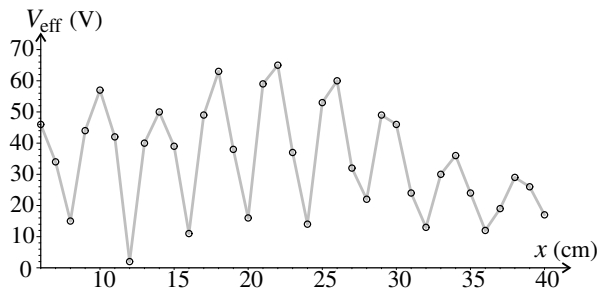


Figure 3.4 – Valeur efficace en fonction de la position du récepteur R .

b) Interprétation de l'expérience

Les signaux émis par E_1 et E_2 sont identiques puisqu'ils sont reliés au même générateur et peuvent s'écrire :

$$s_1(0, t) = s_2(0, t) = A_0 \cos\left(2\pi \frac{t}{T}\right),$$

en choisissant l'origine des temps pour que la phase initiale de ces signaux soit nulle.

On note d_1 la distance entre E_1 et R et d_2 la distance entre E_2 et R . Si l'émetteur E_1 était seul, R recevrait le signal :

$$s_1(d_1, t) = A_1 \cos\left(2\pi \left(\frac{t}{T} - \frac{d_1}{\lambda}\right)\right)$$

où A_1 est une amplitude inférieure à A_0 (on peut vérifier expérimentalement que cette amplitude est inversement proportionnelle à d_1). La phase initiale de ce signal est : $\varphi_1 = -\frac{2\pi d_1}{\lambda}$.

De même si E_2 était seul, R recevrait :

$$s_2(d_2, t) = A_2 \cos\left(2\pi \left(\frac{t}{T} - \frac{d_2}{\lambda}\right)\right).$$

La phase initiale de ce signal est : $\varphi_2 = -\frac{2\pi d_2}{\lambda}$.

Ainsi R reçoit les deux signaux qui sont déphasés de :

$$\Delta\varphi = \varphi_2 - \varphi_1 = 2\pi \frac{d_1 - d_2}{\lambda}.$$

Ce déphasage dépend de $d_1 - d_2$, donc de la position du récepteur par rapport aux deux émetteurs. L'amplitude détectée qui est déterminée par $\Delta\varphi$ selon la formule des interférences (3.1) en dépend donc aussi.

c) Interférences constructive et destructive

L'amplitude détectée est maximale lorsque les signaux des ondes des deux émetteurs sont en phase. On dit alors qu'il y a **interférence constructive**. C'est le cas si :

$$\Delta\varphi = 2m\pi \quad \text{soit} \quad d_1 - d_2 = m\lambda$$

où m est un entier relatif.

L'amplitude du signal détecté est minimale lorsque les signaux des ondes des deux émetteurs sont en opposition de phase. On dit alors qu'il y a **interférence destructive**. C'est le cas si :

$$\Delta\varphi = (2m + 1)\pi \quad \text{soit} \quad d_1 - d_2 = \left(m + \frac{1}{2}\right)\lambda$$

où m est un entier relatif.

d) Échelle angulaire du phénomène d'interférences

La figure 3.5 présente une simulation numérique des interférences entre deux sources ponctuelles. On constate que l'amplitude de l'onde résultante à grande distance dépend essentiellement de la position angulaire.

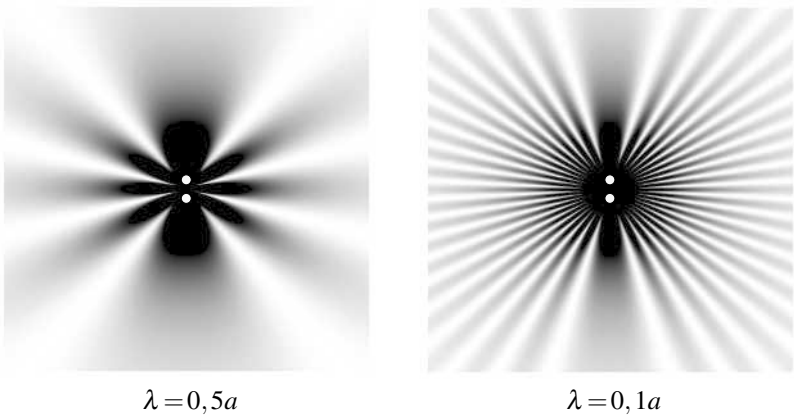


Figure 3.5 – Simulation d'interférences entre les ondes issues de deux sources ponctuelles (points blancs) en phase et distantes de a , pour deux valeurs différentes de la longueur d'onde. Le gris est d'autant plus foncé que l'amplitude de l'onde résultante est grande.

Quelle est l'échelle angulaire du phénomène d'interférences ? Pour la trouver on peut faire un calcul en s'appuyant sur la figure 3.6. Sur cette figure, O est le milieu du segment E_1E_2 et

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

(Ox) un axe orthogonal à E_1E_2 . On repère la position R du récepteur par $d = OR$ et l'angle θ entre (Ox) et le vecteur $\overrightarrow{OR}$.

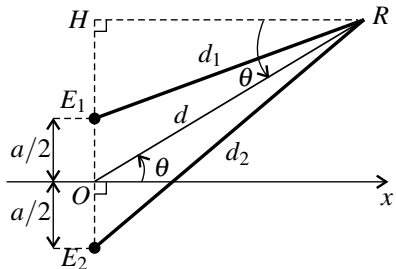


Figure 3.6 – Calcul de $d_2 - d_1$ en fonction de la position angulaire.

Soit H le projeté orthogonal de R sur la droite E_1E_2 . On retrouve l'angle θ dans le triangle rectangle ORH par la propriété des angles alterne-interne, ce qui permet de trouver que :

$$OH = OR \sin \theta = d \sin \theta.$$

On peut écrire ensuite, grâce au théorème de Pythagore :

$$\begin{aligned} d_2^2 - d_1^2 &= E_2R^2 - E_1R^2 = (E_2H^2 + HR^2) - (E_1H^2 + HR^2) \\ &= E_2H^2 - E_1H^2 = \left(d \sin \theta + \frac{a}{2}\right)^2 - \left(d \sin \theta - \frac{a}{2}\right)^2 = 2ad \sin \theta. \end{aligned}$$

D'autre part : $d_2^2 - d_1^2 = (d_2 + d_1)(d_2 - d_1)$. Et si la distance entre le récepteur et les émetteurs est très grande devant la distance entre les émetteurs, soit si $d \gg a$, on peut écrire : $d_1 + d_2 \simeq 2d$. Pour s'en convaincre on peut remarquer que l'égalité est rigoureuse si $\theta = \frac{\pi}{2}$

et que, pour $\theta = 0$, $d_1 + d_2 = 2\sqrt{d^2 + \frac{a^2}{4}} = 2d\sqrt{1 + \frac{a^2}{d^2}} \simeq 2d$ au premier ordre en $\frac{a}{d} \ll 1$. On a ainsi :

$$d_2^2 - d_1^2 \simeq 2d(d_2 - d_1)$$

En comparant les deux expressions de $d_2^2 - d_1^2$ on trouve :

$$d_2 - d_1 \simeq a \sin \theta.$$

On a interférence constructive si :

$$d_2 - d_1 = n\lambda \quad \text{soit si} \quad \sin \theta = n \frac{\lambda}{a},$$

avec n entier. Ainsi, l'échelle angulaire du phénomène d'interférence, à grande distance de deux émetteurs séparés d'une distance a est l'angle θ_i tel que :

$$\sin \theta_i = \frac{\lambda}{a}. \tag{3.2}$$

On remarquera que θ_i diminue si la longueur d'onde diminue, ce qui est bien visible sur la figure 3.5.

e) Échelle de longueur du phénomène d'interférences

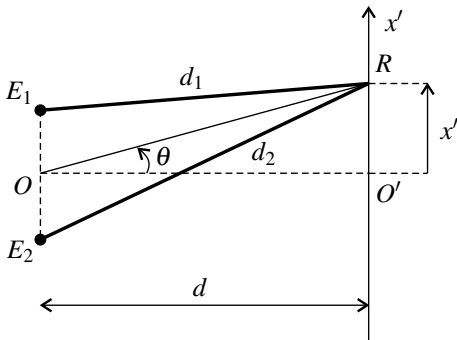


Figure 3.7 – Géométrie de l'expérience du paragraphe a)

Dans l'expérience du paragraphe a), on déplace le récepteur le long d'un axe ($O'x'$) parallèle à la droite E_1E_2 . Il apparaît sur la courbe tracée un distance caractéristique de variation de l'ordre de 5 cm. Peut-on retrouver théoriquement cette distance ?

Sur la figure, 3.7 on remarque que $\tan \theta = \frac{x'}{d}$. On peut donc écrire, pour x petit devant d , donc θ petit :

$$\theta \simeq \frac{x'}{d}.$$

Ainsi, la distance caractéristique du phénomène d'interférence sur l'axe ($O'x'$) est :

$$\ell_i = d\theta_i = \frac{\lambda d}{a}. \tag{3.3}$$

On notera que ℓ_i diminue avec l'écartement a des sources et augmente avec la longueur d'onde et lorsqu'on s'éloigne des sources.

Dans l'expérience précédente on avait $\lambda \sim 1$ cm, $a = 9$ cm et $d = 50$ cm ; la formule donne donc $\ell \sim 5$ cm ce qui est bien l'ordre de grandeur de la distance entre deux maxima successifs.



La valeur de ℓ_i ou θ_i suivant le cas est importante lorsqu'on souhaite observer les interférences. Il faut choisir les paramètres de l'expérience de façon que l'on puisse observer plusieurs maxima et minima d'amplitude en déplaçant le récepteur. Pour une longueur d'onde λ donnée, il faut bien choisir l'écartement a entre les deux émetteurs.

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

2 Ondes stationnaires et modes propres

2.1 Superposition de deux ondes progressives de même amplitude

Dans les deux paragraphes précédents on a considéré différents cas de superposition de signaux sinusoïdaux sans référence à leur dépendance spatiale (signaux pris en un point donné). Dans ce paragraphe, on va étudier la superposition, dans tout l'espace, de deux ondes de même fréquence et même amplitude, se propageant en sens inverse. Ces deux ondes progressives s'écrivent :

$$s_1(x, t) = A \cos(\omega t - kx + \varphi_1) \quad \text{et} \quad s_2(x, t) = A \cos(\omega t + kx + \varphi_2),$$

avec $k = \frac{\omega}{c}$. Elles se propagent le long de l'axe (Ox) en sens opposés.

En utilisant la formule de trigonométrie $\cos p + \cos q = 2 \cos \frac{p+q}{2} \cos \frac{p-q}{2}$ (voir appendice mathématique) qui transforme une somme de cosinus en un produit, on peut écrire l'onde qui résulte de la superposition de ces deux ondes, $s(x, t) = s_1(x, t) + s_2(x, t)$, sous la forme :

$$\begin{aligned} s(x, t) &= A \cos(\omega t - kx + \varphi_1) + A \cos(\omega t + kx + \varphi_2) \\ &= 2A \cos \left(kx + \frac{1}{2}(\varphi_2 - \varphi_1) \right) \times \cos \left(\omega t + \frac{1}{2}(\varphi_1 + \varphi_2) \right). \end{aligned}$$

Cette expression mathématique montre que l'onde $s(x, t)$ est de nature totalement différente des ondes $s_1(x, t)$ et $s_2(x, t)$. En effet, on n'a plus la combinaison linéaire $\omega t - kx$ ou $\omega t + kx$ qui traduit la propagation d'une onde, mais au contraire une *séparation* de la variable spatiale x et de la variable temporelle t . Une telle onde ne se propage pas, mais vibre sur place et elle est appelée **onde stationnaire**.

2.2 Onde stationnaire

a) Définition

Une **onde stationnaire** harmonique est une onde de la forme :

$$s(x, t) = 2A \cos(\omega t + \varphi) \cos(kx + \psi), \tag{3.4}$$

avec $k = \frac{\omega}{c}$, où c est la célérité des ondes progressives de même nature physique. Elle est égale à la superposition de deux ondes progressives sinusoïdales de pulsation ω se propageant en sens inverse le long de l'axe (Ox), ayant la même amplitude A .

La figure 3.8 représente $s(x, t)$ en fonction de x à différents instants. Il est intéressant de la comparer avec la figure 2.14 du chapitre précédent. On constate que cette onde ne se propage pas d'où le qualificatif de stationnaire.

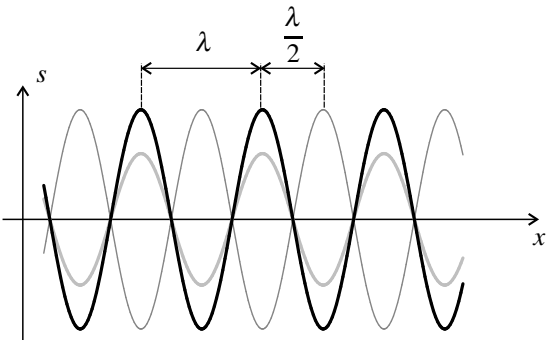


Figure 3.8 – Onde stationnaire à trois instants différents : t_1 (en noir), t_2 (en gris clair) et $t_3 = t_1 + \frac{T}{2}$ (en gris foncé, trait fin).

b) Nœuds et ventres de vibration

L’amplitude de l’onde (3.4) dépend de la position x (alors que ce n’est pas le cas pour une onde progressive) et elle est donnée par :

$$\mathcal{A}(x) = |2A \cos(kx + \psi)|.$$

On constate sur la figure 3.8 qu’il existe des points pour lesquels l’onde stationnaire est nulle à chaque instant. Ces points sont appelés les **nœuds de vibration**. Ce sont les points pour lesquels $\mathcal{A}(x) = 0$; leurs abscisses x vérifient :

$$\begin{aligned} \cos(kx + \psi) = 0 &\iff kx + \psi = \frac{\pi}{2} + n\pi, \text{ avec } n \text{ entier} \\ &\iff x = -\frac{\psi}{k} + \frac{\pi}{2k} + n\frac{\pi}{k}, \text{ avec } n \text{ entier} \\ &\iff x = -\frac{\psi}{k} + \frac{\lambda}{4} + n\frac{\lambda}{2}, \text{ avec } n \text{ entier}, \end{aligned}$$

puisque $\lambda = \frac{2\pi}{k}$. Les nœuds de vibrations sont donc régulièrement espacés et la distance entre deux nœuds consécutifs est $\frac{\lambda}{2}$.

Il existe aussi des points en lesquels l’amplitude $\mathcal{A}(x)$ de l’onde stationnaire est maximale. Ces points sont appelés les **ventres de vibration**. Ce sont les points pour lesquels $\mathcal{A}(x) = 2A$; leurs abscisses x vérifient :

$$\begin{aligned} \cos(kx + \psi) = \pm 1 &\iff kx + \psi = n\pi, \text{ avec } n \text{ entier} \\ &\iff x = -\frac{\psi}{k} + n\frac{\pi}{k}, \text{ avec } n \text{ entier} \\ &\iff x = -\frac{\psi}{k} + n\frac{\lambda}{2}, \text{ avec } n \text{ entier.} \end{aligned}$$

Les ventres de vibration sont donc régulièrement espacés et la distance entre deux ventres

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

consécutifs est $\frac{\lambda}{2}$. De plus, le milieu de deux nœuds de vibrations consécutifs est un ventre de vibration et vice-versa. Tout ceci est visible sur la figure 3.8.

Les **nœuds de vibration** sont les points où la vibration s’annule quelle que soit la date t .

Les **ventres de vibration**, sont les points où l’amplitude de la vibration est maximale. Les nœuds et les ventres sont disposés de manière alternée. La distance entre un nœud et un ventre consécutifs est $\frac{\lambda}{4}$. La distance entre deux nœuds ou de ventres consécutifs est $\frac{\lambda}{2}$.

L’existence de nœuds et de ventres de vibration est une **propriété caractéristique** des ondes stationnaires.

c) Phase initiale de l’onde stationnaire

L’onde stationnaire (3.4) peut s’écrire :

- si $\cos(kx + \psi) > 0$, $s(x, t) = \mathcal{A}(x) \cos(\omega t + \varphi)$,
- si $\cos(kx + \psi) < 0$, $s(x, t) = -\mathcal{A}(x) \cos(\omega t + \varphi) = \mathcal{A}(x) \cos(\omega t + \varphi + \pi)$.

Ainsi, la phase initiale de l’onde stationnaire ne prend que deux valeurs, φ et $\varphi + \pi$. Les vibrations en deux points sont soit en phase, soit en opposition de phase.

Entre deux nœuds, l’ensemble des points vibrent en phase. Lorsqu’on passe un nœud, la phase change de π . On peut observer tout ceci sur la figure 3.8.

2.3 Expérience de la corde de Melde

Le dispositif de la **corde de Melde** est un dispositif classique permettant d’observer des ondes stationnaires le long d’une corde. L’onde est ici observable à l’œil nu, puisqu’il s’agit d’un déplacement latéral de la corde de l’ordre de quelques centimètres. Toutefois, le phénomène étant oscillatoire et très rapide on utilise la technique de la stroboscopie pour le ralentir en apparence et pouvoir le regarder en direct.

a) Dispositif expérimental

On tend une corde horizontale d’environ 2 mètres de longueur entre la lame d’un vibreur et une poulie (figure 3.9). La tension de la corde est déterminée par la masse m que l’on accroche à l’extrémité pendante : si la poulie est sans frottement et que la masse est immobile elle est égale en norme à mg . En déplaçant la poulie on peut faire varier la longueur utile L de la corde entre, en gros, 1,5 m et 2 m.

Le vibreur est alimenté par un générateur de basses fréquences (GBF) et il impose à l’extrémité de la corde un mouvement sinusoïdal d’amplitude de l’ordre de quelques millimètres et dont la fréquence f se règle sur le GBF.

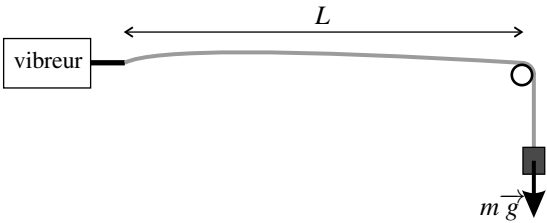


Figure 3.9 – Corde de Melde.

b) Phénomène de résonance

Observation expérimentale Lorsqu'on augmente progressivement f , on constate que la corde prend, pour certaines fréquences particulières, un mouvement d'amplitude très supérieure à l'amplitude du vibreur, pouvant valoir une dizaine de centimètres. On dit que la corde entre en **résonance** pour ces fréquences particulières.

L'aspect de la corde aux différentes résonances est donné sur la figure 3.10 :

- Pour la plus petite fréquence de résonance f_1 , la corde présente l'aspect d'un fuseau. Ce fuseau est la superposition, due à la persistance rétinienne, des différentes formes de la corde au cours du temps. L'amplitude de vibration est maximale au milieu de la corde et quasiment nulle aux deux extrémités (au niveau du vibreur l'amplitude n'est pas nulle, mais beaucoup plus petite que l'amplitude maximale). On observe ainsi un ventre de vibration au centre de la corde et des nœuds à ses extrémités *ce qui prouve que l'onde sur la corde est stationnaire*.
- La deuxième fréquence de résonance est $f_2 \simeq 2f_1$. La corde forme alors deux fuseaux de longueur $\frac{L}{2}$. On a deux ventres de vibrations et trois nœuds de vibration (un au milieu de la corde et deux aux extrémités).
- Les fréquences de résonance suivantes sont $f_3 \simeq 3f_1$, $f_4 \simeq 4f_1$ etc. L'onde forme de plus en plus de fuseaux : pour la fréquence f_n on a n fuseaux de longueur $\frac{L}{n}$, donc n ventres de vibrations et $n + 1$ nœuds (dont les deux extrémités).

Fréquences de résonance On peut exprimer les fréquences f_n de résonance de la corde. On sait que la distance entre deux nœuds de vibration consécutifs d'une onde stationnaire est égale à $\frac{\lambda}{2}$ où λ est la longueur d'onde. C'est la longueur d'un fuseau.

Lorsque la corde forme n fuseaux on a :

$$L = n \frac{\lambda}{2} \quad \text{soit} \quad \lambda = \lambda_n = \frac{2L}{n}.$$

(3.5)

On en déduit la fréquence f_n de l'onde :

$$f_n = \frac{c}{\lambda_n} = \frac{nc}{2L}.$$

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

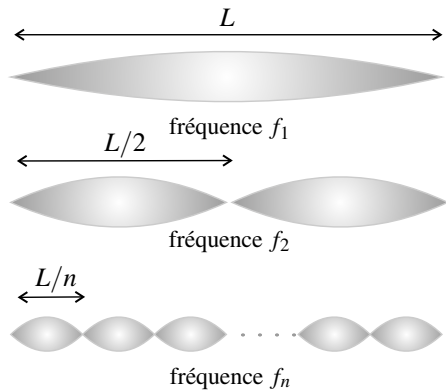


Figure 3.10 – Aspect de la corde pour les différentes fréquences de résonance.

On a ainsi trouvé l’expression théorique des fréquences de résonance de la corde :

$f_n = n f_1$

où

$f_1 = \frac{c}{2L}$

(3.6)



Ces expressions sont à connaître, mais le plus simple est de retenir l’aspect de la figure 3.10 et le raisonnement qui permet de les obtenir.

Vérification expérimentale On connaît l’expression de la célérité des ondes transversales sur la corde :

$$c = \sqrt{\frac{T}{\mu}},$$

où T est la tension de la corde et μ sa masse linéique (remarquer qu’il est intuitif que c augmente avec T et diminue lorsque μ soit l’inertie de la corde augmente). Dans l’expérience de Melde T est égale à mg où m est la masse qui est suspendue à l’extrémité de la corde. La masse linéique de la corde est facile à mesurer : il suffit de déterminer sa masse m' par pesée et de mesurer sa longueur totale L' (supérieure à la longueur utile L), alors : $\mu = \frac{m'}{L'}$. On peut alors calculer $f_1 = \frac{1}{2L} \sqrt{\frac{m}{m'}} g L'$ et comparer cette valeur théorique à la valeur mesurée.

c) Observation stroboscopique

L’observation stroboscopique permet de confirmer le caractère stationnaire de l’onde.

Un **stroboscope** est un dispositif fournissant un éclairage intermittent, composé de brefs éclairs régulièrement espacés d’une durée T_{strobo} qui est la période du stroboscope.

Lorsque la corde est éclairé par le stroboscope, on la voit à des instants $t_i = t_0 + i T_{\text{strobo}}$ où i est un entier . Si T_{strobo} est égale à la période de vibration de la corde, un point est toujours vu

dans la même position, donc la corde paraît fixe. Si T_{strobo} est légèrement plus grande, chaque point a « un peu » bougé entre deux éclairs ce qui fait qu'on voit un mouvement ralenti. Ceci est illustré sur la figure 3.11.

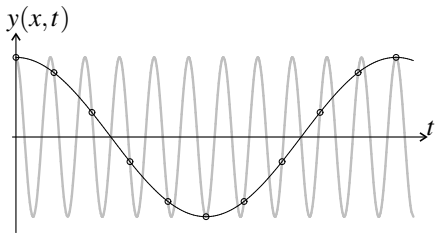


Figure 3.11 – Observation stroboscopique d'un mouvement sinusoïdal avec T_{strobo} légèrement supérieure à la période du mouvement : la courbe en gris représente le mouvement réel, les cercles donnent les instants t_i de visualisation, la courbe en noir représente le mouvement apparent.

En réglant le stroboscope pour avoir un mouvement ralenti (il suffit de le régler sur une fréquence d'éclairage légèrement différente de la fréquence sur laquelle le GBF est réglé) on observe que :

- les extrémités des fuseaux sont des points immobiles,
- à l'intérieur d'un fuseau, les points de la corde vibrent tous en phase,
- deux points situés dans deux fuseaux voisins vibrent en opposition de phase.

Ces observations correspondent bien à une onde stationnaire.

d) Interprétation physique

Pourquoi a-t-on une onde stationnaire sur la corde ? Pourquoi son amplitude est-elle très supérieure à l'amplitude de vibration du vibreur ?

Le vibreur envoie dans la corde une onde progressive qui se propage en direction de la poulie. Cette onde se réfléchit sur l'extrémité de la corde en contact avec la poulie, ce qui donne naissance à une deuxième onde se propageant en sens inverse. Cette deuxième onde se réfléchit sur l'autre extrémité de la corde liée au vibreur, donnant naissance à une troisième onde se propageant dans le même sens que la première, qui se réfléchit sur la poulie, donnant une quatrième onde et ainsi de suite...

L'onde stationnaire le long de la corde résulte de la superposition de ces ondes progressives se propageant dans les deux sens. Son amplitude peut être très supérieure à l'amplitude du vibreur puisqu'il y a *superposition d'un grand nombre d'ondes réfléchies successives*.

Pourquoi l'amplitude de vibration de la corde est-elle maximale pour les fréquences f_n ?

Ceci résulte de la condition d'interférence constructive. En effet, il y a entre deux ondes successives se propageant dans le même sens (la première et la troisième onde ci-dessus, ou la deuxième et la quatrième, par exemple) un décalage dans le temps correspondant à la durée d'un aller-retour le long de la corde soit $\tau = \frac{2L}{c}$. Ces ondes ont donc en tout point un

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

déphasage $\Delta\varphi$ donné par :

$$\Delta\varphi = \omega\tau = \frac{2\omega L}{c} = 2kL.$$

L'interférence entre ces deux ondes est constructive et l'amplitude résultant est maximale si $\Delta\varphi$ est un multiple entier de 2π , soit si :

$$2kL = 2\frac{2\pi}{\lambda}L = 2n\pi \iff \lambda = \frac{2L}{n}.$$

On retrouve ainsi la relation (3.5).

Il se propage le long de la corde un très grand nombre d'ondes provenant des réflexions aux deux extrémités.

Les fréquences de résonance sont celles pour lesquelles il y a *interférence constructive* entre les ondes se propageant dans le même sens.

La superposition entre les ondes se propageant dans les deux sens donne une onde stationnaire.

2.4 Modes propres

On s'intéresse maintenant aux vibrations d'une corde de longueur L finie, fixée entre deux points (il s'agit par exemple d'une corde de guitare). Ces vibrations sont des superpositions de vibrations sinusoïdales appelés **modes propres**.

a) Ondes stationnaires sur une corde de longueur L finie

Dans ce paragraphe on recherche les ondes stationnaires pouvant exister sur une corde de longueur L dont les extrémités sont fixes.

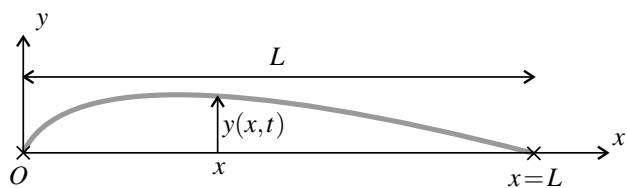


Figure 3.12 – Corde vibrante fixée à ses deux extrémités.

On prend un axe (Ox) le long de la corde, les deux extrémités fixes de la corde étant en $x = 0$ et $x = L$ (figure 3.12). On suppose que la corde reste comprise dans le plan (Oxy) et on appelle $y(x,t)$ le déplacement, supposé transversal, du point de la corde d'abscisse x . Les conditions physiques de fixation de la corde en ses extrémités se traduisent mathématiquement par les **conditions aux limites** :

$$y(0,t) = 0 \quad \text{et} \quad y(L,t) = 0.$$

On cherche une onde stationnaire $y(x,t)$ de la forme (3.4) :

$$y(x,t) = 2A \cos(kx + \psi) \cos(\omega t + \varphi),$$

avec $k = \frac{\omega}{c}$. Les conditions aux limites ci-dessus imposent que les deux extrémités de la corde soient des nœuds de vibration, soit :

$$\cos(\psi) = 0 \quad \text{et} \quad \cos(kL + \psi) = 0.$$

On ne réduit pas la généralité du calcul en supposant que ψ est compris entre $-\pi$ et $+\pi$. La première condition impose $\psi = \pm \frac{\pi}{2}$. Du fait que $\cos\left(kx - \frac{\pi}{2}\right) = -\cos\left(kx + \frac{\pi}{2}\right) = \sin(kx)$, le changement de $\psi = -\frac{\pi}{2}$ en $\psi = \frac{\pi}{2}$ revient à une inversion du signe de $\cos(kx + \psi)$ que l'on peut compenser en changeant φ en $\varphi + \pi$. Ainsi, sans perte de généralité, on prendra : $\psi = -\frac{\pi}{2}$ de sorte que : $\cos(kx + \psi) = \sin(kx)$.

La condition $\cos(kL + \psi) = \sin(kL) = 0$ devient alors $kL = n\pi$. Les valeurs possibles de k sont donc :

$$k_n = n \frac{\pi}{L}, \quad \text{où } n \text{ est un entier.}$$

On en déduit les valeurs possibles de la pulsation ω , en notant c la célérité des ondes se propageant le long de la corde :

$$\omega_n = ck_n = n \frac{\pi c}{L}, \quad \text{où } n \text{ est un entier.}$$

Les ondes stationnaires pouvant exister sur une corde de longueur L fixée en ses extrémités sont :

$$y(x,t) = A \sin\left(n \frac{\pi}{L} x\right) \cos\left(n \frac{\pi c}{L} t + \varphi\right) \quad (3.7)$$

où n est un entier, et A et φ des constantes quelconques.

Ce sont les **modes propres** de vibration de la corde. Les fréquences des modes propres, appelées **fréquences propres**, sont :

$$f_n = n \frac{c}{2L}.$$

Les longueurs d'ondes correspondantes,

$$\lambda_n = \frac{c}{f_n} = \frac{2L}{n},$$

sont les sous-multiples entiers de $2L$.

L'aspect de la corde pour les premiers modes propres est représenté sur la figure 3.13. Le mode propre d'ordre n a $n + 1$ nœuds de vibrations qui sont les deux extrémités de la corde et $n - 1$ autres nœuds.

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

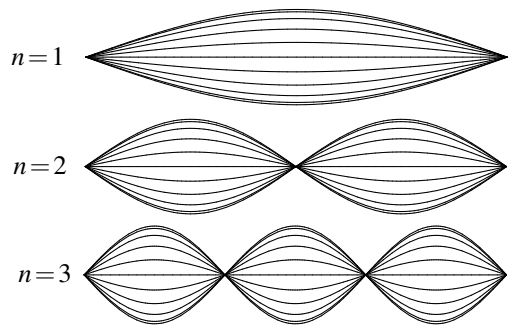


Figure 3.13 – Les trois premiers modes propres d'une corde fixée à ses deux extrémités.

Exemple

Les cordes d’une guitare sont en acier de masse volumique $\rho = 7,87 \cdot 10^3 \text{ kg}\cdot\text{m}^{-3}$. La corde la plus aiguë a un diamètre $\Phi = 3,00 \cdot 10^{-4} \text{ m}$ et une longueur $L = 64,0 \text{ cm}$. Elle est tendue avec une tension $T = 100 \text{ N}$.

La masse linéique de la corde est $\mu = \frac{1}{4} \rho \pi \Phi^2$ et la célérité des ondes sur cette corde est :

$$c = \sqrt{\frac{T}{\mu}} = \sqrt{\frac{T}{\rho \pi (\Phi^2/4)}} = 424 \text{ m}\cdot\text{s}^{-1}.$$

La longueur d’onde du mode fondamental de la corde est : $\lambda_1 = 2L = 1,28 \text{ m}$ et sa fréquence fondamentale :

$$f_1 = \frac{c}{\lambda_1} = \frac{424}{1,28} = 331 \text{ Hz}$$

(pour les lecteurs musiciens, il s’agit d’un mi3).

b) Mouvement général de la corde

Les modes propres sont les vibrations sinusoïdales de la corde. Le mouvement le plus général de la corde n’est pas sinusoïdal, mais c’est une superposition des différents modes propres.

Le mouvement le plus général de la corde est obtenu par **superposition linéaire** de tous ses modes propres, soit :

$$y(x,t) = \sum_{n=1}^{\infty} A_n \sin\left(n \frac{\pi}{L} x\right) \cos\left(n \frac{\pi c}{L} t + \varphi_n\right)$$

où les A_n et φ_n sont des constantes quelconques.

Comme les modes propres ont des fréquences toutes multiples de $f_1 = \frac{c}{2L}$, le mouvement est

périodique de fréquence f_1 .

Les amplitudes A_n et les phases initiales φ_n dépendent des conditions initiales du mouvement. On peut choisir ces conditions de manière à changer le poids des différents modes. Par exemple en imposant un déplacement initial grand là où un mode présente un nœud de vibration on défavorise ce mode qui aura une amplitude faible.

La vibration d’une corde d’instrument de musique est ainsi une superposition de vibrations sinusoïdales de toutes les fréquences propres de la corde. Le son émis contient aussi toutes ces fréquences. C’est un signal périodique non sinusoïdal dont le fondamental a la fréquence $f_1 = \frac{c}{2L}$ et les harmoniques les fréquences $f_n = n f_1$. Si L augmente, la fréquence f_1 diminue, et le son produit est plus grave. C’est un fait bien connu : les cordes d’une contrebasse sont plus longues que celles d’un violon.

SYNTHÈSE

SAVOIRS

- formule des interférences
- conditions pour avoir une amplitude maximale ou minimale
- formule d’une onde stationnaire
- distances entre nœuds et ventres de vibrations
- fréquence des modes propres en fonction de la longueur de la corde et de la célérité
- mouvement général de la corde comme superposition de modes propres

SAVOIR-FAIRE

- additionner deux signaux sinusoïdaux de même fréquence par le calcul analytique ou les vecteurs de Fresnel (MPSI)
- exprimer les conditions d’interférences constructrices ou destructrices
- décrire une onde stationnaire
- caractériser l’onde stationnaire par l’existence de nœuds et de ventres
- retrouver rapidement les fréquences des modes propres

MOTS-CLÉS

- | | | |
|-----------------------------|-----------------------------|------------------|
| • superposition | • interférence destructrice | • corde de Melde |
| • déphasage | • onde stationnaire | • mode propre |
| • interférence constructive | • nœud, ventre | |

S'ENTRAÎNER

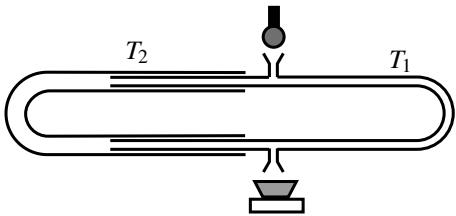
3.1 Utilisation des vecteurs de Fresnel (MPSI) (★)

1. Deux ondes $s_1(x, t) = A_0 \cos(\omega t - kx)$ et $s_2(x, t) = A_0 \cos(\omega t + kx)$ se superposent. Représenter les vecteurs de Fresnel correspondant à ces deux signaux en un point d'abscisse x . En déduire l'amplitude de l'onde somme en x . Déterminer graphiquement les valeurs de x pour lesquelles cette amplitude est maximale ou nulle.

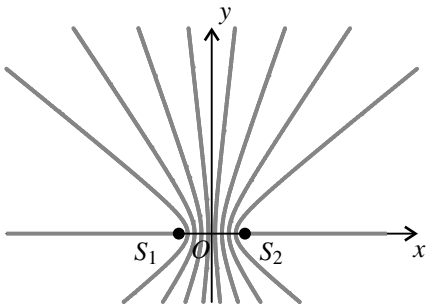
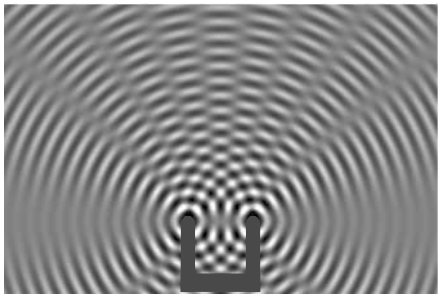
2. Utiliser les vecteurs de Fresnel pour trouver l'amplitude A et la phase initiale du signal somme des trois signaux suivants : $s_0(t) = A_0 \cos(\omega t + \varphi)$, $s_1(t) = rA_0 \cos(\omega t + \varphi + \Delta\varphi)$ et $s_2(t) = rA_0 \cos(\omega t + \varphi - \Delta\varphi)$ où $0 \leq r \leq 1$. Quelles sont les valeurs maximale et minimale de A ?

3.2 Mesure de la vitesse du son (★)

Le trombone de Koenig est un dispositif de laboratoire permettant de faire interférer deux ondes sonores ayant suivi des chemins différents. Le haut-parleur, alimenté par un générateur de basses fréquences, émet un son de fréquence $f = 1500$ Hz. On mesure le signal à la sortie avec un microphone branché sur un oscilloscope. En déplaçant la partie mobile T_2 on fait varier l'amplitude du signal observé. Elle passe deux fois de suite par une valeur minimale lorsqu'on déplace T_2 de $d = 11,5 \text{ cm} \pm 2 \text{ mm}$. Déterminer la valeur de la célérité du son dans l'air à 20°C , température à laquelle l'expérience est faite.



3.3 Interférences sur la cuve à ondes (★★)



La figure ci-dessus représente une cuve à ondes éclairée en éclairage stroboscopique. Deux pointes distantes de a frappent en même temps, à intervalles réguliers, la surface de l'eau, générant deux ondes qui interfèrent. La figure est claire là où la surface de l'eau est convexe

et foncée là où elle est concave. L'amplitude d'oscillation est plus faible là où la figure est moins contrastée.

1. On suppose pour simplifier que des ondes sinusoïdales partent des deux points S_1 et S_2 où les pointes frappent la surface. En notant λ la longueur d'onde, donner la condition pour que l'interférence en un point M situé aux distances d_1 et d_2 respectivement de S_1 et S_2 , soit destructrice. Cette condition fait intervenir un entier m .

2. Pour chaque entier m le lieu des points vérifiant cette condition est une courbe que l'on appelle dans la suite ligne de vibration minimale. Les lignes de vibration minimale sont représentées sur la figure de droite : ce sont des hyperboles (voir annexe mathématique).

a. Les parties $x < -\frac{a}{2}$ et $x > \frac{a}{2}$ de l'axe (Ox) sont des lignes de vibration minimale. En déduire un renseignement sur $\frac{a}{\lambda}$.

b. Sur le segment S_1S_2 quel est l'intervalle de variation de $d_2 - d_1$? Déduire de la figure la valeur de $\frac{a}{\lambda}$.

3. Expliquer pourquoi l'image est bien contrastée au voisinage de l'axe (Oy) .

4. (MPSI) On observe sur la figure que les zones claires et sombres sont alternées en opposition de phase de part et d'autre d'une ligne de vibration minimale. Le but de cette question est de comprendre pourquoi.

a. On suppose la phase initiale de chacune des ondes nulle à la source. Exprimer les phases Φ_1 et Φ_2 en M des ondes 1 et 2 provenant respectivement des sources S_1 et S_2 . En déduire la phase moyenne en M , $\Phi = \frac{1}{2}(\Phi_1 + \Phi_2)$, en fonction de f , t , d_1 , d_2 et λ . Quelle est la nature d'une courbe définie par la condition $\Phi = \text{constante}$?

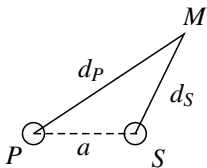
b. On se place en un point M d'une ligne de vibration minimale (L) . Représenter les vecteurs de Fresnel correspondant aux ondes 1 et 2 en M . Faire apparaître sur la figure la phase moyenne Φ .

c. Que devient cette figure si l'on se place en un point M' proche de M , du même côté de L que la source S_1 ? On supposera que Φ a la même valeur en M et M' . Même question pour un point M'' situé du même côté de (L) que la source S_2 .

d. Que peut-on dire des vibrations résultantes en M' et M'' ? On supposera pour simplifier que les deux ondes ont la même amplitude.

3.4 Contrôle actif du bruit en espace libre (★)

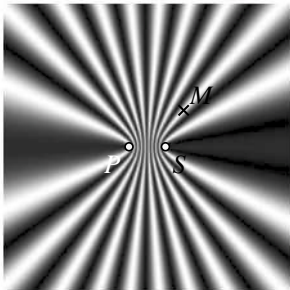
La méthode du contrôle actif du bruit consiste à émettre une onde sonore qui, superposée à l'onde sonore du bruit, l'annule par interférence destructrice. Pour modéliser la méthode on suppose que la source primaire de bruit P est ponctuelle et qu'elle émet une onde sinusoïdale de longueur d'onde λ . On crée une source sonore secondaire ponctuelle S qui est située à distance $PS = a$ de la source primaire et qui émet une onde de même longueur d'onde.



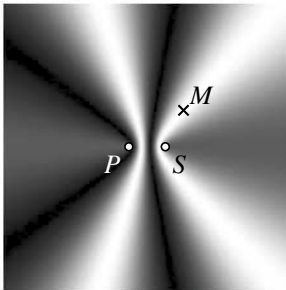
1. On souhaite annuler le bruit en un point M . On pose $PM = d_P$ et $SM = d_S$.

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

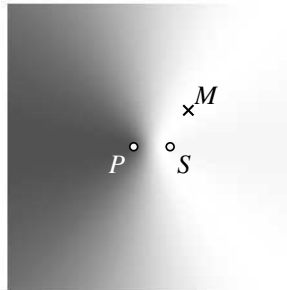
- a. Exprimer le déphasage $\Delta\varphi_0$ que la source secondaire doit présenter par rapport à la source primaire en fonction de λ , d_P , d_S et d'un entier m .
- b. L'amplitude de l'onde d'une source ponctuelle à distance d de la source est $A = \frac{\alpha}{d}$ où α est une constante. Quel doit être le rapport $\frac{\alpha_S}{\alpha_P}$ des constantes d'amplitude relatives aux deux sources ?



$\lambda = 0,2a$



$\lambda = a$



$\lambda = 5a$

2. Les figures ci-dessus obtenues par simulation visualisent l'amplitude de l'onde résultante dans le plan contenant P , S et M : le gris est d'autant plus foncé que l'amplitude de l'onde sonore est élevée. Commenter ce document, notamment l'influence de la longueur d'onde.

3.5 Résonances de la corde de Melde (★)

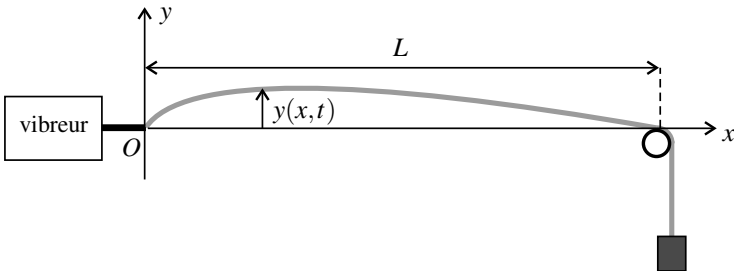
Une corde de Melde de longueur utile L est tendue entre un vibreur, en $x = 0$, et une poulie, en $x = L$. Le vibreur a un mouvement vertical sinusoïdal de pulsation ω :

$$y_{\text{vibreux}} = a_{\text{vibreux}} \cos(\omega t).$$

On cherche le déplacement de la corde sous la forme d'une onde stationnaire :

$$y(x,t) = A \cos(\omega t + \varphi) \cos(kx + \psi)$$

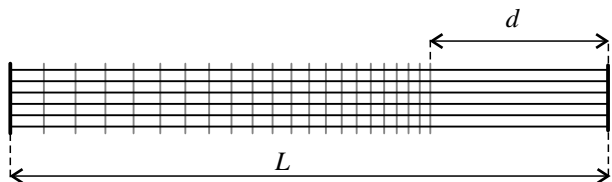
où A , φ et ψ sont des constantes à déterminer et où k est le vecteur d'onde associé à la pulsation ω .



1. En exploitant le fait que la corde est fixe en $x = L$ et liée au vibreur en $x = 0$, obtenir l'expression de $y(x, t)$ en fonction de a_{vibreux} , ω , k , L , x et t .
2. La corde étant en résonance avec un seul fuseau, l'amplitude maximale de vibration de la corde est environ égale à $10a$. Quelle relation y a-t-il entre la longueur d'onde λ de l'onde stationnaire et L ? On fera l'approximation : $\arcsin \frac{1}{10} \simeq \frac{1}{10}$.

3.6 Frettes d'une guitare (★)

Les frettes placées le long du manche d'une guitare permettent au musicien de modifier la hauteur du son produit par la corde. En pressant la corde contre une frette, il diminue sa longueur, provoquant une augmentation de la fréquence fondamentale de vibration de la corde.



1. Retrouver rapidement la fréquence de vibration fondamentale d'une corde de longueur L le long de laquelle les ondes se propagent à la célérité c .
2. La note monte d'un demi-ton lorsque la fréquence est multipliée par $2^{\frac{1}{12}}$. Pour cela, comment doit-on modifier la longueur de la corde ?
3. En plaçant le doigt sur les frettes successives on monte à chaque fois la note d'un demi-ton. Combien de frettes peut-il y avoir au maximum, sachant que la distance d entre la dernière frette et le point d'accrochage de la corde (voir figure) doit être supérieure à $0,25L$?

3.7 Anharmonicité d'une corde de piano (★)

On s'intéresse aux modes propres d'une corde de piano de longueur L , fixée en ses deux extrémités. La relation entre le vecteur d'onde et la pulsation d'une onde se propageant le long de cette corde est : $\omega = ck\sqrt{1 + \alpha k^2}$ où c et α dépendent de la section de la corde et de sa tension mais pas de sa longueur. Le coefficient α est dû à la raideur de la corde (il serait nul pour une corde parfaitement souple comme la corde de Melde).

1. Quelles sont les unités de c et α ?
2. Quelles sont les valeurs possibles de k pour une onde stationnaire existant sur cette corde ? Exprimer les fréquences correspondantes en fonction de c , α , L et d'un entier n .
3. Les cordes d'un piano de concert sont plus longues que les cordes d'un piano de salon. Pourquoi cela améliore-t-il la qualité musicale du son ?

3.8 Fréquences propres d'un tuyau sonore (★★)

La colonne d'air contenue dans un instrument à vent (flûte, clarinette...) ou dans un tuyau d'orgue vibre selon des modes propres correspondant à des conditions aux limites données. Dans une modélisation très simple on envisage deux types de conditions :

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

- si l'extrémité du tuyau est ouverte, la surpression acoustique est nulle à cette extrémité,
- si l'extrémité du tuyau est fermée, l'amplitude de variation de la surpression acoustique est maximale à cette extrémité.

1. On considère un tuyau de longueur L dans lequel la célérité des ondes sonores est c .

a. Déterminer les fréquences des modes propres du tuyau lorsque ses deux extrémités sont ouvertes. Représenter schématiquement la surpression dans le tuyau pour le troisième mode, les modes étant classés par fréquence croissante.

b. Même question si l'une des extrémités du tuyau est ouverte et l'autre fermée.

2. *Première application* : Les grandes orgues peuvent produire des notes très graves. Calculer la longueur d'onde d'un son de fréquence 34 Hz, correspondant au Do^0 , en prenant la valeur de la célérité du son à 0°C dans l'air soit $c = 331 \text{ m}\cdot\text{s}^{-1}$. Calculer la longueur minimale d'un tuyau produisant cette note.

3. *Deuxième application* : On peut modéliser très grossièrement une clarinette par un tube fermé au niveau de l'embouchure et ouvert à l'extrémité de l'instrument.

a. Expliquer pourquoi le son produit par une clarinette ne comporte que des harmoniques impairs.

b. L'instrument est muni d'une « clé de douzième » qui ouvre un trou situé à distance $\frac{L}{3}$ de l'embouchure. Lorsque ce trou est ouvert la surpression est nulle en ce point. Quelles sont dans ce cas les longueurs d'ondes des modes propres du tuyau ? Quel est l'effet de l'ouverture du trou sur la fréquence du son émis par l'instrument ?

APPROFONDIR

3.9 Interférences de deux ondes sonores (★)

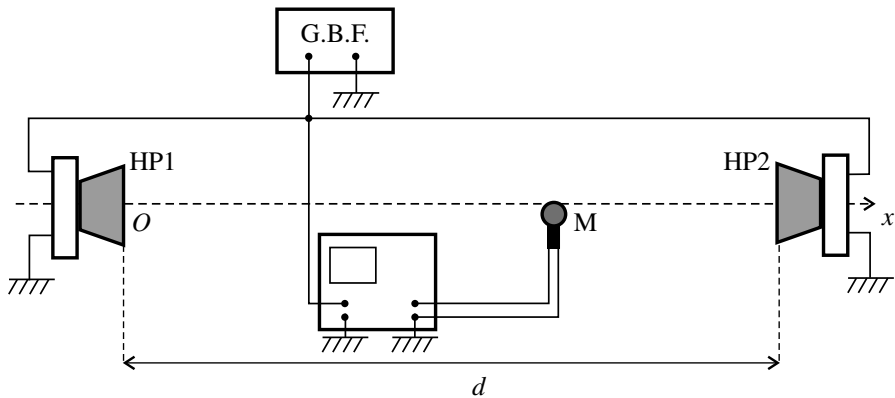
On s'intéresse aux interférences d'ondes sonores produites par deux haut-parleurs identiques, HP1 et HP2, placés face à face, à distance d l'un de l'autre, et alimentés par la même tension sinusoïdale de fréquence $f = 1250 \text{ Hz}$. L'axe (Ox) passe par les centres des haut-parleurs ; le centre de HP1 est en $x = 0$ et le centre de HP2 en $x = d$. Un microphone M de petite dimension peut être déplacé le long de (Ox) .

On envoie sur un oscilloscope le signal du générateur alimentant les haut-parleurs et la tension u délivrée par le micro, le premier signal servant de source de déclenchement.

1. Lorsqu'on déplace le micro autour de la position médiane entre le haut-parleur, soit $x = \frac{d}{2}$, on observe que :

- l'amplitude de la tension u varie et passe successivement par des maxima et minima quasiment nuls, l'écart entre deux positions successives pour lesquelles l'amplitude est minimale étant constant et valant : $e = 13,8 \text{ cm}$;
- la phase de la tension u est fixe entre deux points où l'amplitude s'annule et elle change de π quand on passe par un de ces points.

a. Quel phénomène ces observations évoquent-elles ?



b. Pour modéliser la situation on suppose que les surpressions acoustiques $p_1(x,t)$ et $p_2(x,t)$ ont des amplitudes constantes le long de l'axe (Ox) , toutes les deux égales à P_0 , et qu'elles ont toutes les deux la même phase initiale φ au départ des haut-parleurs. Écrire les expressions de $p_1(x,t)$ et $p_2(x,t)$ en fonction de P_0 , f , c célérité du son, φ , x et t .

c. Obtenir une expression de la surpression $p(x,t)$ résultant de la superposition de ces deux ondes qui explique les observations ci-dessus.

d. Calculer la vitesse du son dans les conditions de cette expérience.

2. Lorsqu'on éloigne le micro de la position médiane entre les haut-parleurs les observations sont différentes. L'amplitude de la tension u augmente et diminue périodiquement mais ne passe plus par zéro. Elle devient de plus en plus grande au fur et à mesure que le micro s'approche d'un haut-parleur. La phase de initiale de u dépend de la position x du micro : près de HP1, elle diminue avec x .

a. Expliquer pourquoi on ne peut plus faire l'hypothèse que les ondes venant des deux haut-parleurs ont la même amplitude.

b. Exprimer le déphasage $\Delta\varphi(x)$ entre ces deux ondes à l'abscisse x en fonction de d , c , f et x .

c. En appelant A_1 et A_2 leurs amplitudes, exprimer l'amplitude $A(x)$ de l'onde résultante.

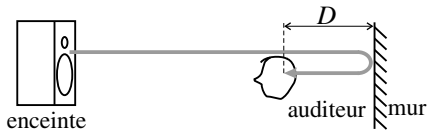
d. Montrer que si on approche le micro près du haut-parleur HP1 : $A(x) \simeq A_1 + A_2 \cos \Delta\varphi$.
On rappelle que : $\sqrt{1+\varepsilon} \simeq 1 + \frac{1}{2}\varepsilon$ pour $\varepsilon \ll 1$.

e. (MPSI) En s'appuyant sur une représentation de Fresnel, expliquer pourquoi la phase initiale de u diminue avec x près de HP1.

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

3.10 Écoute musicale et interférences (★★)

La qualité de l'écoute musicale que l'on obtient avec une chaîne hi-fi dépend de la manière dont les enceintes sont disposées par rapport à l'auditeur. On dit qu'il faut absolument éviter la configuration représentée sur la figure : présence d'un mur à distance D , trop courte derrière l'auditeur.



1. Comme représenté sur la figure, l'onde issue de l'enceinte se réfléchit sur le mur. On note $c = 342 \text{ m}\cdot\text{s}^{-1}$ la célérité du son dans l'air.

a. Exprimer le décalage temporel τ qui existe entre les deux ondes arrivant dans l'oreille de l'auditeur : onde arrivant directement et onde réfléchie.

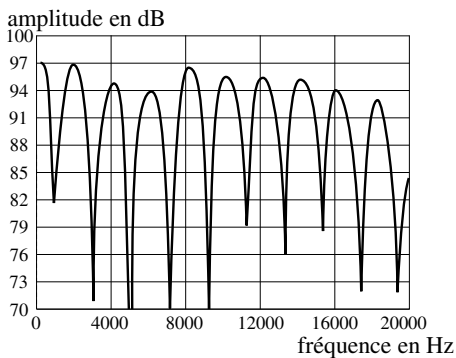
b. En déduire le déphasage $\Delta\phi$ de ces deux ondes supposées sinusoïdales de fréquence f . La réflexion sur le mur ne s'accompagne d'aucun déphasage pour la surpression acoustique, grandeur à laquelle l'oreille est sensible.

c. Expliquer pourquoi il y a un risque d'atténuation de l'amplitude de l'onde pour certaines fréquences. Exprimer ces fréquences en fonction d'un entier n . Quelle condition devrait vérifier D pour qu'aucune de ces fréquences ne soit dans le domaine audible ? Est-elle réalisable ?

d. Expliquer qualitativement pourquoi on évite l'effet nuisible en éloignant l'auditeur du mur.

2. La figure ci-contre donne le résultat d'une expérience dans laquelle on a placé un micro, sensible à la surpression, à une certaine distance D du mur, puis envoyé un signal de fréquence variable et d'amplitude constante A_0 . La courbe, d'allure très caractéristique, est appelée « courbe en peigne ».

L'amplitude en décibels se définit par la relation : $A_{\text{dB}} = 20 \log \frac{A}{A_{\text{ref}}}$, où A_{ref} est une amplitude de référence.

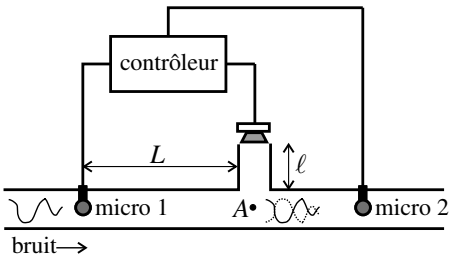


a. Lorsqu'il y a superposition de deux ondes de même amplitude A_0 , quelle est, en décibels, l'augmentation maximale de l'amplitude ? Que peut-on dire de $A_{0,\text{dB}}$ au vu de la courbe ?

b. Calculer numériquement la distance D .

3.11 Contrôle actif du bruit en conduite (★★)

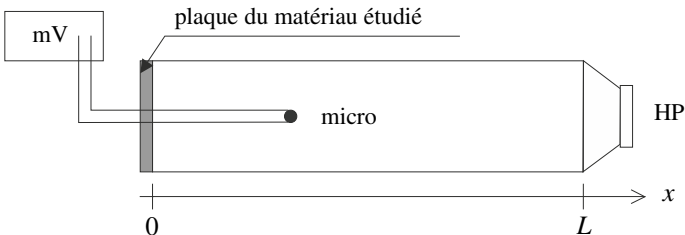
On s'intéresse à un système conçu pour l'élimination d'un bruit indésirable transporté par une conduite. Le bruit est détecté par un premier micro dont le signal est reçu par un contrôleur électronique. Le contrôleur, qui est le centre du système, envoie sur un haut-parleur la tension adéquate pour générer une onde de signal exactement opposé à celui du bruit de manière à ce que l'onde résultante au point A (voir figure) et en aval de A soit nulle.



1. Exprimer, en fonction de L , ℓ et la célérité c du son, le temps disponible pour le calcul du signal envoyé sur le haut-parleur.
2. On suppose le bruit sinusoïdal de pulsation ω . On appelle φ_1 la phase initiale du signal détecté par le micro 1 et φ_{HP} la phase initiale du signal émis par le haut-parleur. Exprimer, en fonction de ω , c , L et ℓ , la valeur que doit avoir $\Delta\varphi = \varphi_{HP} - \varphi_1$.
3. L'onde émise par le haut-parleur se propage dans la conduite dans les deux sens à partir de A. Expliquer l'utilité du micro 2.

3.12 Tube de Kundt (★★)

On considère un tuyau cylindrique, d'axe (Ox) de longueur $L = 1,45$ m, rempli d'air. À l'extrémité $x = L$ est placé un haut-parleur associé à un générateur de basses fréquences qui crée dans le tube une onde progressive se propageant dans le sens négatif de (Ox) . À l'autre extrémité ($x = 0$), l'expérimentateur place une plaque d'un matériau à étudier. Un microphone mobile, relié à un millivoltmètre, peut se déplacer à l'intérieur du tuyau sans perturber les phénomènes étudiés. Il fournit une tension variable, proportionnelle à la surpression de l'onde sonore à la position du micro. On mesure la valeur efficace V de cette tension à l'aide d'un millivoltmètre numérique. V est ainsi proportionnelle à l'amplitude de la surpression au point où se trouve le micro.



1. Donner l'expression de la surpression $p_{HP}(x,t)$ de l'onde émise par le haut-parleur, en notant A_0 son amplitude, f sa fréquence, c la célérité du son. On prendra la phase initiale de cette onde nulle en $x = 0$.
2. La plaque de matériau réfléchit cette onde et envoie dans le tube une onde d'amplitude égale à rA_0 avec $0 \leq r \leq 1$ et de phase initiale en $x = 0$ égale à φ . Donner l'expression de la

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

suppression $p_r(x, t)$ de cette onde.

3. On suppose qu'il n'y a pas d'autre onde que les deux ondes précédentes. Exprimer l'amplitude $A(x)$ de l'onde totale, $p(x, t) = p_{HP}(x, t) + p_r(x, t)$, en fonction de A_0 , r , φ , f , c et x .
4. On a réalisé des mesures avec une plaque de mousse. Les résultats sont consignés dans le tableau ci-dessous. On a noté les positions pour lesquelles l'indication du voltmètre V est minimale : x_1 est la première à partir de $x = 0$, x_i la $i^{\text{ème}}$. $V_{\min}$ et $V_{\max}$ sont les valeurs minimales et maximales de V .

f en Hz	x_1 en cm	x_i en cm	i	$V_{\min}$ en mV	$V_{\max}$ en mV
460	26,6	101,5	3	1,50	6,80
750	13,2	58,8	3	1,10	5,40
845	10,6	51,1	3	1,25	6,30
1016	8,0	41,7	3	0,80	4,05
1042	7,3	40,1	3	1,60	8,35
1185	5,5	49,0	4	0,80	3,95
1400	3,7	28,1	3	1,00	5,05

- a. On pose $\alpha = \frac{V_{\max}}{V_{\min}}$. Déterminer r en fonction de α .
- b. Déterminer l'expression de φ en fonction de x_1 et de la longueur d'onde λ .
- c. Calculer pour chaque fréquence les valeurs de r et de φ . Commenter ces résultats.

3.13 Ondes le long d'une corde de Melde (★★)

On considère une corde de Melde de longueur L . On interprète la vibration de la corde de la manière suivante : le vibreur émet une onde qui se propage en direction de la poulie où elle est réfléchiée ; cette onde réfléchiée se propage en direction du vibreur où elle est elle-même réfléchiée ; l'onde réfléchiée se propage en direction de la poulie où elle se réfléchit, et ainsi de suite.

L'axe (Ox) est parallèle à la corde au repos ; le vibreur est en $x = 0$ et la poulie en $x = L$ (voir figure de l'exercice 6 de ce chapitre).

Le vibreur émet une onde $s_0(x, t)$ telle que $s_0(0, t) = a_0 \cos(\omega t)$. La célérité des ondes sur la corde est c et on note $k = \frac{\omega}{c}$.

On fait les hypothèses simplificatrices suivantes :

- lorsqu'une onde incidente s_i arrive sur la poulie en $x = L$, l'onde réfléchiée s_r vérifie : $s_r(L, t) = -rs_i(L, t)$ où r est un coefficient compris entre 0 et 1 ;
- lorsqu'une onde incidente s_i arrive sur le vibreur en $x = 0$, l'onde réfléchiée s'_r vérifie : $s'_r(0, t) = -r's'_i(0, t)$ où r' est un coefficient compris entre 0 et 1.

1. Exprimer l'onde $s_0(x, t)$.
2. Exprimer l'onde $s_1(x, t)$ qui apparaît par réflexion de l'onde s_0 sur la poulie, puis l'onde $s_2(x, t)$ qui apparaît par réflexion de s_1 sur le vibreur, puis l'onde $s_3(x, t)$ qui apparaît par réflexion de s_2 sur la poulie.

3. À quelle condition les ondes s_0 et s_2 sont-elles en phase en tout point ? Que constate-t-on alors pour les ondes s_1 et s_2 ? La condition précédente est supposée réalisée dans la suite.

4. Justifier l'expression suivante de l'onde totale existant sur la corde :

$$s(x, t) = a_0 (1 + rr' + (rr')^2 + \dots + (rr')^n + \dots) \cos(\omega t - kx) \\ - ra_0 (1 + rr' + (rr')^2 + \dots + (rr')^n + \dots) \cos(\omega t + kx).$$

5. En quels points de la corde l'amplitude de la vibration est-elle maximale ? Exprimer l'amplitude maximale $A_{\max}$ en fonction de a , r et r' . On donne la formule : $\sum_{n=0}^{\infty} (rr')^n = \frac{1}{1 - rr'}$.

6. En quels points l'amplitude est-elle minimale ? Exprimer l'amplitude minimale $A_{\min}$.

7. Expérimentalement on trouve $\frac{A_{\min}}{a_0} \simeq 1$ et $\frac{A_{\max}}{a_0} \simeq 10$. Déterminer r et r' .

CORRIGÉS

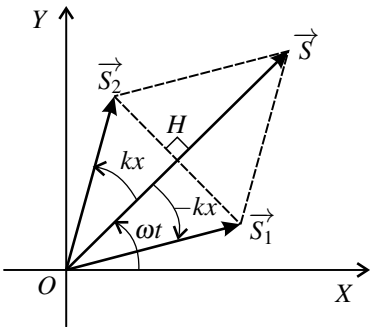
3.1 Utilisation des vecteurs de Fresnel (MPSI)

1. Les vecteurs de Fresnel associés à $s_1(x, t)$ et $s_2(x, t)$ sont représentés ci-contre, pour une valeur quelconque de x .

On a : $\|\vec{S}_1\| = \|\vec{S}_2\| = A_0$.

Le signal $s(x, t)$ est associé à $\vec{S} = \vec{S}_1 + \vec{S}_2$ et son amplitude est :

$$\begin{aligned}\|\vec{S}\| &= 2OH \\ &= 2\|\vec{S}_1\| |\cos(kx)| = 2A_0 |\cos(kx)|.\end{aligned}$$



La figure ci-dessous montre les cas où l'amplitude du signal somme est maximale ou nulle. Elle est maximale quand les vecteurs de Fresnel sont colinéaires, c'est-à-dire quand :

$$kx = 2n\pi \quad \text{ou} \quad kx = (2n + 1)\pi,$$

soit quand :

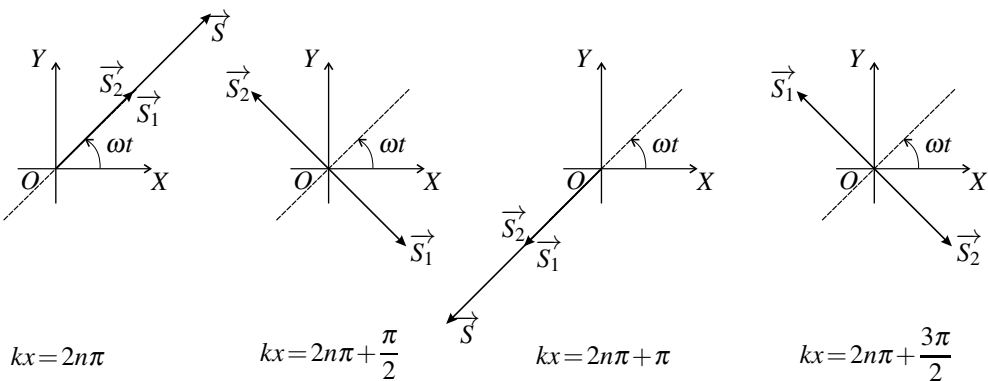
$$x = n\lambda \quad \text{ou} \quad x = (n + \frac{1}{2})\lambda.$$

Elle est nulle si :

$$kx = 2n\pi + \frac{\pi}{2} \quad \text{ou} \quad kx = 2n\pi + \frac{3\pi}{2},$$

soit si :

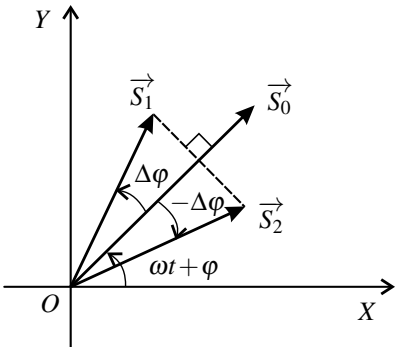
$$x = n\lambda + \frac{\lambda}{4} \quad \text{ou} \quad x = n\lambda + \frac{3\lambda}{4}$$



2. Les vecteurs de Fresnel associés aux signaux $s_0(x,t)$, $s_1(x,t)$ et $s_2(x,t)$ sont disposés comme indiqué sur la figure. Le signal somme $s = s_0 + s_1 + s_2$ a pour amplitude :

$$A = \|\vec{S}\| = A_0 + 2rA_0 \cos(\Delta\varphi).$$

La valeur maximale de A est $A_0(1 + 2r)$ et la valeur minimale $A_0|1 - 2r|$.



3.2 Mesure de la vitesse du son

Le micro reçoit deux ondes sonores qui sont passées par les deux tubes T_1 et T_2 . Ces deux ondes, de même fréquence, interfèrent. Lorsqu'on déplace le tube mobile T_2 on fait varier la longueur du trajet de l'onde qui passe par ce tube, modifiant donc le déphasage entre les ondes reçues par le micro. Plus précisément, lorsqu'on déplace T_2 de la distance d vers la gauche on allonge ce trajet de $2d$. Le retard temporel de cette onde par rapport à l'onde passée par T_1 augmente de $\frac{2d}{c}$, où c est la célérité du son, et son retard de phase augmente donc de $\omega \frac{2d}{c} = 2\pi f \frac{d}{c}$.

D'autre part, l'amplitude détectée par le micro est minimale lorsque les deux ondes sont en opposition de phase (interférence destructive), c'est-à-dire lorsque le retard de phase de l'onde passée par T_2 par rapport à l'onde passée par T_1 est de $2m\pi + \pi$ où m est un entier. Ainsi, en déplaçant T_2 d'une position pour laquelle l'amplitude détectée était minimale à la position suivante pour laquelle elle est de nouveau minimale, on a fait passer ce retard de phase de $2m\pi + \pi$ à $2(m + 1)\pi + \pi$. On a donc :

$$2\pi f \frac{2d}{c} = 2\pi \quad \text{soit} \quad c = 2fd = (345 \pm 6) \text{ m.s}^{-1}.$$

3.3 Interférences sur la cuve à ondes

1. Les ondes partant de S_1 et S_2 sont en phase puisque les deux points frappent en même temps la surface. Lorsqu'elles parviennent en M elles sont déphasées parce que les durées des parcours de S_1 à M et de S_2 à M ne sont pas égales. Le retard temporel en M de l'onde venant de S_2 par rapport à l'onde venant de S_1 est : $\tau = \frac{d_2}{c} - \frac{d_1}{c}$ où c est la célérité de la propagation. Le retard de phase en M de l'onde venant de S_2 par rapport à l'onde venant de S_1 est, en notant f la fréquence des ondes :

$$\Delta\varphi = 2\pi f \tau = \frac{2\pi f}{c}(d_2 - d_1) = \frac{2\pi}{\lambda}(d_2 - d_1).$$

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

L'interférence est destructrice si les deux ondes sont en opposition de phase, c'est-à-dire si $\Delta\varphi = 2m\pi + \pi$ où m est un entier. Il en est ainsi si :

$$d_2 - d_1 = \left(m + \frac{1}{2}\right)\lambda.$$

2. a. Pour M appartenant à (Ox) et en dehors du segment S_1S_2 , $d_2 - d_1 = \pm a$. L'interférence étant destructrice en ces points, on en déduit que $\frac{a}{\lambda}$ est un demi-entier.

b. Entre S_1 et S_2 on a $d_1 + d_2 = a$. Alors $d_2 - d_1 = 2d_2 - a$ qui varie entre $-a$ et $+a$. On compte 8 lignes de vibrations minimales passant entre les sources. Elles correspondent à $\frac{d_2 - d_1}{\lambda} = \pm\frac{1}{2}, \pm\frac{3}{2}, \pm\frac{5}{2}$ et $\pm\frac{7}{2}$. Donc $3,5\lambda < a \leq 4,5\lambda$. Ainsi : $a = 4,5\lambda$.

3. Si M est un point de l'axe (Oy) , on a : $d_1 = d_2$. Ainsi les vibrations en provenance des deux sources arrivent en phase en M , donc ont une interférence constructive. L'amplitude de vibration est ainsi maximale sur l'axe (Oy) ce qui explique le bon contraste de la figure au voisinage de cet axe.

4. a. Les phases des ondes 1 et 2 en M sont respectivement :

$$\Phi_1 = 2\pi ft - \frac{2\pi d_1}{\lambda} \quad \text{et} \quad \Phi_2 = 2\pi ft - \frac{2\pi d_2}{\lambda}.$$

La phase moyenne est : $\Phi = \frac{1}{2}(\Phi_1 + \Phi_2) = 2\pi ft - \frac{\pi}{\lambda}(d_1 + d_2)$.

Une courbe définie par la condition

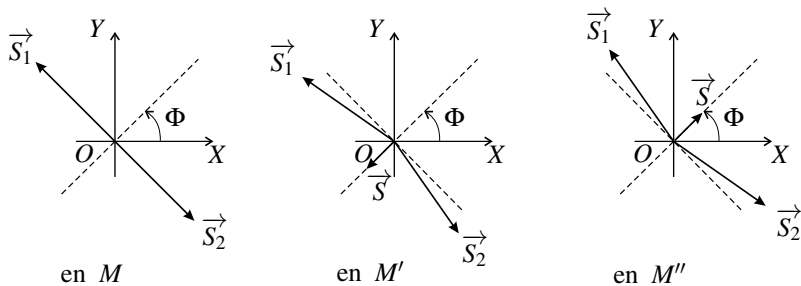
$$\Phi = \text{constante} \iff d_1 + d_2 = \text{constante}$$

est une ellipse dont les foyers sont S_1 et S_2 (voir annexe mathématique).

b. En M , point d'une ligne de vibration minimale, les vecteurs de Fresnel $\vec{S}_1$ et $\vec{S}_2$ relatifs aux deux ondes sont antiparallèles (figure ci-dessous, à gauche). Dans ce qui suit, on supposera que les angles entre les vecteurs $\vec{S}_1$ et $\vec{S}_2$ et l'axe des abscisses (OX) sont respectivement : $\Phi_1 = \Phi + \frac{\pi}{2}$ et $\Phi_2 = \Phi - \frac{\pi}{2}$. Ce pourrait être aussi l'inverse, mais cela ne changerait pas le raisonnement.

En M' , point voisin de M du côté de la source S_1 par rapport à la ligne de vibration minimale, le retard de phase de la vibration 2 par rapport à la vibration 1 est légèrement plus grand que π . L'angle entre l'axe (OX) et le vecteur $\vec{S}_1$ est supérieur à $\Phi + \frac{\pi}{2}$, tandis que l'angle entre (OX) et $\vec{S}_2$ est inférieur à $\Phi - \frac{\pi}{2}$ (figure ci-dessous au centre). Il se passe l'inverse en M'' , point voisin de M du côté de la source S_2 par rapport à la ligne de vibration minimale (figure ci-dessous à droite).

c. Les vibrations résultantes en M' et M'' sont en opposition de phase.



3.4 Contrôle actif du bruit en espace libre

1. a. La différence des temps de propagation entre les deux sources P et S et le point M fait que $\Delta\varphi$, déphasage en M de l'onde secondaire par rapport à l'onde primaire, est différent du déphasage $\Delta\varphi_0$ de la source S par rapport à la source P . Ainsi :

$$\Delta\varphi = \Delta\varphi_0 - \frac{2\pi}{\lambda}(d_S - d_P).$$

 On peut vérifier le signe en se disant que l'onde secondaire est plus en retard sur l'onde primaire au point M qu'au départ de sources si d_S est supérieure à d_P .

On souhaite une interférence destructrice en M soit $\Delta\varphi = \pi + 2m\pi$ où m est un entier. Il faut pour cela que :

$$\Delta\varphi_0 = \frac{2\pi}{\lambda}(d_S - d_P) + \pi + 2m\pi.$$

b. Pour une annulation du bruit en M , il faut de plus que les deux ondes aient la même amplitude en ce point, soit :

$$\frac{\alpha_S}{d_S} = \frac{\alpha_P}{d_P}, \text{ soit } \frac{\alpha_S}{\alpha_P} = \frac{d_S}{d_P}.$$

2. On remarque tout d'abord qu'annuler le bruit en un point ne veut pas dire annuler le bruit dans tout l'espace. Plus la longueur d'onde est grande, plus la zone autour de M où l'amplitude sonore est faible est étendue. Ceci était prévisible puisque l'échelle angulaire θ caractéristique du phénomène d'interférences est donnée par la relation : $\sin \theta = \frac{\lambda}{a}$. En conclusion, la méthode est efficace pour les plus faibles longueurs d'onde (sons graves). Elle est complémentaire avec les méthodes d'isolation phonique qui sont efficaces pour les plus petites longueurs d'onde.

3.5 Résonances de la corde de Melde

1. La corde est fixe en $x = L$ donc :

$$y(L, t) = 0 \Leftrightarrow A \cos(\omega t + \varphi) \cos(kL + \psi) = 0.$$

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

Ceci n'est vrai, à tout instant t , que si : $\cos(kL + \psi) = 0$, ce qui implique : $\psi = \frac{\pi}{2} - kL + m\pi$, où m est un entier. Dans ce cas :

$$\cos(kx + \psi) = \cos\left(k(x - L) + \frac{\pi}{2} + m\pi\right) = (-1)^{m+1} \sin(k(x - L)).$$

D'autre part, la corde est liée au vibreur en $x = 0$ donc :

$$y(0, t) = y_{\text{vibreux}}(t) \Leftrightarrow (-1)^m A \cos(\omega t + \varphi) \sin(kL) = a_{\text{vibreux}} \cos(\omega t).$$

On en déduit que : $\varphi = 0$ et $(-1)^m A = \frac{a_{\text{vibreux}}}{\sin(kL)}$. Finalement, la vibration de la corde est :

$$y(x, t) = \frac{a_{\text{vibreux}}}{\sin(kL)} \cos(\omega t) \sin(k(L - x)).$$

2. L'amplitude maximale de vibration sur la corde est : $10a \simeq \frac{a_{\text{vibreux}}}{\sin(kL)}$. On en déduit que : $\sin(kL) \simeq \frac{1}{10}$. La corde fait un seul fuseau donc L est proche de $\frac{\lambda}{2}$. Ainsi, kL est proche de π et on en déduit que $kL \simeq \pi - \arcsin \frac{1}{10} \simeq \pi - \frac{1}{10}$. Finalement :

$$L \simeq \frac{\lambda}{2} \left(1 - \frac{1}{20\pi}\right) \simeq 0,49\lambda.$$

3.6 Frettes d'une guitare

1. Dans le mode de vibration fondamental, de fréquence f_1 , la corde fait un seul fuseau avec un nœud de vibration à chaque extrémité. La distance entre deux nœuds de vibration étant égale à la demi-longueur d'onde, on a :

$$L = \frac{\lambda}{2} = \frac{c}{2f_1} \text{ soit } f_1 = \frac{c}{2L}.$$

2. Pour monter la note d'un demi-ton il faut multiplier sa longueur par $2^{-\frac{1}{12}}$.

3. Lorsqu'on met le doigt sur la $p^{\text{ième}}$ frette, la note est montée de p demi-tons, donc la longueur de la corde devient : $L_p = L \times 2^{-\frac{p}{12}}$. La condition imposée est :

$$L_p \geq \frac{L}{4} \text{ soit } 2^{-\frac{p}{12}} \geq \frac{1}{4} \text{ soit } p \leq 24.$$

Il y a 24 frettes au maximum. Pour le lecteur musicien : avec 24 frettes on couvre une étendue de 2 octaves avec une corde.

3.7 Anharmonicité d’une corde de piano

1. $[c] = \frac{[\omega]}{[k]} = \frac{T^{-1}}{L^{-1}} = LT^{-1}$. C’est une vitesse.

Le terme sous la racine carrée est sans dimension donc α a la dimension de $\frac{1}{k^2}$ soit $[\alpha] = T^2$.

2. Les extrémités de la corde sont fixes donc sont des nœuds de vibration. Par suite, la longueur L de la corde est un multiple de la demi-longueur d’onde :

$$L = n \frac{\lambda_n}{2} = n \frac{\pi}{k} \quad \text{soit} \quad k_n = \frac{n\pi}{L}.$$

En reportant dans la relation donnée par l’énoncé on trouve : $\omega_n = \frac{n\pi c}{L} \sqrt{1 + \alpha \frac{n^2 \pi^2}{L^2}}$. Finalement, la fréquence du mode propre d’ordre n est :

$$f_n = \frac{nc}{2L} \sqrt{1 + \alpha \frac{n^2 \pi^2}{L^2}}.$$

3. La raideur de la corde fait que les fréquences f_n ne sont pas des multiples entiers de la fréquence fondamentale f_1 . Ceci altère la qualité du son. Le terme lié à la raideur est divisé par L^2 : il diminue si les cordes sont plus longues.

3.8 Fréquences propres d’un tuyau sonore

1. a. Si les deux extrémités sont ouvertes, il y a un nœud de surpression acoustique aux deux extrémités. Or la distance entre deux nœuds est égale à un multiple entier de la demi-longueur d’onde. On a donc nécessairement :

$$L = n \frac{\lambda}{2} = n \frac{c}{2f} \quad \text{soit} \quad f = f_n = n \frac{c}{2L},$$

où n est un entier non nul. Le troisième mode propre ($n = 3$) est représentée sur la figure ci-dessous à gauche.



b. Il y a un nœud de surpression acoustique à l’extrémité ouverte du tuyau et un ventre de surpression acoustique à l’extrémité fermée. Or la distance entre un nœud et un ventre de vibration est égale au quart de la longueur d’onde plus un multiple entier de la demi-longueur d’onde. On a donc :

$$L = \frac{\lambda}{4} + n \frac{\lambda}{2} = \left(n + \frac{1}{2}\right) \frac{c}{2f} \quad \text{soit} \quad f = f'_n = (2n + 1) \frac{c}{4L},$$

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

où n est un entier. Le troisième mode propre ($n = 2$) est représenté sur la figure ci-dessus à droite.

2. $\lambda = \frac{c}{f_{D0^0}} = \frac{331}{32} = 10,3$ m. Pour produire cette note avec un tuyau de longueur minimale, il faut prendre un tuyau fermé à une extrémité et ouvert à l'autre tel que $f_1' = f_{D0^0}$. Ce tuyau a une longueur $L = \frac{\lambda}{4} = 2,6$ m.

3. a. Les modes propres d'un tuyau ouvert à une extrémité et fermé à l'autre sont les multiples impairs de la fréquence du mode fondamental. Quand l'air du tuyau vibre, sa vibration est une superposition linéaire des différents modes propres. Elle produit un son dont le spectre contient les fréquences de ces modes propres, c'est-à-dire un fondamental et seulement des harmoniques d'ordre impair.

b. Le trou impose un nœud de pression à distance $\frac{L}{3}$ de l'embouchure où il y a un ventre de pression. Ainsi, on a la condition :

$$\frac{L}{3} = (2n+1)\frac{\lambda}{4} \quad \text{soit} \quad f = f_n'' = (2n+1)\frac{3c}{4L},$$

où n est un entier. La fréquence fondamentale devient $f_1'' = \frac{3c}{4L} = 3f_1'$. L'ouverture du trou a pour effet de multiplier par 3 la fréquence du son émis par l'instrument. Pour le lecteur musicien : la multiplication de la fréquence par 3 correspond à un saut d'une octave plus une quinte.

3.9 Interférences de deux ondes sonores

1. a. Ces observations font penser aux propriétés de l'onde stationnaire.

b. L'onde 1 se propage dans le sens positif de l'axe (Ox) donc : $p_1(x, t) = p_1\left(0, t - \frac{x}{c}\right)$. Or : $p_1(0, t) = P_0 \cos(2\pi ft + \varphi)$ puisque la phase initiale de l'onde 1 au niveau de HP1 est φ . Il vient donc :

$$p_1(x, t) = P_0 \cos\left(2\pi ft - 2\pi \frac{x}{\lambda} + \varphi\right),$$

en notant $\lambda = \frac{c}{f}$ la longueur d'onde.

L'onde 2 se propage dans le sens négatif de l'axe (Ox) donc : $p_2(x, t) = p_2\left(d, t - \frac{d-x}{c}\right)$, avec $p_2(x, t) = P_0 \cos(2\pi ft + \varphi)$. Ainsi :

$$p_2(x, t) = P_0 \cos\left(2\pi ft + 2\pi \frac{x}{\lambda} - \frac{2\pi d}{\lambda} + \varphi\right).$$

c. La surpression totale est :

$$\begin{aligned} p(x, t) &= P_0 \cos\left(2\pi ft - 2\pi \frac{x}{\lambda} + \varphi\right) + P_0 \cos\left(2\pi ft + 2\pi \frac{x}{\lambda} - \frac{2\pi d}{\lambda} + \varphi\right) \\ &= 2P_0 \cos\left(2\pi \frac{x}{\lambda}\right) \cos\left(2\pi ft - \frac{\pi d}{\lambda} + \varphi\right), \end{aligned}$$

en utilisant la formule de trigonométrie de transformation d'une somme de cosinus en produit (voir annexe mathématique). Il s'agit d'un onde stationnaire, en accord avec les observations expérimentales.

d. Les points d'amplitude minimale sont les nœuds dont on sait qu'ils sont distants d'une demi-longueur d'onde. Ainsi : $e = \frac{\lambda}{2} = \frac{c}{2f}$, donc : $c = 2ef = 2 \times 13,8 \cdot 10^{-2} \times 1250 = 345 \text{ m} \cdot \text{s}^{-1}$.

2. a. L'amplitude de l'onde émise par un haut-parleur n'est pas constante, mais elle diminue lorsqu'on s'en éloigne. L'approximation d'une amplitude constante et identique pour les deux haut-parleurs n'est plus applicable quand on s'approche d'un des deux haut-parleurs.

b. Les ondes sont en phase quand elles partent des haut-parleurs. Pour parvenir au point d'abscisse x , la première onde met le temps $\frac{x}{c}$ et la deuxième le temps $\frac{d-x}{c}$. Elles ont donc un décalage temporel en ce point, la deuxième étant en avance par rapport à la première de $\tau = \frac{x}{c} - \frac{d-x}{c} = \frac{2x-d}{c}$. La deuxième onde présente donc par rapport à la première le déphasage :

$$\Delta\varphi(x) = 2\pi f\tau = \frac{2\pi}{\lambda}(2x-d).$$

c. D'après la formule des interférences à deux ondes, l'amplitude de l'onde résultante est :

$$A(x) = (A_1^2 + A_2^2 + 2A_1A_2 \cos(\Delta\varphi(x)))^{\frac{1}{2}}.$$

d. Près de HP, $A_1 \gg A_2$. Donc :

$$A(x) = A_1 \left(1 + 2\frac{A_2}{A_1} \cos(\Delta\varphi(x)) + \frac{A_2^2}{A_1^2} \right)^{\frac{1}{2}} \simeq A_1 \left(1 + \frac{A_2}{A_1} \cos(\Delta\varphi(x)) \right),$$

en faisant un développement limité du premier ordre en $\frac{A_2}{A_1} \ll 1$.



Dans un calcul au premier ordre, les termes du deuxième ordre, comme le terme $\frac{A_2^2}{A_1^2}$ ci-dessus, peuvent être purement et simplement supprimés.

e. Le vecteur de Fresnel du signal somme a presque la direction du vecteur de Fresnel associé à p_1 , donc le signal somme a pratiquement la phase initiale de p_1 qui décroît avec x .

3.10 Écoute musicale et interférences

1. a. L'onde réfléchi parcourt en plus deux fois la distance D entre l'auditeur et le mur donc : $\tau = \frac{2D}{c}$.

b. C'est la seule cause de décalage entre les deux ondes puisque la réflexion sur le mur ne s'accompagne d'aucun déphasage. L'onde réfléchi présente donc par rapport à l'onde directe le déphasage : $\Delta\varphi = 2\pi f\tau = \frac{4\pi fD}{c}$.

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

c. Il peut y avoir atténuation de l'amplitude si les deux ondes sont en opposition de phase et ont une interférence destructive. C'est le cas si :

$$\Delta\varphi = (2n+1)\pi \quad \text{soit} \quad f = (2n+1) \frac{c}{4D},$$

où n est un entier.

Le domaine audible s'étend de 20 Hz à 20 kHz. Aucune des fréquences précédentes ne se trouve dans le domaine audible si : $\frac{c}{4D} > 20 \text{ kHz}$. Il faut pour cela que $D < \frac{342}{4 \times 20} = 4,3 \text{ mm}$. Il faut que la tête de l'auditeur frôle le mur !

d. Pour D suffisamment grand, l'onde réfléchie par le mur a une amplitude très faible devant l'onde directe.

2. a. L'amplitude est maximale dans le cas d'une interférence constructive et elle vaut $2A_0$. Sa valeur en décibels est :

$$A_{\text{dB}} = 20 \log \frac{2A_0}{A_{\text{ref}}} = 20 \log \frac{A_0}{A_{\text{ref}}} + 20 \log 2 = A_{0,\text{dB}} + 6.$$

Sur la courbe l'amplitude maximale observée est 97dB, donc $A_{0,\text{dB}} \geq 91\text{dB}$.

b. L'écart moyen entre les fréquences pour lesquelles l'amplitude mesurée est minimale est $\Delta f = 2049 \text{ Hz}$. D'après ce qui précède, $\Delta f = \frac{c}{4D}$, soit :

$$D = \frac{c}{4\Delta f} = \frac{342}{4 \times 20149} = 4,2 \cdot 10^{-2} \text{ m}.$$

3.11 Contrôle actif du bruit en conduite

1. Entre l'instant où le signal est détecté par le micro 1 et l'instant où ce signal passe en A s'écoule un temps égal à $\frac{L}{c}$. Pendant ce temps il faut que le contrôleur calcule et produise le signal qu'il envoie dans le haut-parleur et que ce signal se propage jusqu'à A, ce qui prend le temps $\frac{\ell}{c}$. Le temps disponible pour le calcul est donc : $\frac{L-\ell}{c}$.

2. La phase initiale du signal de bruit arrivant en A est $\varphi_{\text{bruit}} = \varphi_1 - \omega \frac{L}{c}$.

La phase initiale du signal de correction arrivant en A est : $\varphi_{\text{corr}} = \varphi_{\text{HP}} - \omega \frac{\ell}{c}$.

Pour avoir une interférence destructive il faut que $\varphi_{\text{corr}} = \varphi_{\text{bruit}} + \pi$, soit que :

$$\Delta\varphi = \varphi_{\text{HP}} - \varphi_1 = \omega \frac{\ell - L}{c} + \pi.$$

3. Le micro 1 capte un signal qui est la superposition du bruit et du signal émis par le haut-parleur se propageant à partir de A vers l'amont. Le micro 2 donne un contrôle du résultat et permet la détermination du meilleur signal de correction.

3.12 Tube de Kundt

- 1. $p_{HP}(x,t) = A_0 \cos(2\pi ft + kx)$ puisque cette onde se propage dans le sens négatif de (Ox) .
- 2. $p_r(x,t) = rA_0 \cos(2\pi ft - kx + \varphi)$.
- 3. Le déphasage des deux ondes à l'abscisse x est : $\Delta\varphi(x) = \varphi - 2kx$. D'après la formule des interférences à deux ondes, l'amplitude en x de l'onde résultante est :

$$A(x) = (A_0^2 + r^2 A_0^2 + 2rA_0^2 \cos(\Delta\varphi(x)))^{\frac{1}{2}} = A_0 (1 + r^2 + 2r \cos(\varphi - 2kx))^{\frac{1}{2}}.$$

- 4. a. $\alpha = \frac{V_{\max}}{V_{\min}} = \frac{A_{\max}}{A_{\min}} = \frac{1+r}{1-r}$. Cette relation s'inverse en : $r = \frac{\alpha - 1}{\alpha + 1}$.
- b. Le premier minimum correspond à : $2kx - \varphi = \pi$, soit : $\varphi = \pi \left(\frac{4x_1}{\lambda} - 1 \right)$.
- c. L'application numérique donne :

f en Hz	460	750	845	1016	1042	1185	1400
r	0,639	0,662	0,669	0,670	0,678	0,663	0,669
ϕ	$-0,64\pi$	$-0,71\pi$	$-0,74\pi$	$-0,76\pi$	$-0,78\pi$	$-0,81\pi$	$-0,85\pi$

Le module du coefficient de réflexion dépend peu de la fréquence alors que sa phase augmente en valeur absolue avec la fréquence.

3.13 Ondes le long d'une corde de Melde

- 1. L'onde s_0 se propage dans le sens de (Ox) avec la célérité c donc :

$$s_0(x,t) = s_0 \left(0, t - \frac{x}{c} \right) = a_0 \cos \left(\omega \left(t - \frac{x}{c} \right) \right) = a_0 \cos(\omega t - kx).$$

- 2. L'onde s_1 provient de l'onde s_0 par réflexion sur la poulie donc : $s_1(L,t) = -rs_0(L,t) = -ra_0 \cos(\omega t - kL)$. De plus, s_1 se propage dans le sens inverse de (Ox) avec la célérité c , donc :

$$s_1(x,t) = s_1 \left(L, t - \frac{L-x}{c} \right) = -ra_0 \cos \left(\omega \left(t - \frac{L-x}{c} \right) - kL \right) = -ra_0 \cos(\omega t + kx - 2kL).$$

On trouve de la même manière : $s_2(t) = -r's_1(0,t) = rr'a_0 \cos(\omega t - 2kL)$, puis :

$$s_2(x,t) = s_2 \left(L, t - \frac{x}{c} \right) = rr'a_0 \cos(\omega t - kx - 2kL).$$

On en déduit : $s_3(L,t) = -rs_2(L,t) = -r^2r'a_0 \cos(\omega t - 3kL)$, puis :

$$s_3(x,t) = s_3 \left(L, t - \frac{L-x}{c} \right) = -r^2r'a_0 \cos(\omega t + kx - 4kL).$$

- 3. La phase initiale de s_0 au point d'abscisse x est égale à $-kx$; la phase initiale de s_2 au même point est égale à $-kx - 2kL$. Les deux ondes sont en phase en tout point à condition que :

$$2kL = 2n\pi, \text{ où } n \text{ est un entier.}$$

Si cette condition est vérifiée, les ondes s_1 et s_3 sont elles aussi en phase en tout point.

CHAPITRE 3 – SUPERPOSITIONS DE SIGNAUX

4. Par réflexion continue sur les extrémités de la corde, il existe une infinité d'ondes se propageant dans les deux sens. Les ondes se propageant dans un sens donné sont toutes en phase.

5. On peut réécrire $s(x, t)$ en utilisant la formule de sommation fournie :

$$s(x, t) = \frac{a_0}{1 - rr'} \cos(\omega t - kx) - \frac{ra_0}{1 - rr'} \cos(\omega t + ks).$$

L'amplitude est maximale en un point où les deux ondes sont en interférence constructive, c'est-à-dire aux points tels que : $kx = -kx + \pi + 2p\pi$ (attention au signe – devant le deuxième terme) où p est un entier soit :

$$x = \frac{(2p+1)\pi}{2k} = \frac{2p+1}{2n}L.$$

L'entier p est compris entre 0 et $n-1$ puisque x est compris entre 0 et L . Ces points sont les ventres de vibration.

La valeur maximale de l'amplitude est : $A_{\max} = \frac{a_0}{1 - rr'} + \frac{ra_0}{1 - rr'} = a_0 \frac{1+r}{1 - rr'}$.

6. L'amplitude est minimale en un point où les deux ondes sont en interférence destructive, c'est-à-dire aux points tels que : $kx = -kx + 2p\pi$ où p est un entier soit :

$$x = \frac{p\pi}{k} = \frac{p}{n}L.$$

L'entier p est compris entre 0 et n puisque x est compris entre 0 et L . Ces points sont les nœuds de vibration.

La valeur minimale de l'amplitude est : $A_{\min} = \frac{a_0}{1 - rr'} - \frac{ra_0}{1 - rr'} = a_0 \frac{1-r}{1 - rr'}$.

7. On a : $\frac{1-r}{1+rr'} \simeq 1$ et $\frac{1+r}{1-rr'} \simeq 10$. C'est un système de deux équations à deux inconnues dont la solution est : $r \simeq 0,8$ et $r' \simeq 1$.

Onde lumineuse

4

Ce chapitre est consacré aux ondes associées aux phénomènes lumineux. Après avoir donné les ordres de grandeurs relatifs aux ondes lumineuses, on caractérise les différentes sources de lumière par leur spectre. On présente ensuite dans le cas de la lumière le phénomène de diffraction qui est commun à tous les types d'onde. Enfin, le caractère particulier des ondes lumineuses, ondes vectorielles, est mis en valeur par l'étude de la polarisation.

1 L'onde lumineuse

1.1 Existence et nature de l'onde lumineuse

La nature de la lumière a fait l'objet d'une controverse scientifique aux $XVII^{\text{ème}}$ et $XVIII^{\text{ème}}$ siècles entre les partisans de deux théories :

- la théorie corpusculaire, initiée par Newton, selon laquelle la lumière est faite de particules émises par les corps lumineux,
- la théorie ondulatoire, initiée par Hooke et développée par Huygens, selon laquelle la lumière est une onde.

Au $XIX^{\text{ème}}$ siècle, la découverte des phénomènes d'interférences lumineuses, par Young et Fresnel notamment, prouva de manière indéniable l'existence d'une onde lumineuse. Toutefois, on ne connaissait pas la nature physique de la grandeur qui se propage avec cette onde. La question fut résolue à la fin du $XIX^{\text{ème}}$, quand Maxwell établit théoriquement l'existence d'ondes électromagnétiques se propageant dans le vide à une vitesse égale à la vitesse de la lumière. Il est depuis établi que la lumière est une **onde électromagnétique**.

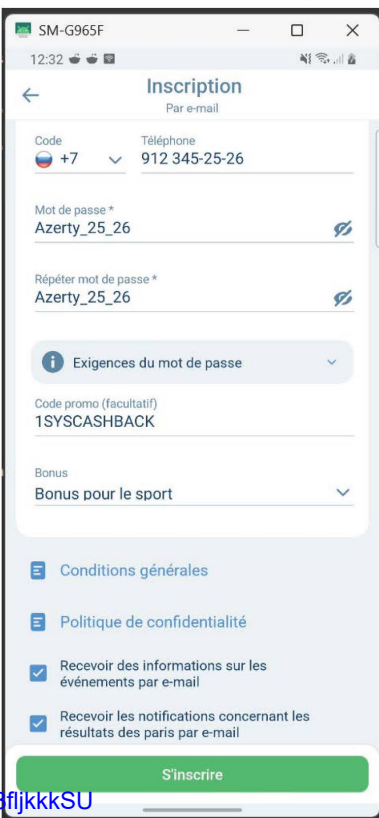
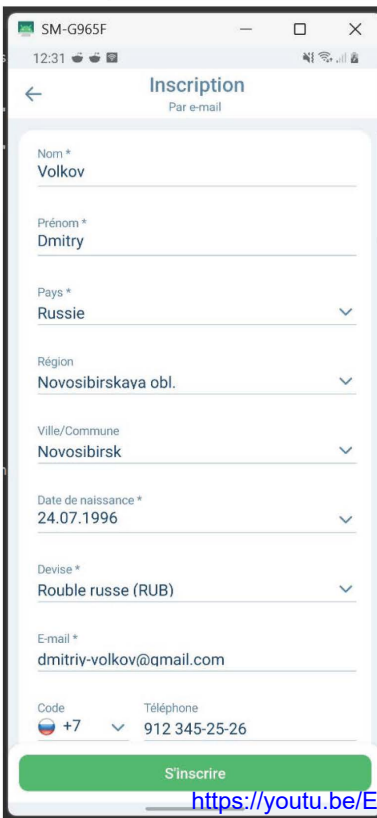
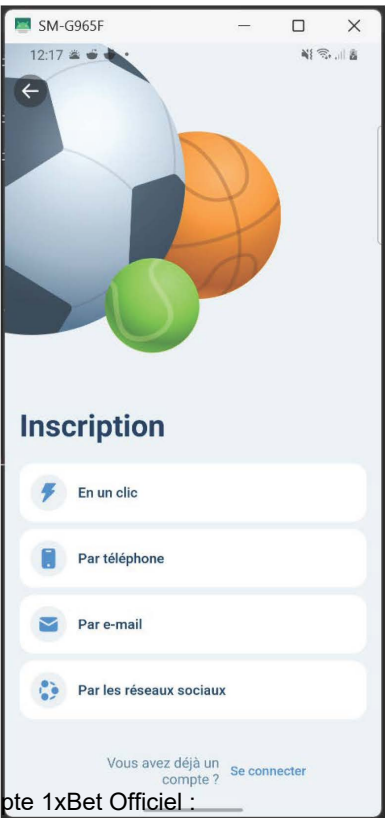
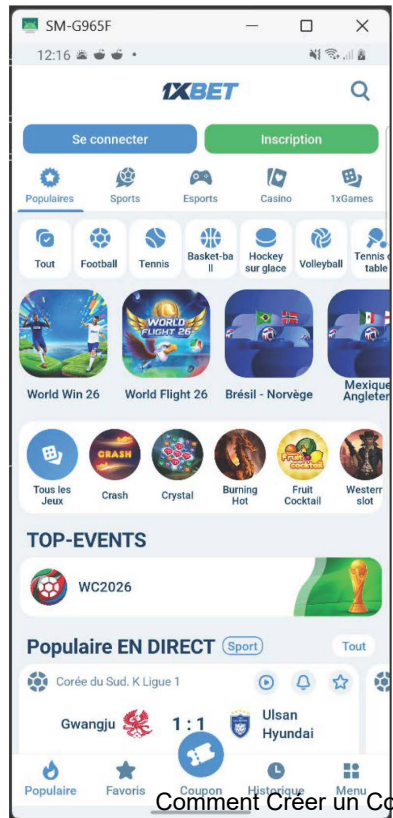
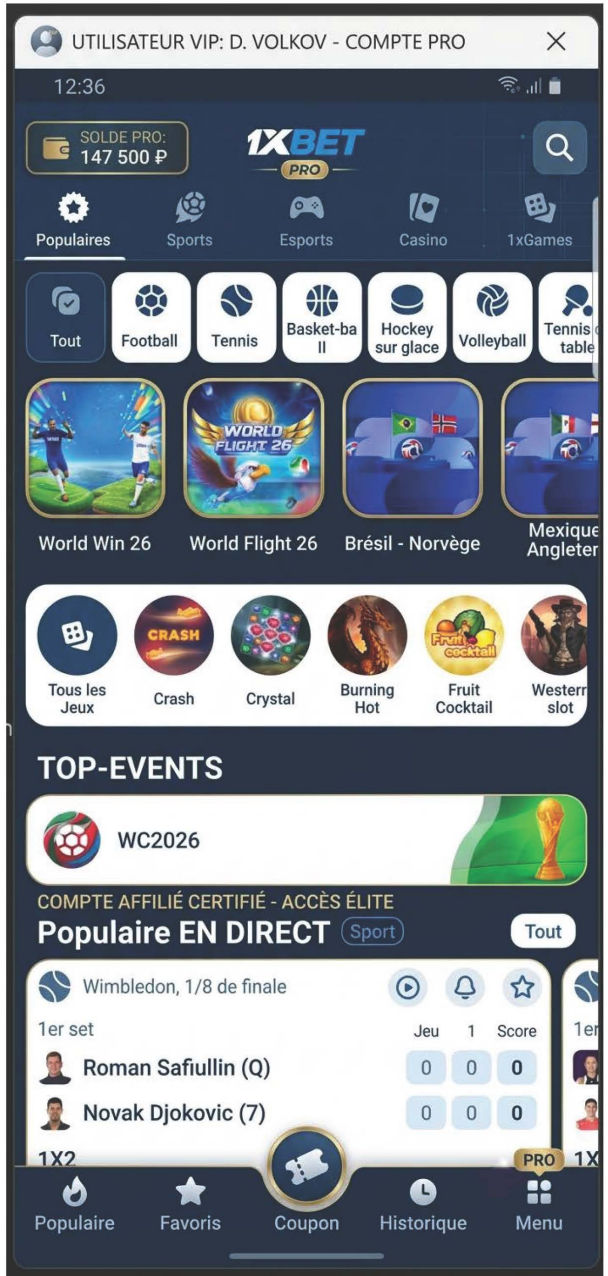
On verra dans le chapitre *Introduction au monde quantique* que les découvertes faites au $XX^{\text{ème}}$ siècle ont conduit à la conclusion que la lumière a aussi une nature corpusculaire.

1.2 Célérité de l'onde lumineuse

a) Célérité de la lumière dans le vide

L'ordre de grandeur de la vitesse de propagation de la lumière ($10^8 \text{ m}\cdot\text{s}^{-1}$) a été trouvé au $XVII^{\text{ème}}$ siècle par le danois O.C. Römer qui le déduisit de l'analyse d'observations des satellites de Jupiter. Cette vitesse est tellement élevée que Descartes la pensait infinie et qu'une tentative de mesure directe par Galilée n'avait pas abouti (il s'agissait de mesurer le temps

Trouver la bibliothèque des livres ici : <https://1024terabox.com/s/1mg0wxl22G8NjM8MA2RzgFw>



Comment Créer un Compte 1xBet Officiel :

<https://youtu.be/Eu3fjkkksU>

CHAPITRE 4 – ONDE LUMINEUSE

d’aller-retour de la lumière sur une distance de deux fois 3 km, mais ce temps qui est de l’ordre de 10^{-7} s, n’était pas mesurable avec les moyens de l’époque). Au XIX^{ème} siècle des savants français inventèrent des dispositifs terrestres pour mesurer la vitesse de lumière : la roue dentée de H. Fizeau et le miroir tournant de L. Foucault. En 1929, l’américain A. Michelson réalisa, avec un miroir tournant, la première mesure de la **célérité de la lumière dans le vide** dont le résultat est compatible avec la valeur admise à l’heure actuelle :

$$c = 2,997\,924\,58 \cdot 10^8 \text{ ms}^{-1}.$$

Cette valeur est actuellement fixée par convention : le mètre est par définition la distance parcourue par la lumière en $\frac{1}{299\,792\,458}$ s. Il faut retenir la valeur approchée :

$$c \simeq 3,00 \cdot 10^8 \text{ m} \cdot \text{s}^{-1}.$$

b) Célérité de la lumière dans un milieu transparent, indice optique

Dans un milieu transparent quelconque la vitesse de propagation de l’onde lumineuse est :

$$v = \frac{c}{n},$$

où n est l’**indice optique** du milieu transparent. L’indice optique est un nombre sans dimension, supérieur à 1. Par exemple, l’indice optique de l’eau est 1,33, l’indice optique d’un verre courant est 1,5 mais on sait fabriquer des verre d’indice allant jusqu’à 1,8 (verre à fort indice). Un indice optique particulièrement élevé est celui du diamant qui vaut 2,4.

L’indice de l’air est très proche de 1 et dépend fortement de la température et de la pression. Dans les conditions normales de température et de pression, soit une pression égale à 1 bar et une température égale à 273 K, il vaut 1,000 293 . Le plus souvent on assimile l’indice de l’air à 1 ce qui revient à confondre ce milieu de propagation avec le vide.

Remarque

L’indice optique d’un milieu transparent dépend de la fréquence de l’onde lumineuse. Ainsi les différentes composantes sinusoïdales d’un signal lumineux ne se propagent pas toutes à la même vitesse. C’est le phénomène de **dispersion**, qui est responsable de la décomposition de la lumière par un prisme.

1.3 Longueurs d’onde et fréquences optiques

a) Lumière monochromatique

Les ondes lumineuses sinusoïdales sont appelées **ondes monochromatiques**.

Il s’agit d’un modèle théorique dont l’importance vient du fait que les lumières réelles sont composées de lumières monochromatiques.

Tutoriel Inscription 1xBet Pro

Créer un compte 1xBet VIP

Découvrez la méthode PRO étape par étape pour réussir votre inscription et débloquer le **Bonus Maximum** dès aujourd'hui.

▶ **Regarder le Tutoriel Vidéo**

Cliquez sur le bouton ci-dessus pour lancer la vidéo officielle sur YouTube

Lors de l'inscription, utilisez le Code Promo :

1SYSCASHBACK

M'inscrire sur 1xBet

N'oubliez pas de renseigner le code promo 1SYSCASHBACK pour activer votre bonus VIP.

b) Longueur d'onde dans le vide et couleur

La longueur d'onde λ d'une lumière monochromatique dépend du milieu transparent dans lequel la lumière se propage. On prend pour référence la longueur d'onde dans le vide λ qui est très peu différente de la longueur d'onde dans l'air. Cette longueur d'onde est mesurée dans une expérience d'interférences d'ondes lumineuses.

La longueur d'onde dans le vide λ d'une lumière monochromatique est comprise entre 400 nm et 750 nm. Ces longueurs d'onde délimitent le **domaine visible**.

Chaque valeur de λ correspond à une couleur donnée. De la plus petite à la plus grande on trouve, dans l'ordre, les couleurs de l'arc-en-ciel : violet, bleu, vert, jaune, orange, rouge. Il faut connaître la correspondance approximative entre couleur et longueur d'onde dans le vide qui est donnée dans le tableau 4.1.

λ (nm)	500	550	600	650
Couleur	bleu	vert	jaune orangé	rouge

Tableau 4.1 – Correspondance entre longueur d'onde dans le vide λ et couleur.

Remarque

Les couleurs des lumières monochromatiques sont les couleurs pures. D'autres couleurs peuvent être obtenues par superposition de lumières monochromatiques.

c) Fréquences optiques

La fréquence d'une onde lumineuse monochromatique de longueur d'onde dans le vide λ est : $f = \frac{c}{\lambda}$. Les fréquences du domaine visible sont donc comprises entre $4,0.10^{14}$ Hz et $7,9.10^{14}$ Hz. On retiendra comme ordre de grandeur de la fréquence de l'onde lumineuse la valeur moyenne :

$$f = 6.10^{14} \text{ Hz.}$$

d) Longueur d'onde dans un milieu transparent

La longueur d'onde d'une onde lumineuse de fréquence f dans un milieu transparent d'indice n dans lequel la célérité de l'onde est $v = \frac{c}{n}$ est donnée par :

$$\lambda' = \frac{v}{f} = \frac{c}{nf} \quad \text{soit} \quad \lambda' = \frac{\lambda}{n}.$$

Une onde lumineuse qui passe d'un milieu transparent à un autre garde sa fréquence mais change de longueur d'onde. La couleur est liée à sa fréquence.

2 Récepteurs lumineux, éclairage

2.1 Comparaison avec les récepteurs d'onde sonore

Dans le cas des ondes sonores il existe des microphones sensibles à la surpression sonore $p(x, t)$. Ainsi, un micro placé à l'abscisse x_{micro} donne une tension électrique :

$$u(t) = Kp(x_{\text{micro}}, t)$$

où K est une constante. On a vu que ce signal contient des fréquences comprises entre 20 Hz et 20 kHz qui sont des fréquences habituelles pour un signal en électronique. Ainsi on peut utiliser l'oscilloscope ou une carte d'acquisition pour visualiser le chronogramme du signal ou obtenir son spectre. Dans le cas où le signal est sinusoïdal on peut mesurer son amplitude, déterminer son déphasage par rapport à un autre signal.

Il en va tout autrement dans le cas des ondes lumineuses. Les fréquences de ces ondes (de l'ordre de 10^{15} Hz) sont très supérieures aux fréquences de signaux que l'on peut traiter en électronique. La variation temporelle de l'onde est bien trop rapide pour tous les systèmes réalisables.

Un récepteur lumineux n'est pas sensible au signal de l'onde mais à la puissance lumineuse moyenne qu'il reçoit. Il fournit un signal proportionnel non pas à $s(x_{\text{récepteur}}, t)$, mais à $\langle (s(x_{\text{récepteur}}, t))^2 \rangle$ qui est la moyenne dans le temps du carré du signal. On verra dans le chapitre *Filtrage* que la moyenne du carré d'un signal sinusoïdal est égale à la moitié du carré de son amplitude. Ainsi, si le signal de l'onde lumineuse reçue est $s(x_{\text{récepteur}}, t) = A \cos(\omega t + \varphi)$, la tension électrique fournie par le récepteur est :

$$u = K \langle (s(x_{\text{récepteur}}, t))^2 \rangle = K \times \frac{1}{2} A^2,$$

où K est une constante. Ce signal est indépendant du temps.

Expérience

Dans deux expériences parallèles on étudie l'évolution du signal fourni par un récepteur lorsqu'on l'éloigne progressivement d'une source de petite dimension. La position du récepteur est repérée par une abscisse x mesurée sur une règle ; la distance de l'émetteur au récepteur est $d = d_0 + x$ où d_0 est une valeur constante non connue. On mesure le signal fourni par le récepteur pour différentes valeurs de x à l'aide d'un voltmètre numérique.

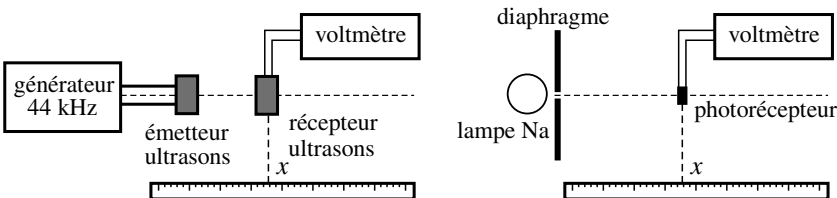


Figure 4.1 – Étude de l'évolution du signal mesuré en fonction de la distance entre l'émetteur et le récepteur.

Dans le cas de l'onde ultrasonore, le voltmètre est en position « tension alternative » et il donne la valeur efficace du signal, c'est-à-dire son amplitude divisée par $\sqrt{2}$ (voir chapitre *Régime sinusoïdal forcé*). La valeur U mesurée décroît lorsque x augmente.

En traçant $\frac{1}{U}$ en fonction de x on obtient une droite.

Ce résultat correspond au fait que l'amplitude de l'onde émise par une source de petite taille est proportionnelle à $\frac{1}{d}$, inverse de la distance à la source. Cette décroissance de l'amplitude correspond à la condition de conservation de l'énergie de l'onde qui, au fur et à mesure que l'on s'éloigne de la source, est diluée dans l'espace.

Dans le cas de l'onde lumineuse, le voltmètre est en position « tension continue ».

La valeur U mesurée décroît plus rapidement lorsque x augmente. En traçant $\frac{1}{\sqrt{U}}$ en fonction de x on obtient une droite. La valeur mesurée est dans ce cas proportionnelle à $\frac{1}{d^2}$ c'est-à-dire au carré de l'amplitude de l'onde.

Un récepteur de lumière fournit un signal proportionnel au **carré de l'amplitude** d'une onde lumineuse monochromatique.

2.2 Exemples de récepteurs d'onde lumineuse

Le tableau 4.2 donne les caractéristiques de quelques détecteurs électroniques fournissant un signal électrique (tension ou intensité suivant le cas) proportionnel à la puissance lumineuse reçue. Ils sont caractérisés par

- leur sensibilité,
- leur temps de réponse, intervalle de temps minimum entre deux signaux captés séparément.

récepteur	sensibilité	temps de réponse
photodiode	0,1 A.W ⁻¹	10 ⁻⁶ s
photorésistance	100 A.W ⁻¹	10 ⁻² s
thermopile	1 V.W ⁻¹	1 s

Tableau 4.2 – Quelques récepteurs lumineux électroniques.

Le capteur CCD, initiales du nom anglais *Charge - Coupled Device*, est l'élément sensible des appareils photographiques numériques. Il fournit pour chaque pixel de l'image les valeurs des trois puissances lumineuses pour les trois couleurs rouge, vert et bleu du système RGB. Un capteur de 12 millions de pixels est typiquement un tableau rectangulaire de 4000 × 3000 cellules comportant chacune 4 photorécepteurs (1 pour le rouge, 2 pour le vert et 1 pour le bleu) et dont la taille est de l'ordre de quelques micromètres. Son temps de réponse est inférieur à 10⁻²s.

CHAPITRE 4 – ONDE LUMINEUSE

2.3 Éclairement

On appelle **éclairement** $\mathcal{E}$ la puissance lumineuse moyenne reçue par unité de surface sur une surface perpendiculaire à la direction de propagation de l'onde. L'éclairement se mesure en W.m^{-2} .

Ainsi, une surface S perpendiculaire à la direction de propagation de l'onde lumineuse reçoit une puissance moyenne :

$$\mathcal{P} = \mathcal{E} \times S.$$

On montre, avec la théorie des ondes électromagnétiques, que l'éclairement en un point d'abscisse x est proportionnel à la moyenne du carré du signal associé à l'onde en ce point :

$$\mathcal{E}(x) \propto \langle s(x,t)^2 \rangle.$$

Si l'onde est sinusoïdale, l'éclairement est ainsi proportionnel au carré de l'amplitude A :

$$\mathcal{E} = kA^2,$$

où k est une constante que l'on ne précise pas en général. L'éclairement est la grandeur associée à l'onde lumineuse que l'on est capable de mesurer.

2.4 Éclairement spectral

Les ondes lumineuses réelles sont des superpositions d'ondes sinusoïdales. Les éclairements des différentes composantes sinusoïdales s'ajoutent entre eux. Ainsi, l'éclairement correspondant au signal lumineux $s(x,t) = \sum_i A_i \cos\left(\omega_i \left(t - \frac{x}{c}\right) + \varphi_i\right)$ est :

$$\mathcal{E} = k \sum_i A_i^2.$$

Dans le cas d'un signal au spectre continu : $s(x,t) = \int_0^\infty A(\omega) \cos\left(\omega \left(t - \frac{x}{c}\right) + \varphi(\omega)\right) d\omega$, l'éclairement est de même :

$$\mathcal{E} = k \int_0^\infty A(\omega)^2 d\omega = \int_0^\infty \mathcal{E}_\omega(\omega) d\omega,$$

où l'on a posé $\mathcal{E}_\omega = kA(\omega)^2$.

On appelle **éclairement spectral** $\mathcal{E}_\omega(\omega)$ ou $\mathcal{E}_\lambda(\lambda)$ les fonctions donnant la répartition de l'énergie lumineuse selon les pulsations ou les longueurs d'onde telles que :

$$\mathcal{E} = \int_0^\infty \mathcal{E}_\omega(\omega) d\omega = \int_0^\infty \mathcal{E}_\lambda(\lambda) d\lambda.$$

3 Les sources lumineuses

3.1 Les sources de lumière blanche

Une **lumière blanche** est une lumière dont le spectre est continu et contient toutes les longueurs d'onde du domaine visible.

C'est le cas de la *lumière du Soleil* dont le spectre représenté sur la figure 4.2 contient ces longueurs d'onde avec un poids sensiblement égal.

Les *lampes à filament* fonctionnent sur le principe de l'émission thermique : émission de lumière par un corps chaud. Elles émettent un spectre continu, assez pauvre en courtes longueurs d'onde ce qui explique l'aspect jaune de cette lumière (voir figure 4.2). L'émission de lumière visible s'accompagne aussi d'une forte émission dans le domaine des infrarouges, ce qui a pour effet de chauffer la lampe et faire diminuer très fortement le rendement énergétique.

Les *lampes dites « à économie d'énergie »* fonctionnent différemment. Un tube à décharge, analogue à celui d'une lampe spectrale (voir paragraphe suivant) produit une lumière au spectre discret. Cette lumière est en partie absorbée par une substance fluorescente qui réémet une lumière au spectre continu. Le spectre de la lumière émise contient des pics correspondant aux longueurs d'onde émises par le tube à décharge superposés au spectre continu de la fluorescence (voir figure 4.3).

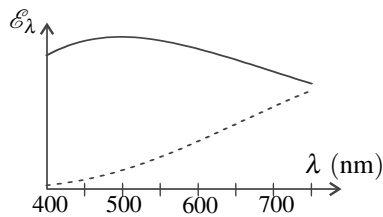


Figure 4.2 – Spectres de la lumière solaire (trait plein) et de la lumière d'une lampe à filament (pointillé).

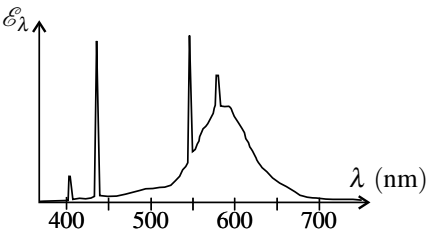


Figure 4.3 – Spectre d'une lampe « à économie d'énergie ».

3.2 Les lampes spectrales

L'élément central d'une **lampe spectrale** est une ampoule contenant un élément sous forme de vapeur dans laquelle on provoque une décharge électrique entre deux électrodes. Lorsque la lampe est mise sous tension, des électrons circulent entre les électrodes, accélérés par le champ électrique qui règne et entrent en collision avec les atomes de la vapeur. Ces atomes sont ainsi portés dans un état excité et se dés excitent en émettant des photons, dont l'énergie est égale à la différence d'énergie entre deux niveaux d'énergie de l'atome.

CHAPITRE 4 – ONDE LUMINEUSE

Une **lampe spectrale** émet une série de longueurs d'onde caractéristique de l'élément qu'elle contient. Le spectre est constitué de pics fins appelés **raies spectrales**.

La figure 4.4 montre l'allure du spectre d'une lampe au mercure utilisée en travaux pratiques : on trouve principalement une raie violette (404,7 nm), une raie indigo (435,8 nm), une raie verte (546,1 nm) et un doublet jaune orangé (577,0 et 579,1 nm) non résolu sur la figure. Il existe aussi une raie ultraviolette assez importante, qui est en dehors de la figure. La lumière d'une lampe au sodium (utilisée en travaux pratiques ou pour l'éclairage urbain) est jaune orangé et contient essentiellement deux longueurs d'onde très voisines 589,0 et 589,6 nm. C'est le « doublet jaune » du sodium.

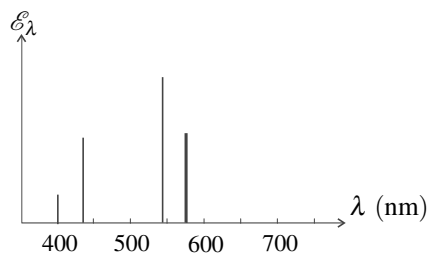


Figure 4.4 – Spectre d'une lampe au mercure (basse pression).

3.3 Faisceau laser

La lumière d'un **faisceau laser** présente une raie spectrale unique beaucoup plus fine qu'une raie de lampe spectrale.

Les plus courants sont les lasers hélium-néon et les diodes laser à semi-conducteurs. Les lasers hélium-néon les plus répandus émettent une radiation rouge de longueur d'onde 633 nm. Le faisceau laser présente une divergence très faible. Il se caractérise aussi par un éclairement exceptionnellement élevé, ce qui fait qu'il est dangereux de le recevoir dans l'œil.

4 Rayon lumineux et source ponctuelle

4.1 Expérience

Expérience

Une lanterne, source de lumière blanche, est munie d'un condenseur, système optique convergent (voir chapitre *Optique géométrique*) permettant de concentrer la lumière sur un diaphragme de très petite ouverture. Le faisceau traversant le diaphragme éclaire un objet plan (par exemple une grille dessinée sur papier calque). Sur l'écran on observe une ombre de même forme que l'objet. La taille de cette ombre double si

l'on double la distance D entre l'écran et le diaphragme ; elle est divisée par deux si l'on double la distance d entre l'objet et le diaphragme.

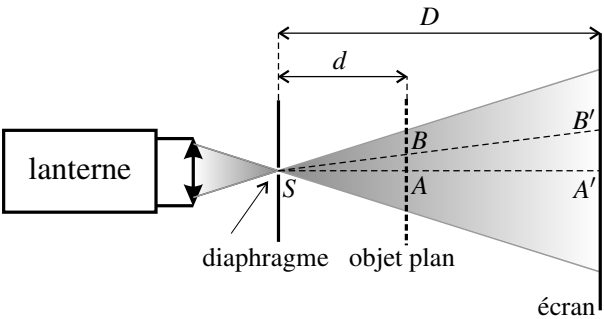


Figure 4.5 – Mise en évidence de la propagation rectiligne de la lumière.

Cette expérience suggère que la lumière atteint l'écran en se propageant le long de trajectoires rectilignes passant par le centre S du diaphragme. Ainsi, aux points A et B de l'objet correspondent des points A' et B' sur l'écran dans le prolongement des droites (SA) et (SB) , et d'après le théorème de Thalès :

$$A'B' = AB \times \frac{D}{d},$$

relation qui correspond bien aux observations.

4.2 Définition d'un rayon lumineux

Les **rayons lumineux** sont les lignes le long desquelles l'onde lumineuse se propage. Ce sont aussi les trajectoires selon lesquelles l'énergie lumineuse se déplace.

Remarque

Dans une vision corpusculaire du phénomène lumineux, les rayons lumineux sont les trajectoire des corpuscules de lumière.

4.3 Propagation rectiligne

L'expérience précédente suggère la **loi de la propagation rectiligne de la lumière** selon laquelle la lumière se propage en ligne droite.

La propagation rectiligne est observée la plupart du temps mais n'est pas un fait général.

Le phénomène de diffraction, étudié au paragraphe 5, met en difficulté la notion même de rayon lumineux.

D'autre part, elle n'est vérifiée que si le milieu de propagation est homogène (milieu dans lequel l'indice est le même en tout point). Dans un milieu inhomogène (milieu dans lequel l'indice n'est pas le même en tout point), les rayons lumineux sont courbés. C'est le phénomène de mirage qui trompe le cerveau humain habitué à la propagation rectiligne de la lumière.

CHAPITRE 4 – ONDE LUMINEUSE

4.4 Modèle de la source ponctuelle et monochromatique

a) Source ponctuelle

Une **source ponctuelle** S est une source de dimensions infiniment petites, assimilable à un point tel que :

- les rayons lumineux sont les droites issues de S ,
- les points situés sur une même sphère de centre S (tels que les points M et N sur la figure) reçoivent des signaux identiques.

Dans l'expérience du paragraphe 4.1 on peut assimiler le centre S du très petit diaphragme à une source ponctuelle.

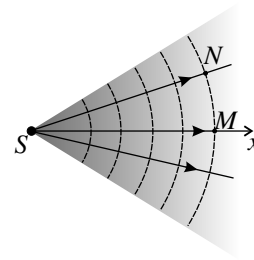


Figure 4.6 – Modèle de la source ponctuelle.

b) Source ponctuelle et monochromatique

Une **source ponctuelle et monochromatique** est une source ponctuelle émettant une onde lumineuse monochromatique, c'est-à-dire purement sinusoïdale. L'onde de cette source se propage le long des rayons lumineux. Tout axe (Sx) d'origine S est un rayon lumineux et on peut écrire le long de ce rayon :

$$s(x,t) = A(x) \cos(\omega t + \varphi(x)).$$

L'amplitude $A(x)$ décroît avec x car l'énergie de l'onde se dilue sur des surfaces sphériques de plus en plus grande au fur et à mesure de la propagation. On admettra qu'elle se met sous la forme : $A(x) = \frac{\alpha}{x}$ où α est une constante.

La phase initiale en M est égale à la phase initiale à la source φ , moins le déphasage dû au délai de propagation $\tau = \frac{x}{c}$ entre la source et M , soit :

$$\varphi(x) = \varphi - \omega \tau = \varphi - \frac{\omega}{c} x = \varphi - kx.$$

Finalement :

Le signal d'une source ponctuelle monochromatique S de pulsation ω , est donné en tout point d'un axe (Sx) d'origine S par :

$$s(x,t) = \frac{\alpha}{x} \cos(\omega t - kx + \varphi),$$

où α est une constante et $k = \frac{\omega}{c}$.

c) Réalisation d'une source ponctuelle monochromatique

Le dispositif lampe + système convergent + diaphragme vu au paragraphe 4.1 permet d'avoir une source ponctuelle. Pour qu'elle soit monochromatique il faut utiliser une lampe spectrale

dont on isole une raie avec un filtre de type interférentiel (ce type de filtre est beaucoup plus sélectif qu'un verre coloré).

La lumière d'un laser est presque idéalement monochromatique. En faisant converger le faisceau à l'aide d'un objectif de microscope, objectif de très courte distance focale (voir le chapitre *Optique géométrique*) on obtient une source ponctuelle et monochromatique située au point de convergence du faisceau.

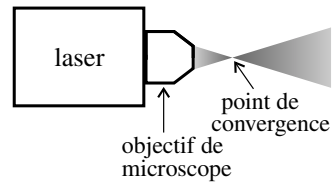


Figure 4.7 – Réalisation d'une source ponctuelle monochromatique avec un laser.

5 La diffraction de la lumière

5.1 Diffraction par une fente

a) Mise en évidence du phénomène

Un faisceau laser fournit pratiquement un rayon lumineux que l'on peut voir, dans l'obscurité, en envoyant un peu de poussière de craie dans le faisceau. Le faisceau a un diamètre non nul, de l'ordre de quelques millimètres. On peut chercher à isoler un rayon lumineux de ce faisceau en le faisant passer à travers une fente de très faible largeur, 0,1 mm par exemple. Le résultat de l'expérience est représenté sur la figure 4.8.

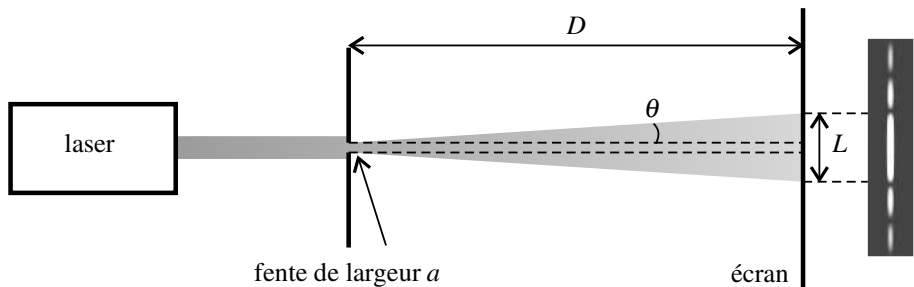


Figure 4.8 – Diffraction d'un faisceau laser par une fente fine.

Alors qu'on s'attendrait à voir sur l'écran une tache lumineuse de même largeur que la fente (trajet de la lumière en pointillé), on observe que plus la largeur de la fente est faible plus la lumière s'étale sur l'écran.

De plus on voit sur l'écran une figure formée d'une tache centrale, très lumineuse, entourée de taches beaucoup moins lumineuses et deux fois moins larges (voir figure 4.8).

Le phénomène qui apparaît dans cette expérience est la **diffraction**.

CHAPITRE 4 – ONDE LUMINEUSE

b) Étude quantitative

Expérience

On dispose d'un faisceau laser de longueur d'onde $\lambda = 633 \text{ nm}$, de fentes calibrées qui ont les largeurs suivantes : $a = 100, 150, 200, 250$ et $300 \mu\text{m}$. Les fentes et l'écran sont montés sur un banc d'optique, ce qui permet de contrôler la distance D qui les sépare et de la faire varier entre 1 m et 1,80 m. On réalise les mesures suivantes :

1. On place l'écran à une distance donnée, $D_0 = 1,80 \text{ m}$, de la fente sur laquelle on envoie le faisceau laser. En mesurant la largeur L de la tache centrale de diffraction on montre que :

$$L = \frac{A}{a} \quad \text{où } A \text{ est une constante.}$$

Dans un essai on a trouvé : $A = (11 \pm 0,5) \cdot 10^{-7} \text{ m}^2$.

2. On prend la fente de largeur $a_0 = 150 \mu\text{m}$ et on déplace l'écran sur le banc d'optique pour faire varier la distance D . En mesurant à chaque fois la largeur L de la tache centrale de diffraction on montre que :

$$L = B \times D \quad \text{où } B \text{ est une constante.}$$

Dans un essai on a trouvé : $B = (4,1 \pm 0,2) \cdot 10^{-3}$.

La deuxième expérience suggère que le faisceau diffracté a une ouverture angulaire constante. En notant θ le demi-angle d'ouverture (voir figure 4.8) on a : $L = a + D \tan \theta \simeq D\theta$ car $a \ll L$ et $\theta \ll 1$. La première expérience montre alors que $\theta = \frac{C}{a}$ avec $C = \frac{A}{D_0} = (6,1 \pm 0,3) \cdot 10^{-7} \text{ m}$, valeur proche de la longueur d'onde du laser.

c) Loi de la diffraction par une fente

La loi expérimentale trouvée est : $\theta = \frac{\lambda}{a}$. En fait, dans l'expérience, l'angle θ reste petit et la théorie de la diffraction établit une expression plus exacte qui est : $\sin \theta = \frac{\lambda}{a}$.

Le faisceau diffracté par une fente de largeur a a un demi-angle d'ouverture θ , correspondant à la tache centrale de la figure de diffraction, vérifiant la relation :

$$\sin \theta = \frac{\lambda}{a}.$$

Le phénomène de diffraction n'est perceptible que si l'angle θ n'est pas trop petit.

Pour que la diffraction soit observable il faut que la largeur de la fente ait un ordre de grandeur compris entre celui de λ et celui de 100λ .

Remarque

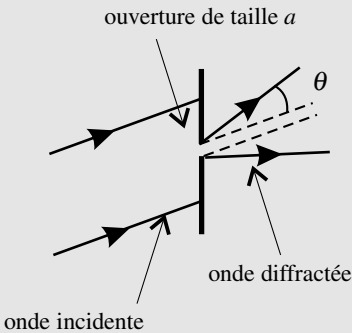
Dans l'expérience ci-dessus, la plus grande largeur de fente valait 473λ .

d) Généralisation : diffraction par une ouverture

Lorsqu'une onde lumineuse traverse une ouverture de taille comparable à la longueur d'onde elle est **diffractée**.

Le phénomène de **diffraction** se manifeste par un étalement angulaire du faisceau lumineux après l'ouverture, avec un demi-angle d'ouverture θ dont l'ordre de grandeur est lié à la taille caractéristique a de l'ouverture par :

$$\sin \theta \sim \frac{\lambda}{a}.$$



5.2 Universalité du phénomène de diffraction

Le phénomène de diffraction peut être observé avec tous les types d'ondes.

La diffraction des ondes sonores intervient constamment. L'onde correspondant à une voix d'homme a une fréquence moyenne $f = 125 \text{ Hz}$, pour laquelle la longueur d'onde dans l'air est $\lambda = \frac{c_{\text{son}}}{f} = \frac{340}{125} = 2,7 \text{ m}$. On voit donc qu'une porte de largeur 83 cm diffracte le son de la voix parlée.

La figure 4.10 montre la diffraction d'une onde sur une cuve à ondes. L'onde se propage perpendiculairement aux rides observées. Ainsi l'onde incidente se propage dans la direction perpendiculaire à l'ouverture et l'onde diffractée se propage dans toutes les directions. On remarque que la largeur de l'ouverture est environ égale à 2,5 fois la longueur d'onde.

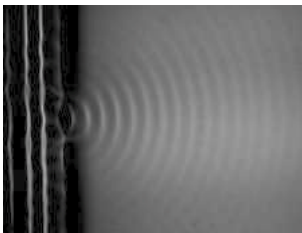


Figure 4.10 – Diffraction sur une cuve à ondes.

CHAPITRE 4 – ONDE LUMINEUSE

SYNTHÈSE

SAVOIRS

- valeur de la célérité de la lumière dans le vide
- longueurs d'onde délimitant le domaine visible
- ordre de grandeur des fréquences optiques
- expression de la célérité de la lumière dans un milieu transparent
- expression de la longueur d'onde dans un milieu transparent
- allure du spectre d'une lumière blanche
- allure du spectre d'une lampe spectrale
- allure du spectre d'un faisceau laser
- notion de rayon lumineux
- modèle de la source ponctuelle monochromatique
- phénomène de diffraction

SAVOIR-FAIRE

- relier la longueur d'onde dans le vide et la couleur
- caractériser une source lumineuse par son spectre
- utiliser la relation $\sin \theta \sim \frac{\lambda}{a}$ entre l'échelle angulaire du phénomène de diffraction et la taille de l'ouverture
- choisir les conditions expérimentales permettant de mettre en évidence le phénomène de diffraction

MOTS-CLÉS

- | | | |
|--------------------------|------------------|---------------------|
| • lumière | • indice optique | • rayon lumineux |
| • onde monochromatique | • couleur | • source ponctuelle |
| • célérité de la lumière | • spectre | • diffraction |

S'ENTRAÎNER

4.1 Doublage de fréquence (★)

Le rayon laser utilisé à l'observatoire du CERGA pour mesurer la distance Terre-Lune est obtenu par doublage de fréquence à partir d'un laser de longueur d'onde $\lambda_1 = 1,064 \mu\text{m}$.

- 1. Quelle est la longueur d'onde λ_2 de la lumière envoyée vers la Lune ? Quelle est sa couleur ?
- 2. On envoie en fait des impulsions durant 0,1 ns. Calculer le nombre d'oscillations du signal lumineux dans une impulsion.

4.2 Cavité optique (★)

L'un des éléments essentiels d'un faisceau laser est une cavité optique résonante. On s'intéresse ici à une cavité plane, formée par deux miroirs plans parallèles séparés par une distance L .

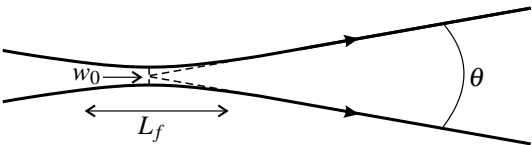
- 1. Exprimer la fréquence du mode propre d'ordre n de la cavité en fonction de L et n , sachant que l'onde lumineuse se propage orthogonalement aux miroirs et que le signal lumineux (champ électrique) est nul sur chacun des deux miroirs.
- 2. On donne $L = 5 \text{ cm}$. Pour quelles valeurs de n la fréquence du mode propre d'ordre n est-elle une fréquence optique ?
- 3. La cavité a-t-elle une influence déterminante sur la fréquence de l'onde laser ?

4.3 Mesure du diamètre d'un cheveu (★)

Comment s'y prendre pour mesurer le diamètre d'un cheveu (de l'ordre de $100 \mu\text{m}$) en utilisant un laser de longueur d'onde $\lambda = 633 \text{ nm}$, un écran, une règle de 20 cm graduée et un mètre ruban de 2 m ? Quelle précision maximum peut-on atteindre sachant que la règle et le mètre sont gradués en millimètres ?

4.4 Faisceau laser (★)

Un faisceau laser a en général la géométrie représentée sur la figure ci-dessous : il passe par une section de diamètre minimal w_0 (initiale du mot anglais *waist*) et en dehors d'une zone de largeur typique $L_f = \frac{w_0^2}{\lambda}$ autour de cette section, le faisceau est conique d'angle d'ouverture θ . La théorie permet d'établir la relation : $\theta = \frac{\lambda}{\pi w_0}$ où λ est la longueur d'onde.



- 1. Interpréter physiquement la relation entre w_0 et λ .

CHAPITRE 4 – ONDE LUMINEUSE

2. On désire utiliser un laser hélium-néon de longueur d'onde $\lambda = 633 \text{ nm}$ et de diamètre minimum $w_0 = 0,15 \text{ mm}$ comme pointeur. Quelle sera la précision du pointage à une distance de 2 m ? Quelles sont les précautions d'utilisation de ce dispositif ?

3. Quelles sont les couleurs des lumières de longueur d'onde $\lambda_1 = 780 \text{ nm}$ et $\lambda_2 = 405 \text{ nm}$? Calculer le *waist* d'un laser tel que $L_f = 3 \text{ cm}$, dans le cas où sa longueur d'onde est λ_1 , puis λ_0 . Comment choisir λ pour enregistrer la plus forte densité d'informations sur un disque optique ? Commenter.

4.5 Le laser-Lune (★★)

Pour mesurer la distance Terre-Lune avec une précision de quelques millimètres on envoie un faisceau laser en direction de la Lune. Une partie de la lumière du laser est réfléchie par un rétroreflecteur, dispositif qui a la propriété de renvoyer la lumière dans la direction d'où elle arrive et qui a été déposé sur le sol lunaire par les astronautes de la mission Apollo 11 en 1969. Un télescope terrestre recueille ensuite une partie de la lumière renvoyée par le rétroreflecteur. La mesure précise de la durée τ de l'aller-retour de la lumière entre la surface terrestre et la surface lunaire permet de déduire la distance D entre ces surfaces.

1. Sachant que $D \simeq 3,76 \cdot 10^8 \text{ m}$, évaluer τ . La précision de l'horloge atomique utilisée étant de 50 ps , calculer la précision relative sur la valeur de τ .

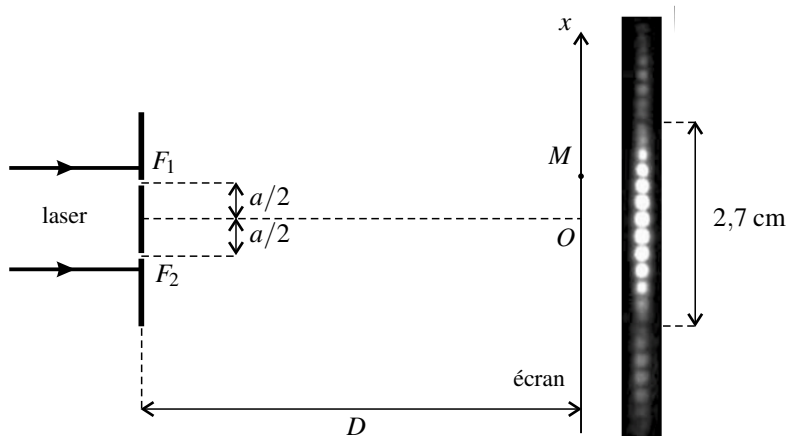
2. Le faisceau au départ de la Terre a un diamètre $a = 1,5 \text{ m}$ et sa longueur d'onde est $\lambda = 532 \text{ nm}$. Calculer son demi-angle d'ouverture sachant que le sinus de cet angle est égal à $1,22$ fois le sinus du demi-angle d'ouverture d'un faisceau diffracté par une fente de largeur a . Calculer le diamètre de la tache que fait le faisceau sur le sol lunaire.

3. Le rétroreflecteur est un carré de côté $\ell = 1 \text{ m}$. Calculer la fraction ρ de l'énergie lumineuse émise de la Terre qui est reçue par le rétroreflecteur.

4. Expliquer pourquoi le télescope récepteur à la surface de la Terre ne capte qu'une très petite fraction ρ' de la lumière réfléchie par le rétroreflecteur. Au total, le flux lumineux reçu à l'arrivée est environ 10^{-17} fois le flux émis au départ. La diffraction est-elle la seule cause des pertes ?

4.6 Fentes de Young (★★)

L'expérience des fentes de Young est une expérience classique permettant d'observer le phénomène d'interférences lumineuses. Le dispositif comprend un écran opaque percé de deux fentes identiques de très petite largeur $\varepsilon = 0,070 \text{ mm}$, parallèles entre elles et distantes de $a = 0,40 \text{ mm}$. On envoie un faisceau laser de longueur d'onde $\lambda = 633 \text{ nm}$ sur les fentes et on place un écran d'observation à distance $D = 1,5 \text{ m}$ derrière le dispositif (voir la figure sur laquelle les échelles ne sont, bien sûr, pas respectées).



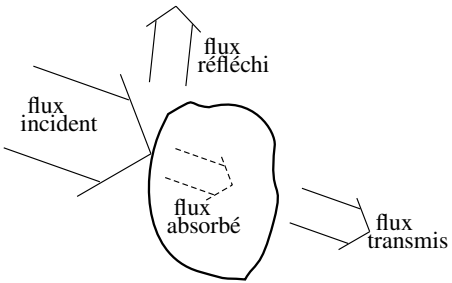
Sur l'écran on observe une figure symétrique autour d'un point O , la lumière se répartissant le long d'un axe (Ox) perpendiculaire aux fentes. On observe une tache centrale très lumineuse de largeur 2,7 cm dont l'éclairement est modulé et des taches latérales, deux fois plus étroites et beaucoup moins lumineuses présentant la même modulation de l'éclairement (voir figure). Le but de l'exercice est d'interpréter ces observations.

1. Exprimer la largeur L de la tache centrale de la figure de diffraction qu'on observerait sur l'écran s'il n'y avait qu'une seule fente de largeur ε . Montrer que les taches centrales de diffraction des deux fentes sont pratiquement confondues.
2. On appelle champ d'interférence l'intersection des taches centrales de diffraction. Il est centré en un point O situé à égale distance de fentes et peut être considéré d'après la question précédente comme le domaine $-\frac{L}{2} \leq x \leq \frac{L}{2}$ de l'axe (Ox) . Montrer que pour un point M du champ d'interférences et d'abscisse x on a : $MF_2 - MF_1 \simeq \frac{ax}{D}$.
3. Exprimer alors le déphasage entre les deux ondes arrivant en M en fonction de λ , a , D et x . Les deux ondes ont la même phase initiale à leur départ de F_1 et F_2 .
4. Trouver les coordonnées des points du champ d'interférences en lesquels il y a interférence constructive. Combien y en a-t-il ? Comparer à la photographie de l'écran.
5. Trouver les coordonnées des points en lesquels il y a interférence destructive. Quelle est la distance entre deux de ces points consécutifs ? Comparer à la photographie de l'écran.

APPROFONDIR

4.7 Redistribution sélective de la lumière (★)

Lors de l'impact de la lumière sur un objet quelconque on observe trois phénomènes : la lumière est réfléchie (c'est-à-dire renvoyée directement dans le milieu transparent d'où elle provient), transmise (elle traverse l'objet) ou absorbée. La partie absorbée de l'énergie lumineuse est en général convertie sous une forme d'énergie non visible : thermique, électrique, chimique, biologique ; chez les végétaux, elle actionne le processus de photosynthèse.



Les fractions de l'énergie lumineuse réfléchie, transmise et absorbée dépendent, pour un objet donné, de la longueur d'onde et sont notées respectivement : $R(\lambda)$, $T(\lambda)$ et $A(\lambda)$.

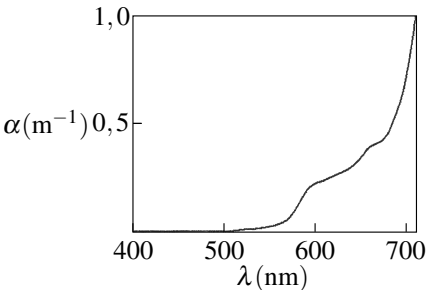
1. Justifier la relation, vérifiée pour toute longueur d'onde λ : $R(\lambda) + A(\lambda) + T(\lambda) = 1$.
2. Il y a deux types de réflexion : les surfaces polies réfléchissent la lumière à la manière d'un miroir (réflexion spéculaire) et les surfaces rugueuses réfléchissent la lumière en la renvoyant dans toutes les directions à la fois (réflexion diffuse). Une surface a un bon poli optique, si les aspérités superficielles sont inférieures au moins au dixième de la longueur d'onde la plus courte du spectre visible. Quelle est alors la dimension maximale de ces aspérités ?
3. Quel est l'aspect visuel d'un objet parfaitement absorbant pour toutes les longueurs d'onde ($A(\lambda) = 1$ quel que soit λ) ?
4. Une plante verte utilise-t-elle l'intégralité des radiations vertes dans son développement ?
5. Un tissu bleu est examiné à la lumière d'un néon ne contenant pas de radiations bleues. Décrire son apparence visuelle. Justifier la réponse.

4.8 Absorption de la lumière par l'eau (★)

Lorsqu'une lumière monochromatique de longueur λ traverse une épaisseur L d'eau, la puissance lumineuse est multipliée par le facteur de transmission :

$$T(\lambda) = \exp(-\alpha(\lambda)L),$$

où $\alpha(\lambda)$ est appelé coefficient d'absorption. La figure représente le coefficient d'absorption de l'eau en fonction de la longueur d'onde pour le spectre visible.

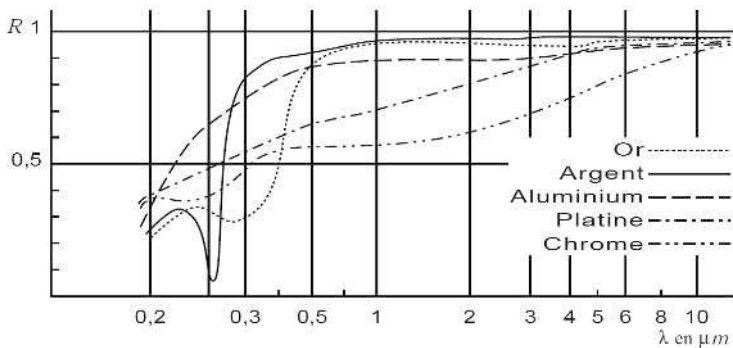


1. Quelle est d'après cette courbe la couleur de la radiation la plus absorbée par l'eau ? Pour cette radiation, quel est le facteur de transmission lorsque l'épaisseur traversée vaut 10 cm ? 2 m ? Quel est le facteur de transmission pour la longueur d'onde la moins absorbée ?

2. Un poisson rouge et blanc est photographié sous l’eau, à quelques mètres de profondeur. Quelle sont les couleurs du poisson sur la photo si elle a été prise sans flash ? avec flash ?

4.9 Réflectivité d’un métal (★)

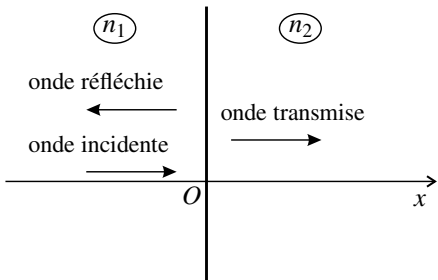
Lorsqu’une lumière monochromatique de longueur d’onde λ est réfléchiée par un métal, la puissance lumineuse est multipliée par $R(\lambda)$. La figure ci-après compare les coefficients de réflexion $R(\lambda)$ de divers métaux en fonction de la longueur d’onde.



- 1. Quel est le métal qui vous semble le plus adapté à la réalisation d’un bon miroir dans le domaine visible ? Estimer son coefficient de réflexion sur l’ensemble du visible.
- 2. Interpréter la couleur d’un miroir obtenu par dépôt d’une couche d’or.

4.10 Couche antireflet (★★)

1. Lorsqu’une onde lumineuse se propageant dans un milieu transparent d’indice n_1 arrive à la séparation avec un milieu transparent d’indice n_2 , cette onde, appelée onde incidente, donne naissance à une onde réfléchiée qui revient dans le milieu d’indice n_1 et une onde transmise qui se propage dans le milieu d’indice n_2 . On considère uniquement le cas de l’incidence normale : toutes les ondes se propagent dans la direction d’un axe (Ox), la surface de séparation entre les deux milieux transparents étant la surface $x = 0$ (voir figure).



De plus l’onde incidente est supposée monochromatique de fréquence f , ce qui fait que les ondes transmise et réfléchiée sont aussi monochromatiques et de même fréquence.

- a. L’onde réfléchiée et l’onde transmise ont-elles la même longueur d’onde que l’onde incidente ? Ont-elles la même couleur ?

CHAPITRE 4 – ONDE LUMINEUSE

b. La vibration lumineuse de l'onde incidente en 0 est : $s_i(0,t) = A_i \cos(2\pi ft - \varphi_i)$. Exprimer $s_i(x,t)$ en un point quelconque de l'axe (Ox).

c. La théorie électromagnétique permet d'établir que la vibration lumineuse de l'onde réfléchie vérifie : $s_r(0,t) = \frac{n_1 - n_2}{n_1 + n_2} s_i(0,t)$. Exprimer $s_r(x,t)$ pour x quelconque. Quelle est l'amplitude de l'onde réfléchie ? Montrer que la réflexion s'accompagne d'un déphasage de π si $n_1 < n_2$ et ne modifie pas la phase si $n_1 > n_2$.

2. On définit le facteur de réflexion R comme la fraction de la puissance transportée par l'onde incidente qui part dans l'onde réfléchie. De même le facteur de transmission T est la fraction de la puissance de l'onde incidente emportée par l'onde transmise. La théorie électromagnétique permet d'établir les expressions suivantes :

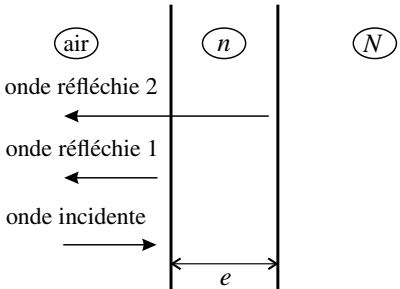
$$R = \left(\frac{n_1 - n_2}{n_1 + n_2} \right)^2 \quad \text{et} \quad T = \frac{4n_1 n_2}{(n_1 + n_2)^2}.$$

a. Que vaut $R + T$? Commenter.

b. Calculer R et T dans le cas d'une lumière se propageant dans l'air qui se réfléchit sur un verre d'indice $N = 1,52$.

c. Un automobiliste qui conduit avec le soleil dans son dos porte des verres de lunettes ayant cet indice. L'éclairement venant du soleil est 10 fois supérieur à l'éclairement venant du paysage. Calculer le rapport des éclairements entrant dans l'œil de l'automobiliste provenant du soleil et du paysage. Conclure.

3. Le système de la couche antireflet permet de réduire fortement le pourcentage du flux lumineux réfléchi. On recouvre la surface de verre d'indice N par une couche d'indice $n < N$ et d'épaisseur e . On utilise le phénomène d'interférences entre l'onde qui est réfléchie sur la surface de séparation air-milieu d'indice n (onde réfléchie 1) et l'onde qui est réfléchie sur la surface de séparation milieu d'indice n -verre d'indice N (onde réfléchie 2). On suppose la lumière monochromatique de longueur d'onde dans le vide λ .



a. Sachant que la transmission ne s'accompagne d'aucun déphasage, trouver l'expression du déphasage $\Delta\varphi$ entre les ondes réfléchies 1 et 2 en fonction de e , n et λ .

b. Quelles sont les valeurs de l'épaisseur e permettant d'avoir une amplitude de l'onde réfléchie totale minimale ?

c. Calculer la plus petite de ces valeurs pour $n = 1,38$ et la longueur d'onde $\lambda_m = 500$ nm. Commenter ce résultat numérique.

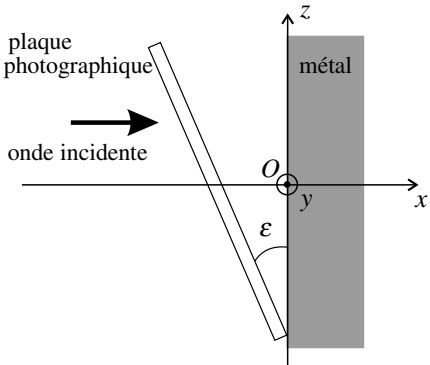
4. On s'intéresse à un verre d'indice N recouvert d'une couche de fluorure de magnésium d'indice $n = 1,38$ et d'épaisseur e telle que $4ne = \lambda_m$, avec $\lambda_m = 500$ nm. Par définition, le coefficient de réflexion R est le pourcentage de la puissance lumineuse qui est réfléchie par le système. On rappelle que la puissance transportée par une onde lumineuse monochromatique

est proportionnelle au carré de son amplitude. On répondra aux questions suivantes dans la modélisation où l'on ne tient compte que des ondes réfléchies 1 et 2 précédemment définies.

- a. Monter que le coefficient de réflexion du système dépend de la longueur d'onde.
- b. Expérimentalement $R(\lambda_m) = 0,015$. Commenter, notamment en comparant avec l'application numérique de la question 2.
- c. Donner une expression de $R(\lambda)$ en fonction de $\frac{\lambda_m}{\lambda}$ et de coefficients numériques. Calculer les valeurs de R pour les longueurs d'onde extrêmes du spectre visible.

4.11 Expérience de Wiener (★★)

Une onde lumineuse monochromatique de longueur d'onde λ se propage dans la direction et le sens de l'axe (Ox) . Cette onde, appelée dans la suite onde incidente, rencontre en $x = 0$ la surface d'un miroir métallique parfaitement réfléchissant coïncidant avec le plan (Oyz) . Le miroir renvoie une deuxième onde, appelée dans la suite onde réfléchie, qui se propage dans la direction de l'axe (Ox) dans le sens négatif. L'onde réfléchie a même amplitude S_0 que l'onde incidente et elle présente en $x = 0$ un retard de phase φ en O par rapport à l'onde incidente.



On dispose sur le trajet des ondes un film photosensible plan très mince, parfaitement transparent durant l'expérience, incliné par rapport au plan d'un angle très petit ($\epsilon = 1.10^{-3}$ rad). Une fois développé, le film photosensible montre une alternance régulière de bandes claires (qui ont reçu une amplitude lumineuse forte) et sombres (qui ont reçu une amplitude lumineuse nulle), parallèles à la direction (Oy) . La distance entre deux bandes sombres est $D = 0,27$ mm. On constate également que le bord du film qui était en contact avec le conducteur est sombre.

1. Exprimer les signaux $s_i(x,t)$ et $s_r(x,t)$ des ondes incidente et réfléchie respectivement, en fonction de S_0 , λ , c (célérité de la lumière), φ , x et t . On supposera que la phase initiale de l'onde incidente est nulle en $x = 0$. Que peut-on dire de l'onde totale $s(x,t) = s_i(x,t) + s_r(x,t)$?
2. Expliquer l'aspect de la plaque photographique après développement. Quelle relation existe entre D , λ et ϵ ? Déterminer numériquement la longueur d'onde utilisée dans cette expérience.
3. Montrer que l'expérience n'est concluante que si le film photographique est très fin (on donnera un ordre de grandeur de l'épaisseur maximale acceptable) et commenter.
4. L'expérience a été menée en 1890 par Wiener afin de déterminer si la plaque photographique est sensible au champ électrique ou au champ magnétique de l'onde lumineuse. La théorie électromagnétique montre que pour le champ électrique le déphasage à la réflexion est $\varphi = \pi$ tandis que pour le champ magnétique il est $\varphi = 0$. Quelle est la conclusion de l'expérience ?

CHAPITRE 4 – ONDE LUMINEUSE

4.12 Décalage vers le rouge (★★)

Soit un détecteur D fixe en un point O de l'espace, une source S se déplace à la vitesse algébrique V sur un axe (Ox) orienté du détecteur vers la source. L'abscisse $x(t)$ représente la distance entre D et S . La source S émet un signal de période T ; la célérité du signal dans le milieu qui sépare S de D est c .

1. Soit t (respectivement t') et $t + T$ (respectivement $t' + T'$) les instants correspondant à l'émission par la source (respectivement la réception par le détecteur) du début et de la fin d'une période du signal. Calculer t' et $t' + T'$ en fonction des données.
2. Si on suppose que le temps caractéristique de variation de la vitesse de la source est très grand devant T , exprimer T' en fonction de T , V et c . Dans le cas où $V \ll c$, exprimer ν' la fréquence mesurée par le détecteur en fonction de la fréquence émise ν , de V et de c .
3. On constate que le spectre de la lumière reçue des étoiles est décalé vers le rouge par rapport à celui que l'on obtiendrait sur Terre avec une source composée d'atomes identiques à ceux constituant les étoiles observées. Quelle information peut-on en déduire ?

CORRIGÉS

4.1 Doublage de fréquence

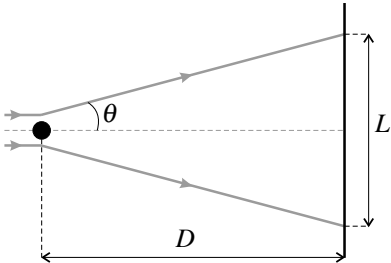
1. La relation entre la longueur d'onde λ et la fréquence f est $\lambda = \frac{c}{f}$ où c est la vitesse de la lumière. Si la fréquence est doublée, la longueur d'onde est divisée par deux. Donc : $\lambda_2 = \frac{1}{2}\lambda_1 = 532 \text{ nm}$. Cette longueur d'onde correspond à une couleur verte.
2. La période de l'onde est $T_2 = \frac{\lambda_2}{c} = 1,77 \cdot 10^{-15} \text{ s}$. Le nombre d'oscillations dans une impulsion durant $\tau = 0,1 \cdot 10^{-9} \text{ s}$ est $\frac{\tau}{T_2} \simeq 6 \cdot 10^4$.

4.2 Cavité optique

1. Le mode propre est une onde stationnaire. Le signal étant nul sur les miroirs, cela veut dire qu'il y a un nœud de vibration sur chaque miroir. La distance entre deux nœuds successifs étant $\frac{\lambda}{2}$ on en déduit que $L = n \frac{\lambda}{2}$ où n est un entier (c'est le nombre de nœuds entre les miroirs plus 1). Ainsi : $L = n \frac{c}{2f_n}$ où f_n est la fréquence du mode propre d'ordre n , soit : $f_n = n \frac{c}{2L}$.
2. Le domaine visible en longueur d'onde s'étend de $\lambda_{\min} = 400 \text{ nm}$ à $\lambda_{\max} = 750 \text{ nm}$. La longueur d'onde $\lambda_n = \frac{2L}{n}$ est dans le domaine visible si $\frac{2L}{\lambda_{\max}} \leq n \leq \frac{2L}{\lambda_{\min}}$ soit si $1,3 \cdot 10^5 \leq n \leq 2,5 \cdot 10^5$.
3. La différence entre deux fréquences propres de la cavité est $\Delta f = \frac{c}{2L} = 3 \cdot 10^9 \text{ Hz}$. C'est cinq ordres de grandeurs en dessous de l'ordre de grandeur des fréquences optiques ($10^{14} - 10^{15} \text{ Hz}$). C'est pourquoi on a trouvé de très grandes valeurs pour n à la question précédente. Ainsi, la condition de résonance de la cavité ne porte que sur le 5^e chiffre significatif de la valeur de la fréquence. La fréquence de l'onde laser est déterminée par la différence entre les niveaux d'énergie des atomes qui se trouvent dans la cavité laser.

4.3 Mesure du diamètre d'un cheveu

On peut mesurer le diamètre d'un cheveu en faisant diffracter la lumière laser par le cheveu. Le cheveu de diamètre d diffracte comme une fente de même largeur. En mettant l'écran à distance D derrière le cheveu on voit une figure de diffraction dont la tache centrale, la plus lumineuse, a pour largeur $L = 2D \tan \theta$ où θ est le demi-angle d'ouverture du faisceau du faisceau diffracté (voir figure).



CHAPITRE 4 – ONDE LUMINEUSE

La relation $\sin \theta = \frac{\lambda}{d} \sim \frac{633 \cdot 10^{-9}}{100 \cdot 10^{-6}} \sim 10^{-3}$ montre que $\sin \theta \ll 1$ donc $\theta \ll 1$ et $\tan \theta \simeq \theta \simeq \sin \theta = \frac{\lambda}{d}$. Ainsi : $L \simeq \frac{2\lambda D}{d}$ soit $d \simeq \frac{2\lambda D}{L}$.

Il faut donc mesurer L avec la règle graduée. L'incertitude sur cette valeur est au mieux $\Delta L \simeq 0,3 \text{ mm}$ (plus grande si les conditions notamment d'obscurité ne sont pas très bonnes). Avec $\theta \sim 10^{-3}$ et $D \sim 1,5 \text{ m}$ on trouve $L \sim 2 \text{ cm}$ donc l'incertitude relative sur L est : $\frac{\Delta L}{L} \simeq 1,5\%$.

La distance D est mesurée avec une incertitude relative guère inférieure à 1%. La longueur d'onde du laser est connue avec une précision bien inférieure à 1%.

Finalement, l'incertitude relative sur le diamètre du cheveu est : $\frac{\Delta d}{d} = \frac{\Delta L}{L} + \frac{\Delta D}{D} \simeq 3\%$.

On peut augmenter sensiblement la précision en faisant la mesure de L pour différentes valeurs de D . En exploitant n mesures on peut diviser cette incertitude par $\sqrt{n}$.

4.4 Faisceau laser

1. La relation entre λ et w_0 implique que plus le laser est étroit à son minimum de section, plus son angle de divergence est grand. On reconnaît le phénomène de diffraction. La relation est analogue à la relation $\sin \theta = \frac{\lambda}{a}$ de la diffraction par une fente.

2. Pour visualiser la géométrie du faisceau il faut calculer $L_f = \frac{(0,15 \cdot 10^{-3})^2}{633 \cdot 10^{-9}} = 3,5 \cdot 10^{-2} \text{ m}$ et l'angle d'ouverture $\theta = \frac{633 \cdot 10^{-9}}{\pi 0,15 \cdot 10^{-3}} = 1,34 \cdot 10^{-3} \text{ rad}$.

La distance à l'écran $D = 2 \text{ m}$ est grande devant L_f donc il se trouve dans la zone conique du faisceau. La largeur de la tache vue sur l'écran est approximativement :

$$D\theta = 2 \times 1,34 \cdot 10^{-3} = 2,7 \text{ mm}.$$

3. λ_1 correspond à une lumière du proche infrarouge et λ_2 à une lumière bleue. La relation donnée par l'énoncé implique $w_0 = \sqrt{\lambda L_f}$. On trouve pour le laser infrarouge $w_{0,1} = 1,5 \cdot 10^{-4} \text{ m}$ et pour le laser bleu $w_{0,2} = 1,1 \cdot 10^{-4} \text{ m}$. Le laser de plus petite longueur d'onde, avec son waist plus petit, permet de lire des détails plus petits sur un disque optique. On met plus d'information sur un *blue ray* que sur un CD.

4.5 Le laser-Lune

1. La durée de l'aller-retour de la lumière entre la surface de la Terre et la surface de la Lune est : $\tau = \frac{2D}{c} = 2,51 \text{ s}$. L'incertitude relative sur cette mesure est : $\frac{\Delta \tau}{\tau} = 2 \cdot 10^{-11}$. Une telle précision n'est obtenue qu'avec une horloge atomique.

2. D’après l’énoncé le demi-angle d’ouverture θ du faisceau est tel que :

$$\sin \theta = 1,22 \frac{\lambda}{a} = \frac{532 \cdot 10^{-6}}{1,5} = 4,3 \cdot 10^{-7}.$$

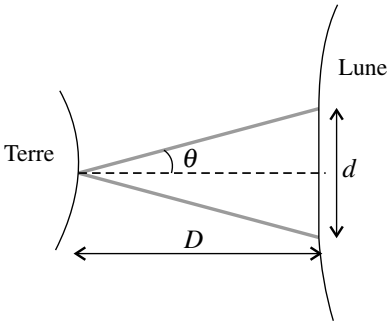
La petitesse du sinus montre que θ est très petit aussi. La tache faite par ce faisceau à la surface de la Lune est un disque de diamètre : $d = a + 2D \tan \theta$. L’angle θ étant très faible, $\tan \theta \simeq \sin \theta = 1,22 \frac{\lambda}{a}$, de plus on peut négliger a dans l’expression de d .

Finalemt : $d \simeq 2D \times 1,22 \frac{\lambda}{a} \simeq 320 \text{ m}.$

3. La fraction de l’énergie lumineuse émise depuis la Terre et reçue par le rétroréflexeur est :

$$\rho = \frac{\ell^2}{\pi d^2/4} = 1,2 \cdot 10^{-5}.$$

4. La lumière qui revient sur Terre est aussi diffractée, avec un demi-angle d’ouverture de l’ordre de $\frac{\lambda}{\ell}$. Puisque ℓ est du même ordre que a , la perte due à la diffraction est du même ordre de grandeur au retour qu’à l’aller. Si la diffraction était la seule cause des pertes on devrait recevoir une fraction $10^{-5} \times 10^{-5} = 10^{-10}$ des photons. On en reçoit beaucoup moins donc il y a donc d’autres causes de pertes, notamment la perturbation de l’atmosphère terrestre qui dévie légèrement le faisceau, l’absorption par l’atmosphère, etc.



4.6 Fentes de Young

1. L’onde laser est diffractée par la fente ouverte : le faisceau prend une extension angulaire caractérisée par l’angle θ tel que :

$$\sin \theta = \frac{\lambda}{\varepsilon} = 9,0 \cdot 10^{-4}.$$

La tache centrale de diffraction sur l’écran a pour largeur :

$$L = 2D \tan \theta \simeq 2D \sin \theta = \frac{2\lambda D}{a} = 2,7 \text{ cm}.$$

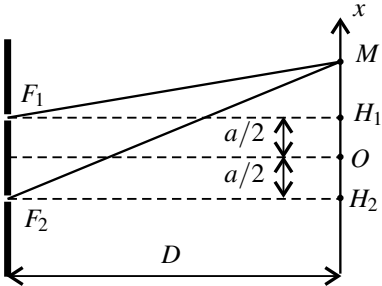
C’est précisément la dimension de la tache centrale observée sur l’écran.

Les taches centrales de diffraction correspondant aux deux fentes sont décalées de $a = 0,40 \text{ mm}$.

On a : $\frac{L}{a} = 67,5 \gg 1$, donc les deux taches sont pratiquement confondues.

2. En utilisant le théorème de Pythagore (voir figure) il vient :

$$\begin{aligned} F_2M^2 - F_1M^2 &= (F_2H_2^2 + H_2M^2) - (F_1H_2^2 + H_1M^2) \\ &= H_2M^2 - H_1M^2 = \left(x + \frac{a}{2}\right)^2 - \left(x - \frac{a}{2}\right)^2 = ax. \end{aligned}$$



CHAPITRE 4 – ONDE LUMINEUSE

D'autre part :

$$F_2M^2 - F_1M^2 = (F_2M + F_1M)(F_2M - F_1M) \simeq 2D(F_2M - F_1M),$$

en faisant l'approximation $F_1M + F_2M \simeq 2D$ justifiée par le fait que $D \gg a$ et $D \gg L > OM$. En comparant les deux expressions on trouve :

$$F_2M - F_1M \simeq \frac{ax}{D}.$$

3. Les deux ondes partent en phase des fentes mais les durées de propagation entre les fentes et M diffèrent de $\tau = \frac{1}{c}(F_2M - F_1M) = \frac{ax}{cD}$. Par suite les deux ondes arrivant en M ont un déphasage $\Delta\phi = \omega\tau = 2\pi f \frac{ax}{cD} = \frac{2\pi ax}{\lambda D}$.

4. L'interférence est constructive si $\Delta\phi = 2n\pi$ où n est un entier. Ceci s'écrit $x = n \frac{\lambda D}{a}$, soit numériquement $x = n \times 2,4$ mm. Les valeurs de l'entier n sont limitées car M est dans le champ d'interférences. La condition $-\frac{L}{2} \leq x \leq \frac{L}{2}$ devient, si l'on tient compte de l'expression de L , $-\frac{a}{\varepsilon} \leq n \leq \frac{a}{\varepsilon}$ soit numériquement $-5,7 \leq n \leq 5,7$. Il y a donc dans le champ d'interférences 11 points en lesquels l'interférence est constructive, c'est-à-dire où l'éclairement est maximal. Ce sont les milieux des 11 taches claires que l'on compte dans le champ d'interférences sur la photo.

5. L'interférence est destructive si $\Delta\phi = (2n+1)\pi$ où n est un entier. Ceci s'écrit $x = \left(n + \frac{1}{2}\right) \frac{\lambda D}{a}$. La distance entre deux de ces points est $\frac{\lambda D}{a} = 2,4$ mm. Ce résultat est en accord avec la distance que l'on peut mesurer sur la photographie (qui est en vraie grandeur) entre deux points plus sombres.

4.7 Redistribution sélective de la lumière

1. Cette relation traduit la conservation de l'énergie .
2. La longueur d'onde la plus courte du domaine optique valant 400 nm, la taille maximale des aspérités pour une surface optique polie est 40 nm.
3. Un tel objet ne réfléchit pas la lumière : $R(\lambda) = 0$ quel que soit λ . Il apparaît donc comme noir.
4. La plante verte réfléchit au moins une partie des radiations vertes donc elles ne les utilise pas toutes pour la photosynthèse. (La chlorophylle absorbe fortement dans le bleu et dans le rouge mais pas dans le vert.)
5. Le tissu est bleu parce qu'il réfléchit uniquement des longueurs d'onde bleues. La lumière du néon ne comportant pas ces longueurs d'onde dans son spectre, le tissu paraît noir sous cet éclairage.

4.8 Absorption de la lumière par l'eau

1. L'absorption est d'autant plus forte que $\alpha(\lambda)$ est grand. Les longueurs d'onde les plus absorbées par l'eau sont au dessus de 600 nm et correspondent à la couleur rouge.

Pour $\lambda_1 \simeq 710$ nm (plus grande longueur d'onde sur le graphe), $\alpha(\lambda_1) \simeq 1 \text{ m}^{-1}$. Pour une épaisseur d'eau traversée de 10 cm, $T(\lambda_1) \simeq \exp(-0,1) = 0,90$ et pour une épaisseur de 2 m, $T(\lambda_1) \simeq \exp(-2) = 0,14$.

Pour les longueurs d'onde les plus faibles (couleur bleu), l'eau n'absorbe pas : $\alpha(\lambda) \simeq 0$ $T(\lambda) \simeq 1$.

2. Les parties rouges du poisson réfléchissent dans le rouge. Mais à quelques mètres de profondeur la lumière du soleil est fortement appauvrie en rouge, elles apparaîtront donc noires sur la photo. Les parties blanches réfléchissent tout le spectre, elles apparaîtront vertes (couleur complémentaire du rouge absent de la lumière reçue par le poisson). Ainsi, le poisson sera noir et vert sur la photo en lumière solaire.

Si on utilise un flash, on éclaire le poisson avec un spectre complet (on peut négliger l'absorption par la dizaine de centimètres d'eau entre l'appareil et le poisson) et le poisson aura ses véritables couleurs sur la photo.

4.9 Réflectivité d'un métal

1. Le meilleur métal pour faire un miroir est l'argent car il a le plus fort pouvoir réflecteur et ce pouvoir réflecteur est assez uniforme sur le spectre visible. Le coefficient de réflexion moyen est environ de 0,9.

2. L'or ne réfléchit pas bien les longueur d'onde les plus petites du spectre visible (entre 400 et 450 nm). Il réfléchit donc peu la lumière bleue, ce qui explique sa couleur jaune (couleur complémentaire du bleu).

4.10 Couche anti-reflet

1. a. La longueur d'onde dépend du milieu transparent dans lequel la lumière se propage. Ainsi, l'onde incidente et l'onde réfléchie ont la même longueur d'onde $\lambda_1 = \frac{c}{n_1 f}$ et l'onde transmise a une longueur d'onde différente : $\lambda_2 = \frac{c}{n_2 f}$. La couleur est associée à la fréquence de l'onde. Elle est la même pour les trois ondes.

b. L'onde incidente se propage à la vitesse $\frac{c}{n_1}$ dans le sens positif de l'axe (Ox). On a donc :

$$s_i(x,t) = s_i\left(0,t - \frac{n_1 x}{c}\right) = A_i \cos\left(2\pi f t - 2\pi \frac{n_1 x}{c} - \varphi_i\right).$$

c. L'onde réfléchie se propage à la vitesse $\frac{c}{n_1}$ dans le sens négatif de l'axe (Ox), donc :

$$s_r(x,t) = s_r\left(0,t + \frac{n_1 x}{c}\right) = \frac{n_1 - n_2}{n_1 + n_2} A_i \cos\left(2\pi f t + 2\pi \frac{n_1 x}{c} - \varphi_i\right).$$

L'amplitude de l'onde réfléchie est $A_r = \frac{|n_1 - n_2|}{n_1 + n_2} A_i$ et sa phase initiale à l'origine est φ_i si $n_1 > n_2$ et $\varphi_i + \pi$ si $n_1 < n_2$. La réflexion sur un milieu plus réfringent que le milieu de l'onde incidente provoque un déphasage de π sur l'onde réfléchie.

CHAPITRE 4 – ONDE LUMINEUSE

2. a. $R + T = \frac{(n_1 - n_2)^2 + 4n_1n_2}{(n_1 + n_2)^2} = 1$. Cette relation traduit la conservation de l'énergie lumineuse qui est tout entière dans l'onde réfléchie et dans l'onde transmise.

b. L'application numérique donne $R = 4,26 \cdot 10^{-2}$. Le rapport entre l'éclairement du reflet du soleil dans les lunettes et l'éclairement du paysage est : $10R = 0,426$. Le reflet du soleil n'est pas négligeable et peut être gênant pour le conducteur.

3. a. L'aller retour de l'onde réfléchie 2 dans le matériau d'indice n la retarde par rapport à l'onde réfléchie 1 de $\tau = \frac{2e}{c/n}$. Il n'y a pas d'autre cause de déphasage entre les deux ondes.

En effet :

- les deux ondes subissent un déphasage de π puisqu'elles se réfléchissent sur des milieux d'indice plus fort que leur milieu de propagation ($n > 1$ et $N > n$) ;
- la transmission du milieu d'indice n à l'air ne change pas la phase de l'onde réfléchie 2.

Ainsi le déphasage entre les deux ondes est :

$$\Delta\varphi = 2\pi f\tau = 2\pi f \frac{ne}{c} = \frac{4\pi ne}{\lambda}.$$

b. Pour avoir une interférence destructrice entre les ondes réfléchies 1 et 2, il faut que $\Delta\varphi = (2m + 1)\pi$ où m est un entier quelconque, soit que : $e = \left(m + \frac{1}{2}\right) \frac{\lambda}{2n}$.

c. La plus petite valeur de e possible est $\frac{\lambda}{4n} = 90,6$ nm avec les valeurs. Cette épaisseur correspond à une centaine de couches atomiques. La couche antireflet est fragile. On peut la fabriquer avec la technique d'évaporation sous vide.

4. a. Le coefficient de réflexion est $R = \frac{A^2}{A_i^2}$ où A_i est l'amplitude de l'onde incidente et A l'amplitude de l'onde résultant de l'interférence des ondes réfléchies 1 et 2. A est donnée par la formule des interférences (voir chapitre précédent) :

$$A^2 = A_1^2 + A_2^2 + 2A_1A_2 \cos(\Delta\varphi),$$

en notant A_1 et A_2 les amplitudes respectives des ondes réfléchies 1 et 2 et $\Delta\varphi$ leur déphasage. Le déphasage dépendant de la longueur d'onde de la lumière, A^2 et R en dépendent aussi.

b. L'interférence est destructrice pour $\lambda = \lambda_m$ mais pas totalement destructrice parce que les deux ondes n'ont pas la même amplitude ($R(\lambda_m)$ est non nul). Si l'on compare au résultat de la question 2., la couche antireflet divise par trois le coefficient de réflexion ce qui est satisfaisant.

c. À la question 2. on a calculé $\frac{A_1^2}{A_i^2} = 0,042$. La valeur de $R(\lambda_m)$ donne $\frac{(A_1 - A_2)^2}{A_i^2} = 0,015$. On en tire $\frac{A_2^2}{A_i^2} = (\sqrt{0,042} - \sqrt{0,015})^2 = 0,007$. Par ailleurs, $\Delta\varphi = \pi \frac{\lambda_m}{\lambda}$. Finalement :

$$R(\lambda) = 0,049 + 0,034 \cos\left(\pi \frac{\lambda_m}{\lambda}\right).$$

On a d'après cette formule : $R(400 \text{ nm}) = 0,025$ et $R(750 \text{ nm}) = 0,032$.

4.11 Expérience de Wiener

1. L'onde incidente se propage dans le sens positif de l'axe (Ox) . Elle a une amplitude S_0 , une longueur d'onde λ et une phase initiale nulle en $x = 0$ donc s'écrit :

$$s_i(x,t) = S_0 \cos \left(\frac{2\pi}{\lambda}(x - ct) \right).$$

L'onde réfléchie se propage dans le sens négatif de l'axe (Ox) . Elle a une amplitude S_0 , une longueur d'onde λ et une phase initiale $-\varphi$ en $x = 0$ donc s'écrit :

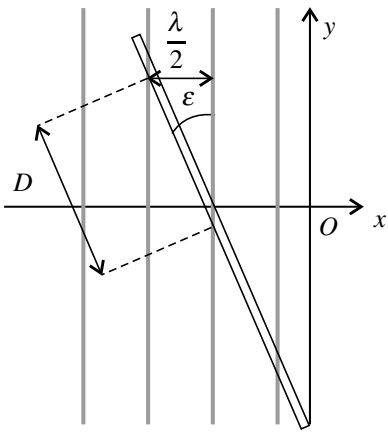
$$s_r(x,t) = S_0 \cos \left(\frac{2\pi}{\lambda}(x + ct) - \varphi \right).$$

L'onde résultante $s(x,t) = s_i(x,t) + s_r(x,t)$ est la superposition de deux ondes progressives de même amplitude se propageant dans la même direction mais en sens opposés : il s'agit d'une onde stationnaire. On peut l'écrire :

$$\begin{aligned} s(x,t) &= S_0 \cos \left(\frac{2\pi}{\lambda}(x - ct) \right) + S_0 \cos \left(\frac{2\pi}{\lambda}(x + ct) - \varphi \right) \\ &= 2S_0 \cos \left(\frac{2\pi}{\lambda}x - \frac{\varphi}{2} \right) \cos \left(\frac{2\pi c}{\lambda}t - \frac{\varphi}{2} \right). \end{aligned}$$

2. Après le développement, les points de la plaque qui ont été sur un ventre de vibration sont clairs et les points qui ont été sur un nœud de vibration sont sombres.

Les ventres et les nœuds sont sur des plans d'équations $x = \text{constante}$ perpendiculaires à l'axe (Ox) . Les plans ventraux sont représentés en gris sur la figure ci-contre. La distance entre deux plans ventraux est $\frac{\lambda}{2}$. La distance entre deux franges claires sur la plaque développée est d'après la figure : $D = \frac{\lambda/2}{\sin \epsilon} \simeq \frac{\lambda}{2\epsilon}$. Avec les valeurs numériques données, on trouve : $\lambda = 2D\epsilon \simeq 2 \times 1.10^{-3} \times 0.27.10^{-3} = 0,54.10^{-6} \text{ m}$. C'est bien une longueur d'onde du domaine visible.



3. On voit bien sur la figure que les traits clairs ont une largeur provenant de l'épaisseur de la plaque photographique. Si e est cette épaisseur, l'épaisseur des traits clairs est $\frac{e}{\sin \epsilon}$. Il faut qu'elle reste inférieure à, par exemple, $\frac{D}{10}$ soit que $\frac{e}{\sin \epsilon} < \frac{1}{10} \frac{\lambda}{2 \sin \epsilon}$ soit que $e < \frac{\lambda}{20} \sim 25 \text{ nm}$. La plaque photographique doit être très fine.

4. L'amplitude de l'onde stationnaire en $x = 0$ est $\mathcal{A}(0) = 2S_0 \left| \cos \left(\frac{\varphi}{2} \right) \right|$. Elle est nulle puisque le bord du film n'a pas été impressionné au cours de l'expérience. Cela correspond à

CHAPITRE 4 – ONDE LUMINEUSE

$\varphi = \pi$ et non à $\varphi = 0$. Donc le signal optique auquel le film photographique est sensible est le champ électrique.

4.12 Effet Doppler, décalage vers le rouge

1. Le signal de début de période, émis à l'instant t , doit parcourir jusqu'au détecteur la distance $x(t)$; il est donc reçu à l'instant : $t' = t + \frac{x(t)}{c}$. Le signal de fin de période, émis à l'instant $t + T$, doit parcourir jusqu'au détecteur la distance $x(t + T)$; il est donc reçu à l'instant : $t' + T' = t + T + \frac{x(t + T)}{c}$.

2. Si la vitesse V de la source varie peu entre les instants t et $t + T$, on a : $x(t + T) \simeq x(t) + VT$. On a alors :

$$T' = \left(t + T + \frac{x(t + T)}{c} \right) - \left(t + \frac{x(t)}{c} \right) = T + \frac{1}{c}(x(t + T) - x(t)) \simeq \left(1 + \frac{V}{c} \right) T.$$

On en déduit la fréquence f' mesurée par le détecteur :

$$f' = \frac{1}{T'} = \left(1 + \frac{V}{c} \right)^{-1} \frac{1}{T} = \left(1 + \frac{V}{c} \right)^{-1} f \simeq \left(1 - \frac{V}{c} \right) f,$$

en faisant l'approximation $(1 + \varepsilon)^{-1} \simeq 1 - \varepsilon$ pour $\varepsilon = \frac{V}{c} \ll 1$. Ce changement de fréquence est une manifestation de l'effet Doppler.

3. Le « décalage vers le rouge », nommé *redshift* en anglais, correspond à une augmentation des longueurs d'onde, par l'effet Doppler, donc à une diminution des fréquences. Le calcul précédent a montré que c'est ce qui se passe si la distance entre la source et le détecteur augmente au cours du temps. On peut donc déduire de cette observation que les étoiles lointaines s'éloignent de nous, phénomène dû à l'expansion de l'Univers.

Optique géométrique

5

L'objet de l'optique géométrique est la détermination géométrique des rayons lumineux. Ses applications sont les différents instruments d'optique : loupe, appareil photo, lunette astronomique...

1 Approximation de l'optique géométrique

En optique géométrique on étudie les rayons lumineux sans faire référence à l'onde lumineuse. On fait l'hypothèse que ces rayons sont indépendants entre eux.

Ceci est contredit par le phénomène de diffraction. Par exemple un rayon lumineux traversant une fente diffractante est dévié parce que des rayons voisins sont arrêtés par la fente. En optique géométrique on néglige le phénomène de diffraction qui ne concerne que des rayons lumineux « à la marge ». Il faut pour cela que les ouvertures traversées par la lumière aient toutes une dimension a supérieure à 1000λ où λ est la longueur d'onde de la lumière. La déviation des rayons lumineux due à la diffraction, de l'ordre de $\frac{\lambda}{a}$, sera dans ce cas négligeable.

L'approximation de l'optique géométrique consiste à négliger tout phénomène de diffraction. Cela revient à considérer que la longueur d'onde λ de la lumière est quasiment nulle.

Exemple

Le diamètre de l'objectif d'un téléphone portable est à peu près $d = 2 \text{ mm}$. La longueur d'onde λ de la lumière visible étant de l'ordre de $0,5 \mu\text{m}$, on a $\frac{d}{\lambda} \sim 4000$. On peut appliquer l'optique géométrique dans ce cas.

Un grand télescope a un diamètre $d = 4,2 \text{ m}$. Avec une telle dimension il est naturel de vouloir négliger la diffraction. Cependant, on souhaite distinguer des étoiles situées dans des directions séparées d'un angle très faible. La dispersion angulaire due à la diffraction, $\theta \simeq \frac{\lambda}{d} \sim 10^{-7} \text{ rad}$ pour $\lambda = 0,5 \mu\text{m}$, limite le pouvoir de résolution de l'instrument.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

2 Lois de Descartes

Les **lois de Descartes** sont les lois fondamentales de l’optique géométrique.

2.1 Lois de Descartes pour la réflexion

a) Réflexion d’un rayon lumineux

La **réflexion** consiste en un changement de direction d’un rayon lumineux « rebondissant » sur une surface réfléchissante (voir figure 5.1). On distingue le **rayon incident**, rayon lumineux avant la réflexion, et le **rayon réfléchi**, rayon lumineux après la réflexion.

La surface réfléchissante doit être parfaitement lisse à l’échelle de la longueur d’onde de la lumière. Il s’agit le plus souvent d’une surface métallique polie. L’usage est d’appeler cette surface miroir et de la représenter avec des hachures à l’arrière.

Le point I où le rayon incident rencontre la surface réfléchissante est appelé **point d’incidence**. La droite passant par I , orthogonale à la surface réfléchissante, est appelée normale en I . Le plan contenant le rayon incident et la normale en I est le **plan d’incidence** (voir figure 5.1).

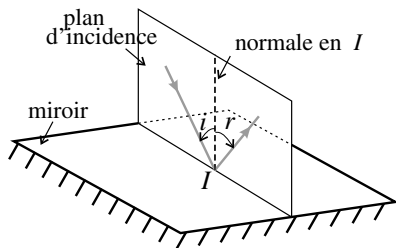


Figure 5.1 – Réflexion sur un miroir.

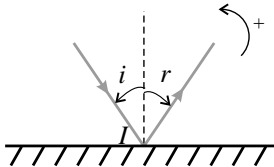


Figure 5.2 – Représentation de la réflexion dans le plan d’incidence.

b) Lois de la réflexion

Les angles entre la normale et les rayons incident et réfléchi sont notés i et r et appelés respectivement **angle d’incidence** et **angle de réflexion**. Ces angles doivent être algébrisés (voir annexe mathématique). Il est très important de les prendre toujours de la normale vers les rayons. Leurs signes dépendent du sens positif choisi et sont opposés.

Les **lois de Descartes pour la réflexion** s’énoncent ainsi :

- 1. le rayon réfléchi est compris dans le plan d’incidence,
- 2. les angles d’incidence i et de réflexion r , sont liés par la relation :

$$r = -i. \tag{5.1}$$

2.2 Lois de Descartes pour la réfraction

a) Réflexion et réfraction sur un dioptre

La surface de séparation entre les deux milieux transparents est appelée en optique **dioptre**. Un rayon lumineux arrivant sur un dioptre, appelé rayon incident, donne en général naissance à deux rayons (voir figure 5.3) :

- un rayon réfléchi qui repart dans le milieu du rayon incident,
- un rayon transmis qui entre dans l'autre milieu transparent.

Le rayon incident suit les lois de Descartes données au paragraphe précédent. Le rayon transmis n'est pas dans le prolongement du rayon incident. C'est le phénomène de **réfraction**. Le rayon transmis est aussi appelé **rayon réfracté**.

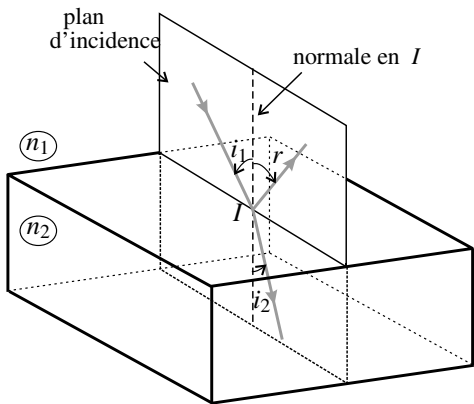


Figure 5.3 – Réflexion et réfraction sur un dioptre.

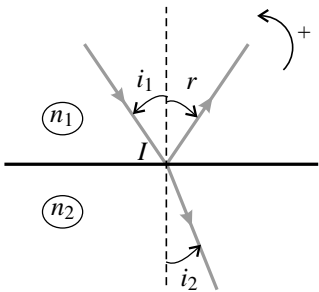


Figure 5.4 – Représentation dans le plan d'incidence.

b) Lois de la réfraction

On note n_1 l'indice optique du milieu de propagation du rayon incident et n_2 l'indice optique du milieu de propagation du rayon réfracté. Les angles que font ces deux rayons avec la normale au dioptre sont notés i_1 et i_2 et appelés respectivement angle d'incidence et angle de réfraction. Ces angles sont algébriques et doivent toujours aller de la normale vers les rayons. Leur signes dépendent du sens positif choisi et sont identiques.

Les **lois de Descartes pour la réfraction** s'énoncent ainsi :

1. le rayon réfracté est dans le plan d'incidence,
2. l'angle d'incidence i_1 et l'angle de réfraction i_2 sont liés par la relation :

$$n_1 \sin i_1 = n_2 \sin i_2. \tag{5.2}$$

c) Déviation due à la réfraction

Le sens de cette déviation du rayon transmis dépend de l'ordre des indices de deux milieux :

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

- Si deuxième milieu est plus réfringent que le premier, c'est-à-dire si $n_2 > n_1$, le rayon réfracté se rapproche de la normale (voir figure 5.5).
- Si le deuxième milieu est moins réfringent que le premier, soit si $n_2 < n_1$, le rayon réfracté s'écarte de la normale (voir figure 5.6).

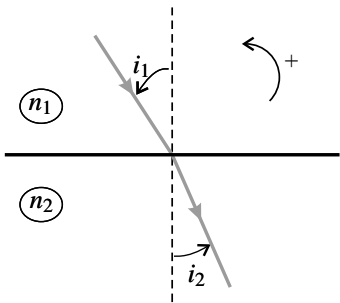


Figure 5.5 – Réfraction avec un milieu (2) plus réfringent.

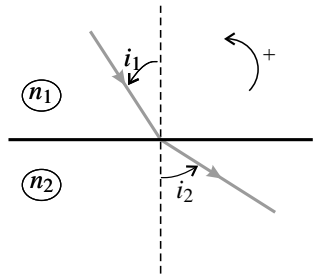


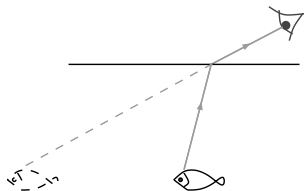
Figure 5.6 – Réfraction avec un milieu (2) moins réfringent.



Il faut veiller à respecter cette règle lorsqu'on dessine un rayon réfracté.

Exemple

Un pêcheur utilisant un harpon observe un poisson. Il voit le poisson dans le prolongement des rayons lumineux arrivant dans son œil mais, à cause de la réfraction, l'animal est en fait plus proche de lui. Pour le toucher il faut donc qu'il vise plus près que l'endroit où il voit le poisson.



d) Phénomène de réflexion totale

Dans le cas où $n_2 < n_1$, $\sin i_2 = \frac{n_1}{n_2} \sin i_1 > \sin i_1$. Si $\sin i_1 = \frac{n_2}{n_1}$, alors $\sin i_2 = 1$ ce qui veut dire que $i_2 = \frac{\pi}{2}$. Si $\sin i_1 > \frac{n_2}{n_1}$ la loi de Descartes donnerait $\sin i_2 > 1$ ce qui est impossible.

Dans le cas où le deuxième milieu est moins réfringent, le rayon réfracté n'existe que si l'angle d'incidence i_1 est inférieur à l'angle R_{lim} , appelé **angle de réfraction limite** défini par :

$$\sin R_{\text{lim}} = \frac{n_2}{n_1}.$$

Si i_1 est supérieur à R_{lim} , il n'y a plus qu'un rayon réfléchi, c'est le phénomène de **réflexion totale**.

Exemple

1. L'angle limite de réfraction pour que la lumière sorte d'un verre d'indice $n_1 = 1,53$ dans l'air d'indice $n_2 = 1,00$ est : $R_{\text{lim}} = \arcsin\left(\frac{1,00}{1,53}\right) = 40,8^\circ$. Sur la figure 5.7, la face AB du prisme, dont la section est un triangle rectangle isocèle et dont l'indice est n_1 , se comporte donc comme un miroir pour la lumière puisque l'angle d'incidence dépasse l'angle limite de réfraction.
2. La lumière guidée par une fibre optique est « piégée » à l'intérieur de la fibre grâce au phénomène de réflexion totale : pour cela l'angle d'incidence i doit constamment dépasser l'angle de réfraction limite.

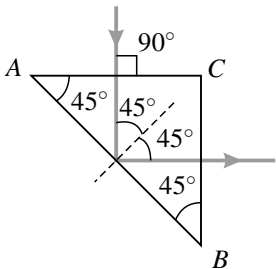


Figure 5.7 – Prisme à réflexion totale.

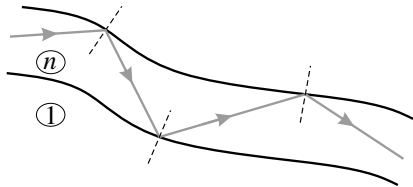


Figure 5.8 – Propagation dans une fibre optique.

3 Miroir plan

La suite de ce chapitre est consacrée à la formation d'images par des systèmes optiques. Un système optique est un ensemble de composants optiques (dioptries ou miroirs) rencontrés successivement par les rayons lumineux. Le système optique le plus simple est le miroir plan.

3.1 Image d'un point objet par un miroir plan

Un **miroir plan** est une surface réfléchissante plane. Comme chacun sait, on « se voit » dans un miroir. L'optique géométrique explique cette constatation en termes de rayons lumineux.

a) Image par un miroir plan d'un point objet réel

Un **point objet réel** est un point d'où partent les rayons lumineux qui arrivent sur un système optique.

On considère un point objet réel A situé devant un miroir plan. Un observateur reçoit les rayons lumineux issus de A après qu'ils ont été réfléchis par (M) (voir figure 5.9). La loi de la réflexion fait que ces rayons semblent provenir du point A' symétrique de A par rapport au plan de (M) (voir figure 5.9).

On va prouver ce résultat. Soit un rayon issu de A se réfléchissant sur (M) en I avec un angle

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

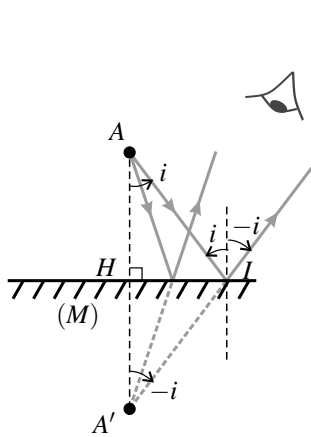


Figure 5.9 – Image d'un point objet réel par un miroir plan.

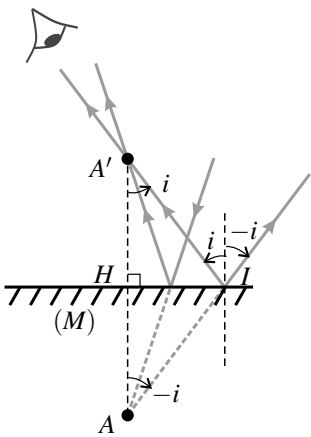


Figure 5.10 – Image d'un point objet virtuel par un miroir plan.

d'incidence quelconque i et soit A' le point d'intersection du prolongement du rayon réfléchi avec la droite normale au miroir passant par A . Sur la figure 5.9 on retrouve l'angle i en A par la propriété des angles alterne-interne et l'angle $-i$ en A' . Il vient alors :

$$\tan i = \frac{HI}{HA} = \frac{HI}{HA'} \quad \text{donc} \quad HA' = HA.$$

Ainsi, A' est le symétrique de A par rapport au plan du miroir. Il en est ainsi quel que soit l'angle d'incidence i donc tous les rayons réfléchis semblent venir de A' .

Pour l'œil qui reçoit les rayons réfléchis, tout se passe comme si ces rayons provenaient de A' . On dit que A' est l'image de A dans le miroir. C'est une image virtuelle.

Un **point image virtuel** A' est un point d'où semblent provenir les rayons lumineux sortant d'un système optique. Le prolongement des rayons au delà de la surface du système passe par A' .

b) Image par un miroir plan d'un point objet virtuel

Un **point objet virtuel** A est un point vers lequel convergent les rayons lumineux arrivant sur un système optique. Le prolongement des rayons au delà de la surface du système passe par A .

Sur la figure 5.10 le point A est un point objet virtuel pour le miroir (M) . La géométrie est la même que sur la figure 5.9 avec des rayons lumineux parcourus dans l'autre sens. Les rayons réfléchis passent tous par le symétrique A' de A par rapport au plan du miroir. On dit que A' est l'image de A par le miroir. C'est une image réelle.

Un **point image réel** est un point par lequel passent les rayons lumineux sortant d'un système optique.

En plaçant un écran perpendiculaire aux rayons en A' on voit un point lumineux net. Une image réelle peut être recueillie sur un écran.

3.2 Image d'un objet par un miroir plan

Un objet est un ensemble de points images. L'objet le plus simple que l'on puisse envisager est un vecteur $\overrightarrow{AB}$. Son image est le vecteur $\overrightarrow{A'B'}$ symétrique par rapport au miroir, soit un vecteur de même taille et de même sens. Le **grandissement** d'un miroir plan est égal à 1.

Un miroir plan donne de tout objet une image symétrique par rapport au plan du miroir.

4 Systèmes centrés et approximation de Gauss

4.1 Systèmes optiques centrés

Un **système optique centré** est un système optique dont les éléments constitutifs (dioptries, miroirs) ont un axe de symétrie commun. Cet axe est appelé **axe optique**. Il est orienté dans le sens de propagation de la lumière.

Les plans perpendiculaires à l'axe optique sont appelés **plans de front** (voir figure 5.12).

4.2 Approximation de Gauss

a) Exemple d'une lentille demi-boule

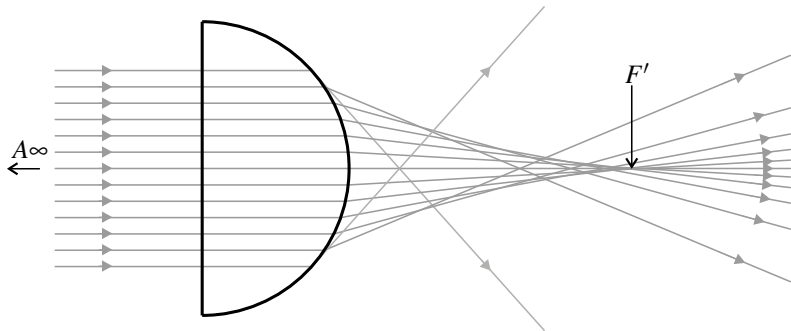


Figure 5.11 – Lentille demi-boule éclairée par un faisceau parallèle.

La figure 5.11 montre des **rayons incidents** parallèles traversant une demi-boule transparente d'indice $n = 1,5$, placée dans l'air d'indice $n_{\text{air}} = 1$. Les rayons sortant de la demi-boule sont les **rayons émergents**.

La demi-boule est un système centré. Les rayons incidents sont parallèles à son axe de symétrie, c'est-à-dire son axe optique. Ces rayons proviennent d'un point objet réel A situé à

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

distance infinie dans la direction des rayons. On constate que les rayons émergents ne passent pas tous par un point image. Cependant, les cinq rayons les plus proches de l'axe optique passent quasiment par le point F' indiqué par la flèche sur la figure. La demi-boule ne donne pas une image rigoureuse comme le miroir plan, mais une image approchée si l'on se restreint aux rayons proches de l'axe.

b) Rayons paraxiaux, conditions de Gauss

Un rayon incident sur un système centré est dit **paraxial** quand les deux conditions suivantes sont respectées :

- le rayon est proche de l'axe optique du système,
- le rayon est peu incliné par rapport à l'axe optique.

Un système centré est utilisé dans les **conditions de Gauss** si tous les rayons incidents sont des rayons paraxiaux.

4.3 Propriétés d'un système centré dans les conditions de Gauss

Les systèmes centrés s'utilisent toujours dans les conditions de Gauss. Ils ont alors les propriétés de **stigmatisme** et d'**aplanétisme** approchés

a) Stigmatisme approché

On admet le résultat suivant :

Un **système centré** utilisé **dans les conditions de Gauss** donne de tout point objet (pouvant être réel ou virtuel) une image ponctuelle approchée (pouvant être réelle ou virtuelle).

Ceci est représenté sur la figure 5.12 dans le cas de points objets réels et de points images réels. L'image de tout objet A appartenant à l'axe optique est un point A' appartenant aussi à l'axe optique.

Les points objet et image sont dits **conjugués** par le système optique.

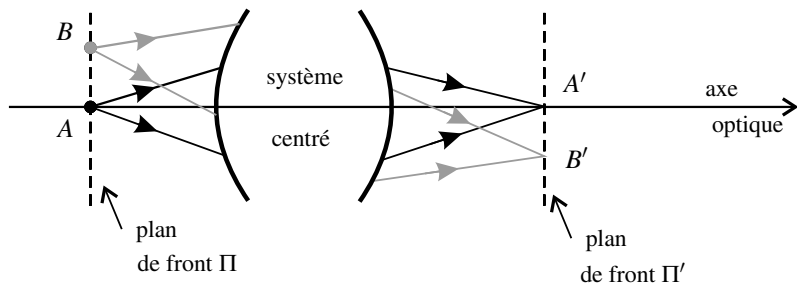


Figure 5.12 – Stigmatisme et aplanétisme dans les conditions de Gauss.

b) Aplanétisme approché

Un point B du plan de front passant par A a son image B' dans le plan de front passant par A' (voir figure 5.12).

Les plans de front Π et Π' passant par A et par A' sont conjugués par le système optique. De plus, par symétrie les vecteurs $\overrightarrow{A'B'}$ et $\overrightarrow{AB}$ sont colinéaires. On peut montrer qu'il existe un nombre sans dimension γ tel que $\overrightarrow{A'B'} = \gamma \overrightarrow{AB}$. γ est appelé **grandissement linéaire** et on le définit traditionnellement par la relation

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}}.$$

les mesures algébriques étant prises le long de deux axes orientés dans le même sens. Le grandissement γ est indépendant du point B mais *il dépend des plans Π et Π' conjugués*.

La signification physique du grandissement γ est la suivante :

- si $|\gamma| > 1$ l'image est plus grande que l'objet, si $|\gamma| < 1$ l'image est plus petite que l'objet ;
- si $\gamma > 0$ l'image est droite, c'est-à-dire de même sens que l'objet, si $\gamma < 0$ l'image est renversée par rapport à l'objet.

c) Le stigmatisme approché est-il suffisant ?

Les systèmes centrés dans les conditions de Gauss n'ont qu'un stigmatisme approché. Ainsi, lorsqu'on recueille l'image réelle d'un point sur un écran on obtient non un point, mais une tache lumineuse plus ou moins grosse. Est-ce que cela peut être satisfaisant ? Quelle taille maximale de la tache peut-on accepter ?

Pour répondre à ces questions il faut réfléchir à l'utilisation de l'image. Par exemple, dans le cas d'un appareil photographique, l'objectif forme l'image sur le capteur CCD qui est le récepteur photosensible de l'appareil. Ce capteur est une mosaïque de cellules élémentaires correspondant chacune à un pixel, c'est-à-dire un « point » sur l'image finale. Ces cellules ont une taille non nulle, de l'ordre de quelques microns. Le capteur ne fait pas la distinction entre un point parfait et une tache plus petite que la cellule élémentaire. En effet, si deux photons arrivent sur la même cellule élémentaire, ils sont identifiés comme arrivant au même point. Le stigmatisme approché de l'objectif est donc suffisant si la taille des taches images est inférieure à la taille d'une cellule élémentaire du capteur.

4.4 Foyers objet, foyers image

a) Point à l'infini

Un point à l'infini est un point situé à distance infinie du système. Les rayons passant par un point à l'infini sont parallèles entre eux. Le point est localisé à l'infini, dans la direction commune aux rayons.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

b) Foyers principaux

On appelle **foyer image principal** l'image du point objet situé à l'infini dans la direction de l'axe optique. Ce point est noté F' . Les rayons incidents parallèles à l'axe optique donnent des rayons émergents qui passent tous par F' (voir figure 5.13).

On appelle **foyer objet principal** le point objet dont l'image est située à l'infini dans la direction de l'axe optique. Ce point est noté F . Les rayons incidents issus de F donnent des rayons émergents tous parallèles à l'axe optique (voir figure 5.14).

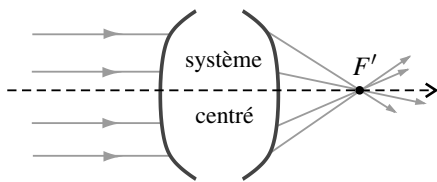


Figure 5.13 – Foyer image principal.

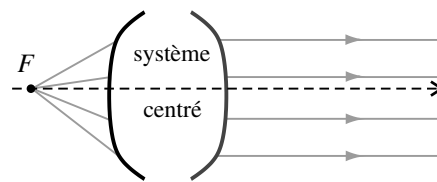


Figure 5.14 – Foyer objet principal.



Contrairement à ce que les notations pourraient laisser croire, le foyer principal image F' n'est pas le point conjugué du foyer image F !

c) Foyers secondaires

Par aplanétisme, tout point objet à l'infini mais hors de l'axe optique a son image dans le plan de front passant par F' . Ce plan est appelé **plan focal image**. Les points du plan focal image sont des **foyers image secondaires** et notés F'_s . Un faisceau incident de rayons parallèles entre eux, mais non parallèle à l'axe optique, est transformé en un faisceau convergent vers un foyer image secondaire (voir figure 5.15).

De même, on appelle **plan focal objet** le plan de front passant par F . Les points F_s de ce plan sont des **foyers objet secondaires**. Un faisceau incident issu d'un foyer objet secondaire donne un faisceau de rayons parallèles entre eux, mais non parallèle à l'axe optique. (voir figure 5.16).

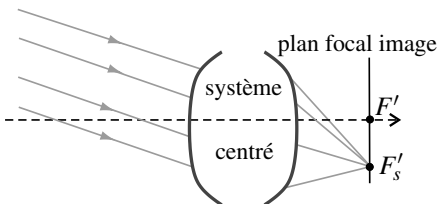


Figure 5.15 – Foyer image secondaire.

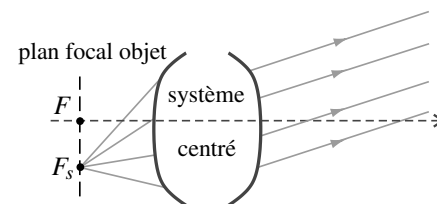


Figure 5.16 – Foyer objet secondaire.

d) **Système afocal**

Un système qui transforme un faisceau incident parallèle en un faisceau émergent parallèle, conjugue un point objet à l’infini avec un point image à l’infini. Il n’a donc pas de foyers et est qualifié de **système afocal** (voir figure 5.17).

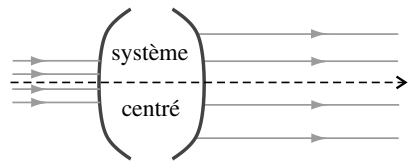


Figure 5.17 – Système afocal.

5 **Lentilles minces**

Les lentilles minces sont les composants élémentaires de la plupart des systèmes centrés.

5.1 **Présentation des lentilles**

a) **Les différentes sortes de lentilles**

Une **lentille** est un composant optique constitué par un milieu transparent, homogène et isotrope, délimité par deux dioptries sphériques ou un dioptre sphérique et un dioptre plan (voir figures 5.18 et 5.19).

Les lentilles sont des systèmes centrés. On distingue deux types de lentilles :

- les **lentilles convergentes** qui sont les lentilles à bords minces (voir figure 5.18),
- les **lentilles divergentes** qui sont les lentilles à bords épais (voir figure 5.19).

Lentilles à bords minces		
biconvexe	plan convexe	ménisque convergent

Figure 5.18 – Différentes lentilles à bords minces, convergentes.

Lentilles à bords épais		
biconcave	plan concave	ménisque divergent

Figure 5.19 – Différentes lentilles à bords épais, divergentes.

Dans les conditions normales d’utilisation de la lentille, il existe un point *O* situé sur l’axe optique, tel que pour tout rayon passant par *O* à l’intérieur de la lentille, le rayon sortant de la lentille est parallèle au rayon entrant (voir figure 5.20). Ce point est appelé **centre optique** de la lentille.

Une **lentille mince** est une lentille dont l’épaisseur *e* sur l’axe est petite comparée aux rayons de courbures de ses faces. Dans ce cas, on néglige complètement cette épaisseur et on représente les lentilles par un trait sur lequel on fait seulement figurer le centre optique *O*. Les

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

symboles des lentilles convergente et divergente sont donnés sur la figure 5.21.

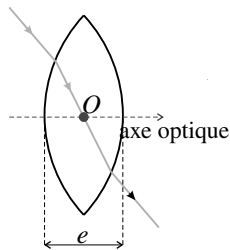


Figure 5.20 – Centre optique d’une lentille.

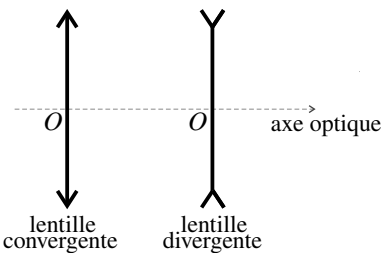


Figure 5.21 – Symbole des lentilles minces.

b) Lentilles convergentes

Si on éclaire une lentille convergente par un faisceau parallèle à l’axe optique, on observe que les rayons réfractés convergent en un point de l’espace image qui est donc le foyer principal image F' (voir figure 5.23). De même, si on place une source ponctuelle au point F , symétrique de F' par rapport à O , on observe qu’après la lentille, le faisceau est parallèle à l’axe, donc F est le foyer principal objet (voir figure 5.22).

Pour une **lentille convergente**, les foyers principaux objet et image sont symétriques par rapport au centre optique. Le foyer objet est situé avant le centre optique O et le foyer image après O dans le sens de propagation de la lumière.

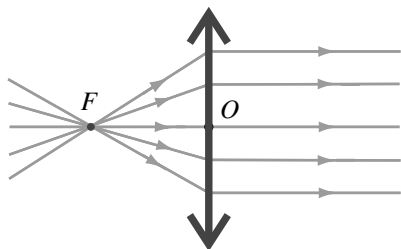


Figure 5.22 – Foyer principal objet d’une lentille convergente.

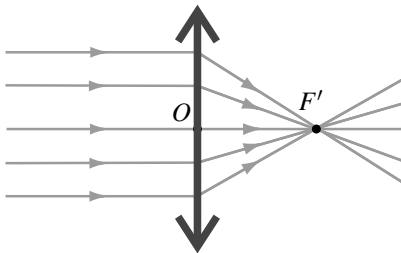


Figure 5.23 – Foyer principal image d’une lentille convergente.

On définit les distances focales :

- la **distance focale objet** $f = \overline{OF}$, qui est négative pour une lentille convergente,
- la **distance focale image** $f' = \overline{OF'}$, qui est positive pour une lentille convergente.

Les distances focales f et f' sont l’opposée l’une de l’autre.

On définit aussi la **vergence** V d'une lentille par :

$$V = \frac{1}{f'} = -\frac{1}{f}. \tag{5.3}$$

La vergence d'une lentille convergente est positive. La vergence est homogène à l'inverse d'une distance et l'unité de la vergence est la *dioptrie* (symbole δ) qui correspond à des m^{-1} . Ainsi une lentille de distance focale image $f' = 50 \text{ cm}$ a une vergence de 2δ .

c) Lentilles divergentes

Si l'on éclaire une lentille divergente avec un faisceau parallèle à l'axe, on observe que le faisceau diverge après passage à travers la lentille (voir figure 5.25). Si on prolonge les rayons du faisceau divergent, ils semblent provenir d'un point de l'espace objet qui est donc l'image du faisceau parallèle. Ce point est le foyer principal image F' de la lentille. C'est un point image virtuel.

Si l'on éclaire la lentille par un faisceau convergent (voir figure 5.24), dont les prolongements de rayon se croisent au point symétrique de F' par rapport à la lentille, le faisceau émergent est parallèle à l'axe. Ce point est donc le foyer principal objet F . C'est un point objet virtuel.

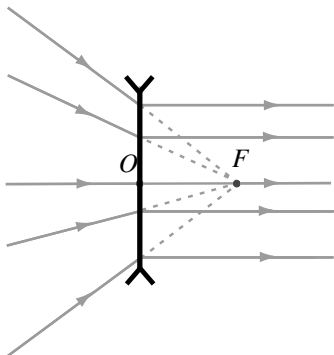


Figure 5.24 – Foyer objet d'une lentille divergente.

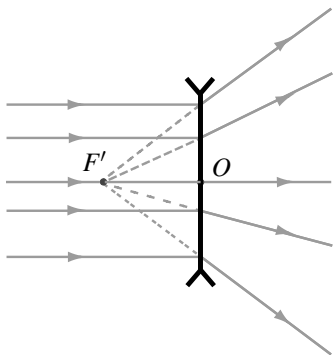


Figure 5.25 – Foyer image d'une lentille divergente.

Pour une **lentille divergente** les foyers principaux objet et image sont symétriques par rapport au centre optique. Ce sont des points *virtuels* : le foyer objet est situé après le centre optique O et le foyer image avant O dans le sens de propagation de la lumière. La distance focale objet $f = \overline{OF}$ est négative, la distance focale image $f' = \overline{OF'}$ est positive et $f = -f'$. La vergence d'une lentille divergente est négative.

5.2 Constructions géométriques

Pour trouver l'image d'un objet par une lentille on peut procéder graphiquement, en réalisant une construction géométrique. Il est important de maîtriser cette méthode qui est présentée dans ce paragraphe.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

a) Les règles de construction

On connaît le rayon émergent d'une lentille mince pour trois types de rayons incidents remarquables :

1. un rayon incident passant par le centre optique O n'est pas dévié ;
2. un rayon incident parallèle à l'axe donne un rayon émergent qui passe (réellement ou virtuellement) par le foyer image F' ;
3. un rayon incident passant (réellement ou virtuellement) par le foyer objet F donne un rayon émergent parallèle à l'axe optique.

De plus, d'après la définition des foyers secondaires :

4. deux rayons incidents parallèles donnent des rayons émergents qui se croisent (réellement ou virtuellement) dans le plan focal image ;
5. deux rayons incidents qui se croisent (réellement ou virtuellement) dans le plan focal objet donnent des rayons émergents parallèles entre eux.

b) Méthode de construction

Les trois premières règles permettent de déterminer graphiquement l'image de n'importe quel point objet B hors de l'axe optique. Il suffit de tracer les trois rayons remarquables passant (réellement ou virtuellement) par B . Les trois rayons émergents passent tous (réellement ou virtuellement) par l'image (réelle ou virtuelle) B' de B .

Dans le cas d'un point objet A situé sur l'axe optique, les trois rayons remarquables sont confondus avec l'axe optique. Pour déterminer l'image A' de A on prend un point B dans le plan de front passant par A et on construit son image B' . A' est l'intersection du plan de front passant par B' avec l'axe optique.

Dans ces constructions les parties virtuelles des rayons sont représentées en pointillés.

c) Application à une lentille convergente

Les figures 5.26 à 5.29 montrent les constructions géométriques dans les quatre cas possibles :

1. l'objet est réel et situé avant F ;
2. l'objet est réel et situé entre F et O ;
3. l'objet est virtuel donc situé après O ;
4. l'objet est situé à l'infini.



Il est vivement conseillé de refaire ces constructions sur une feuille : il faut représenter d'abord la lentille avec son axe, son centre optique et ses deux foyers (dans le bon ordre), puis placer dans la zone choisie un objet $\overrightarrow{AB}$ avec A sur l'axe optique et B dans le plan de front passant par A , puis tracer les trois rayons remarquables passant par B pour en déduire l'image $\overrightarrow{A'B'}$.

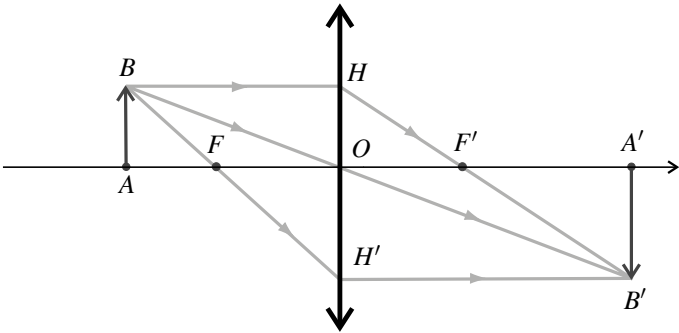


Figure 5.26 – Construction de l'image d'un objet réel situé avant le foyer objet par une lentille convergente. Les trois rayons émergents se croisent en B' , l'image est réelle.

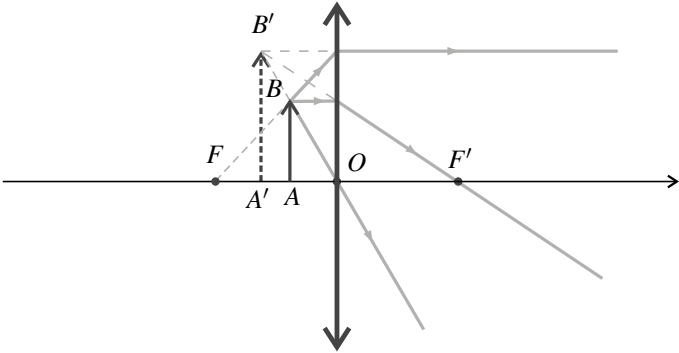


Figure 5.27 – Construction de l'image d'un objet réel situé entre le foyer objet et le centre par une lentille convergente. Les prolongements des trois rayons émergents se croisent en B' , l'image est virtuelle.

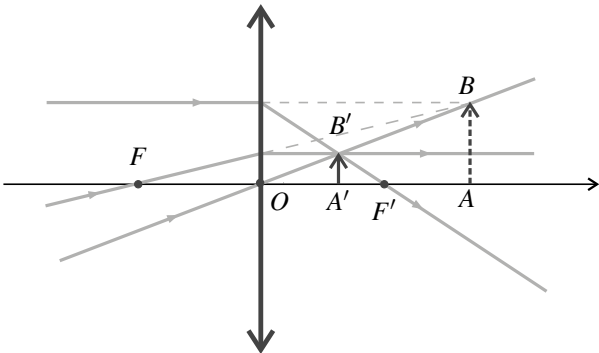


Figure 5.28 – Construction de l'image d'un objet virtuel situé après le centre par une lentille convergente. Les trois rayons émergents se croisent en B' , l'image est réelle.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

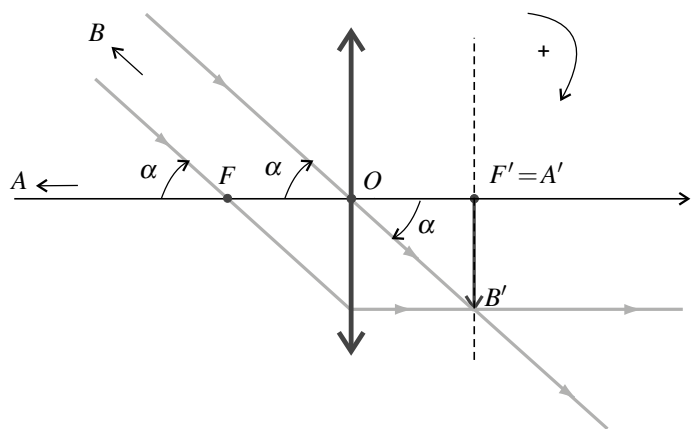


Figure 5.29 – Construction de l’image d’un objet à l’infini. Les deux rayons émergents se croisent en B' , l’image est réelle.

Les résultats sont les suivants :

1. un objet réel avant F donne une image réelle, renversée, située après F' ;
2. un objet réel entre F et O donne une image virtuelle, droite, plus grande, c’est le cas de la loupe ;
3. un objet virtuel après O donne une image réelle, plus petite, droite ;
4. un objet à l’infini a une image réelle, dans le plan focal image.

Dans le cas où l’objet est à l’infini on ne peut tracer que deux rayons remarquables venant de B . Ils sont parallèles entre eux et inclinés d’un angle α par rapport à l’axe optique. B' appartient au plan focal image ce qui est normal puisque B est à l’infini. Dans le triangle $OF'B'$ on peut écrire :

$$\tan \alpha = -\frac{\overline{A'B'}}{\overline{OF'}} \quad \text{soit} \quad \overline{A'B'} = -f' \tan \alpha \simeq -f' \alpha,$$

puisque, dans les conditions de Gauss, l’angle α est très petit.

d) Application à une lentille divergente

Les figures 5.30 à 5.33 montrent les constructions géométriques dans les quatre cas possibles :

1. l’objet est réel donc situé avant O ;
2. l’objet est virtuel et situé entre O et F ;
3. l’objet est virtuel et situé après F ;
4. l’objet est situé à l’infini.

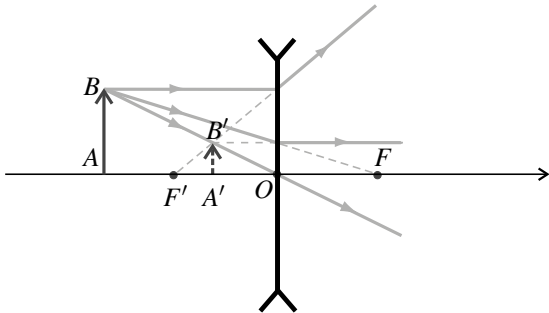


Figure 5.30 – Construction de l'image d'un objet réel situé avant le centre par une lentille divergente. Les prolongements des trois rayons émergents se croisent en B' , l'image est virtuelle.

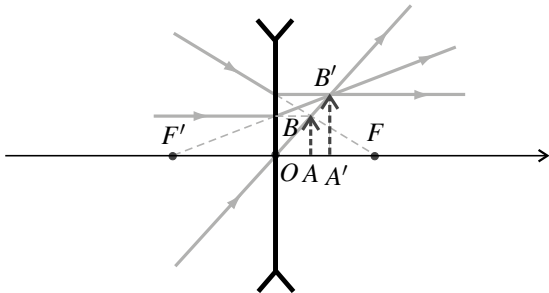


Figure 5.31 – Construction de l'image d'un objet virtuel situé entre le centre et le foyer objet par une lentille divergente. Les trois rayons émergents se croisent en B' , l'image est réelle.

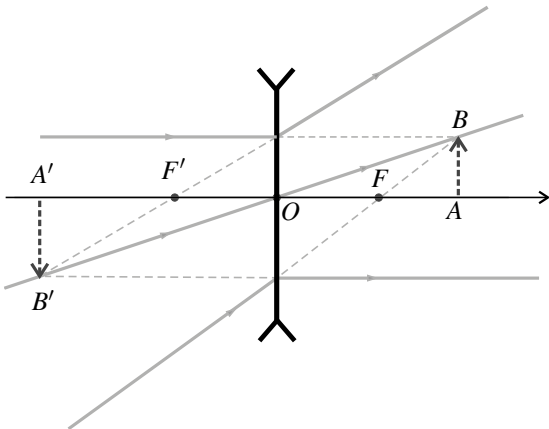


Figure 5.32 – Construction de l'image d'un objet virtuel situé après le foyer objet par une lentille divergente. Les prolongements des trois rayons émergents se croisent en B' , l'image est virtuelle.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

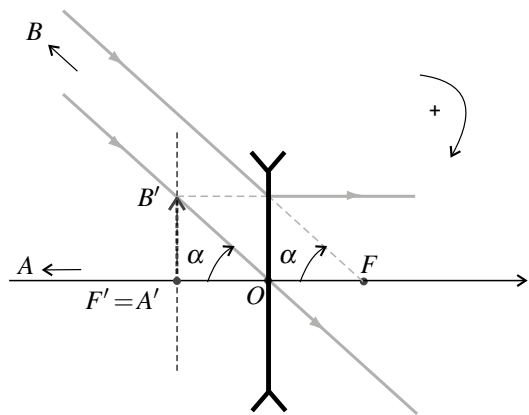


Figure 5.33 – Construction de l'image d'un objet à l'infini. Les prolongements des deux rayons émergents se croisent en B' , l'image est virtuelle.

Les résultats sont les suivants :

1. un objet réel avant O donne une image virtuelle, droite, plus petite ;
2. un objet virtuel entre O et F donne une image réelle, droite, plus grande ;
3. un objet virtuel après F donne une image virtuelle, renversée ;
4. un objet à l'infini a une image virtuelle, située dans le plan focale image.

À la différence de la lentille convergente, une lentille divergente ne peut pas donner une image réelle d'un objet réel.

Dans le cas où l'objet est à l'infini on ne peut tracer que deux rayons remarquables venant de B . Ils sont parallèles entre eux et inclinés d'un angle α par rapport à l'axe optique. B' appartient au plan focal image ce qui est normal puisque B est à l'infini. Dans le triangle $OF'B'$ on peut écrire :

$$\tan \alpha = -\frac{\overline{A'B'}}{\overline{OF'}} \quad \text{soit} \quad \overline{A'B'} = -f' \tan \alpha \simeq -f' \alpha,$$

puisque, dans les conditions de Gauss, l'angle α est très petit.

e) Comment reconnaître rapidement une lentille ?

En pratique on peut savoir facilement si une lentille est convergente ou divergente :

- Lorsqu'on regarde à travers une lentille divergente on voit quelle que soit la position de l'objet (qui est réel bien sûr) une image à l'endroit de taille réduite : c'est le cas 1 du paragraphe d) ;
- lorsqu'on regarde à travers une lentille convergente un objet lointain l'image est le plus souvent floue, il faut éloigner la lentille de son œil pour voir une image nette de taille réduite et à l'envers : c'est le cas 1 étudié au paragraphe c). Si la distance focale de la lentille est trop grande on ne trouve pas d'image nette.

Dans le cas d'une lentille convergente on peut déterminer rapidement une valeur approchée de sa distance focale. Il faut pour cela regarder à travers la lentille un objet placé assez près de la lentille : on voit une image droite et agrandie, c'est le cas 2 du paragraphe c) page 162 qui correspond à l'utilisation de la lentille comme loupe. Lorsqu'on éloigne la lentille de l'objet l'image s'agrandit de plus en plus puis devient floue. À ce moment on est à la limite du cas 2 et du cas 1 : la distance entre l'objet et la lentille est égale à la distance focale f' .

5.3 Relations de conjugaison

a) Principe

Les images trouvées par les constructions géométriques peuvent se calculer en utilisant les **formules de conjugaison**. Ces formules vont par deux :

- une **formule de position** liant les positions de deux points A et A' conjugués appartenant à l'axe optique ;
- une **formule du grandissement** exprimant le grandissement $\gamma = \frac{\overline{A'B'}}{\overline{AB}}$ entre les plans de front conjugués passant par A et A' qui permet de positionner l'image B' d'un point B quelconque du plan de front passant par A .

Les formules de conjugaison ne sont pas à apprendre par cœur. Elles doivent être fournies aux candidats lors des épreuves de concours.

b) Relations avec origine aux foyers, dites relations de Newton

Ces relations font apparaître les mesures algébriques $\overline{FA}$ et $\overline{F'A'}$ qui repèrent les positions du point objet A et du point image A' respectivement par rapport au foyer objet F et au foyer image F' .

La formule de position est :

$$\overline{F'A'} \times \overline{FA} = -f'^2 = -f^2 = ff', \tag{5.4}$$

et la formule du grandissement :

$$\gamma = -\frac{\overline{F'A'}}{f'} = -\frac{f}{\overline{FA}}. \tag{5.5}$$



Ces formules font intervenir des mesures algébriques. Lorsqu'on les applique il faut contrôler les signes en observant une figure.

c) Relations avec origine au centre, dites relations de Descartes

Ces relations font apparaître les mesures algébriques $\overline{OA}$ et $\overline{OA'}$ qui repèrent les positions du point objet A et du point image A' par rapport au centre optique O de la lentille.

La formule de position est :

$$\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{f'} = V, \tag{5.6}$$

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

et la formule du grandissement :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OA'}}{\overline{OA}}. \quad (5.7)$$

5.4 Complément : démonstration des formules de conjugaison

On va établir les formules de conjugaison à partir de la figure 5.26.

On peut exprimer le grandissement en faisant intervenir le point H projeté de B sur la lentille et le point H' projeté de B' sur la lentille (voir figure 5.26). Il vient alors, en utilisant le théorème de Thalès dans les triangles ABF et $OH'F$:

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OH'}}{\overline{AB}} = \frac{\overline{FO}}{\overline{FA}} = -\frac{f}{\overline{FA}}.$$

De même, en utilisant le théorème de Thalès dans les triangles HOF' et $F'A'B'$:

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{A'B'}}{\overline{OH}} = \frac{\overline{F'A'}}{\overline{F'O}} = -\frac{\overline{F'A'}}{f'}.$$

On vient de trouver les relations de grandissement de Newton. En combinant ces deux relations on obtient la relation de position.

$$\overline{F'A'} \times \overline{FA} = f f'.$$

Revenant à l'observation de la figure 5.26, on peut écrire le grandissement d'une autre manière, en utilisant le théorème de Thalès dans les triangles OAB et $OA'B'$:

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OA'}}{\overline{OA}},$$

qui est la formule du grandissement de Descartes. Enfin, en modifiant la formule de position de Newton grâce à la loi de Chasles pour faire apparaître le centre O on trouve :

$$(\overline{F'O} + \overline{OA'}) (\overline{FO} + \overline{OA}) = -f'^2.$$

En remplaçant $\overline{F'O}$ par $-f'$ et $\overline{FO}$ par $-f = f'$ et en développant, les termes en f'^2 se simplifient et on obtient :

$$f' \overline{OA'} - f' \overline{OA} + \overline{OA} \overline{OA'} = 0.$$

Il ne reste plus qu'à diviser l'expression par $f' \overline{OA} \overline{OA'}$, pour trouver la relation de position de Descartes.

6 Applications des lentilles

6.1 Projection d'une image

a) Position du problème

On veut projeter l'image d'un objet rétroéclairé sur un écran (diapositive par exemple). On souhaite avoir une image agrandie le plus possible, aussi lumineuse et nette que possible, avec une distance D entre l'objet et l'écran qui est imposée par les conditions extérieures.

Pour cela on doit utiliser une lentille nécessairement convergente car il faut obtenir une image réelle d'un objet réel. Comment choisir la distance focale de cette lentille ? Comment obtenir une image lumineuse et uniformément éclairée ?

b) Choix de la lentille

Condition pour avoir une image nette On voit une image nette si la lentille forme l'image de l'objet exactement sur l'écran. Pour cela il faut placer la lentille convergente au bon endroit entre l'objet et l'écran. La figure 5.34 illustre cette situation : A est un point de l'objet conjugué avec le point A' sur l'écran.

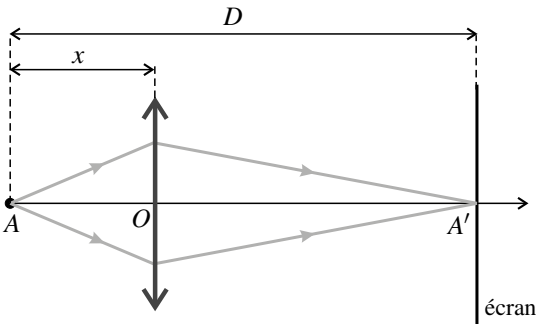


Figure 5.34 – Principe de la projection sur un écran.

Pour calculer la bonne distance x entre l'objet et la lentille on peut appliquer la formule de conjugaison de Descartes. On trouve en observant la figure (en faisant attention aux signes des mesures algébriques) : $\overline{OA} = -x$ et $\overline{OA'} = D - x$. Il vient :

$$\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{f'} \Leftrightarrow \frac{1}{D-x} + \frac{1}{x} = \frac{1}{f'}.$$

Ceci conduit à l'équation du second degré en x : $x^2 - xD + Df' = 0$. Cette équation n'a de solution que si son discriminant Δ est positif, or : $\Delta = D^2 - 4f'D = D(D - 4f')$. On en déduit une condition sur la distance focale de la lentille :

Pour projeter un objet sur un écran avec une lentille convergente, il faut que la distance D entre l'objet et l'écran et la distance focale image f' de la lentille vérifient :

$$D > 4f'.$$

Condition pour avoir un grandissement suffisant D'après la formule du grandissement :

$$\gamma = \frac{\overline{OA'}}{\overline{OA}} = \frac{D-x}{-x} = -\left(\frac{D}{x} - 1\right).$$

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

Ce grandissement est négatif : l'image sur l'écran est renversée (ceci apparaît sur la figure 5.26). La taille de l'image est proportionnelle à $|\gamma|$ qui augmente si x diminue : *pour avoir une grande image il faut que la lentille soit le plus près possible de l'objet.*

L'équation du second degré précédente a pour solutions :

$$x_1 = \frac{1}{2} \left(D - \sqrt{D^2 - 4Df'} \right) \quad \text{et} \quad x_2 = \frac{1}{2} \left(D + \sqrt{D^2 - 4Df'} \right).$$

Ces deux racines correspondent à des positions de la lentille de part et d'autre du milieu entre l'objet et l'écran. Pour une image agrandie il faut choisir la plus petite des deux soit x_1 . En effet, dans l'autre position l'écran est plus près de la lentille que l'objet et donc l'image est réduite, d'après la formule $|\gamma| = \frac{OA'}{OA}$.

D'autre part, x_1 diminue si f diminue.

Pour avoir une image agrandie il faut placer la lentille plus près de l'objet que de l'écran.
Pour une distance objet-écran fixée, l'image est d'autant plus grande que la distance focale de la lentille convergente utilisée est petite.

Prise en compte des conditions de Gauss Pour que l'image soit de bonne qualité il faut respecter les conditions de Gauss : les rayons doivent être proches de l'axe optique de la lentille et peu inclinés par rapport à cet axe. Or, en diminuant la distance objet-lentille, pour une taille de lentille donnée, on augmente l'angle d'inclinaison des rayons qui frappent la lentille sur ses bords. Il y a donc un compromis à trouver entre la taille de l'image et sa netteté.

La lentille doit avoir une distance focale suffisamment petite pour avoir un grandissement suffisant mais pas trop petite pour respecter les conditions de Gauss.

Par ailleurs on améliore la qualité des conditions de Gauss en utilisant, si nécessaire, un diaphragme circulaire qui arrête les rayons trop éloignés de l'axe optique. Ce dispositif est visible sur la figure 5.35 ci-dessous.

c) Éclairage de l'objet

Pour éclairer l'objet de manière optimale on utilise une lanterne munie d'un **condenseur**. Un condenseur est une lentille convergente (ou un ensemble de lentilles) permettant de concentrer les rayons de la source de lumière vers l'objet pour qu'il soit plus lumineux (voir figure 5.35). Le meilleur réglage du condenseur est celui pour lequel la lumière converge sur la lentille de projection. On a ainsi un éclairage de l'objet uniforme et une grande quantité de lumière qui traverse la lentille.

Le montage de projection complet est montré avec son réglage sur la figure 5.35.

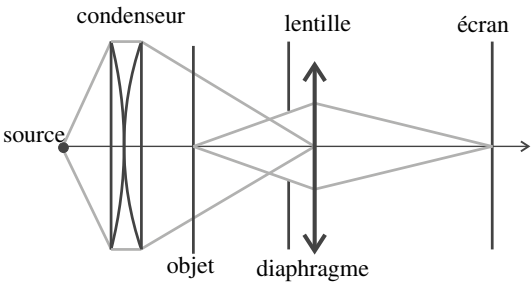


Figure 5.35 – Montage pratique pour la projection d’une image.

6.2 Le microscope

Le microscope est constitué de deux parties optiques :

- l’**objectif** qui est du côté de l’objet ;
- l’**oculaire** qui est du côté de l’œil.

Pour comprendre le fonctionnement il faut savoir qu’un œil normal voit correctement sans se fatiguer lorsque l’objet observé est situé à l’infini (voir le dernier paragraphe de ce chapitre). Avec un instrument d’optique, l’objet pour l’œil est en fait l’image finale donnée par l’instrument. Il faut donc retenir que :

L’image finale donnée par un instrument d’optique oculaire doit être **à l’infini**.

a) Modélisation d’un microscope

On modélise très simplement un microscope par deux lentilles minces :

- l’objectif est équivalent à une lentille convergente L_1 de distance focale $f'_1 = 5\text{ mm}$,
- l’oculaire est équivalent à une lentille convergente L_2 de distance focale $f'_2 = 2,5\text{ cm}$.

Puisqu’on a deux lentilles il faut préciser leur position relative. On se donne donc la distance entre le foyer image de l’objectif F'_1 et le foyer objet de l’oculaire F_2 , soit $\Delta = F'_1F_2 = 25\text{ cm}$. Ces valeurs sont indicatives, elles varient d’un microscope à l’autre.

b) Étude de la mise au point

La mise au point consiste à déplacer l’objet jusqu’à voir son image nette dans l’oculaire sans fatigue. Quelle doit être la distance entre l’objectif et l’oculaire ?

Soit AB l’objet observé. L’objectif donne de AB une image A_1B_1 ; l’oculaire donne de A_1B_1 une image à l’infini. Ceci est résumé par le schéma suivant :

$$AB \xrightarrow{\text{objectif}} A_1B_1 \xrightarrow{\text{oculaire}} \infty.$$

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

 Ce type de schéma « des images successives » est très utile lorsqu'on étudie un instrument avec 2 ou plusieurs lentilles.

Puisque l'image finale est à l'infini, A_1B_1 doit se trouver dans le plan focal objet de l'oculaire, donc : $A_1 = F_2$.

On peut en déduire la position de AB en utilisant une formule de conjugaison. Il est judicieux d'utiliser la relation de Newton pour L_1 qui s'écrit :

$$\overline{F_1'F_2} \times \overline{F_1A} = -f_1'^2, \quad \text{ce qui donne : } \overline{F_1A} = -\frac{f_1'^2}{\Delta} = -0,1 \text{ mm.}$$

On obtient enfin $\overline{O_1A}$ par la loi de Chasles :

$$\overline{O_1A} = \overline{O_1F_1} + \overline{F_1A} = -f_1' + \overline{F_1A} = -5,1 \text{ mm.}$$

C'est la distance entre la lentille objectif et l'objet lorsque la mise au point est réussie.

c) Construction géométrique

Pour effectuer la construction, il faut chercher l'antécédent de F_2 par L_1 , ce qui a été réalisé sur la figure 5.36, en utilisant les trois rayons remarquables pour L_1 passant par B_1 .

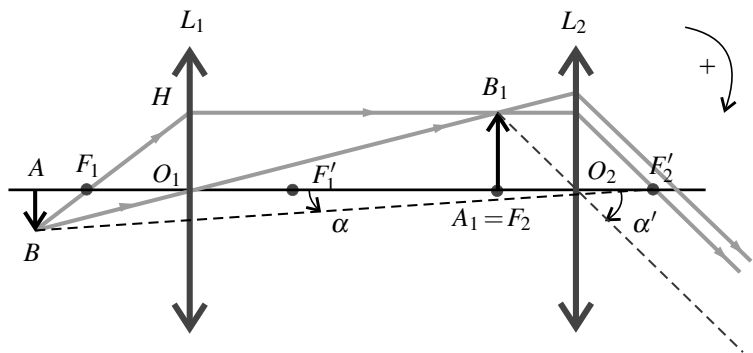


Figure 5.36 – Microscope.

d) Grossissement

Par définition :

Le grossissement G est le rapport de l'angle α' sous lequel l'objet est vu à travers le microscope sur l'angle α sous lequel il est vu à l'œil nu.

Le microscope est conçu de manière à ce que l'on place son œil au foyer image de l'oculaire. L'angle sous lequel l'objet est vu dans le microscope est l'angle d'inclinaison α' des rayons émergents (voir figure 5.36). Comme on a pris le sens horaire comme sens positif pour les

angles, α' est positif. De plus cet angle est petit puisque l'on doit respecter les conditions de Gauss. On peut alors écrire :

$$\alpha' \simeq \tan \alpha' = \frac{\overline{A_1 B_1}}{\overline{F_2 O_2}} = \frac{\overline{A_1 B_1}}{f'_2}.$$

On peut exprimer $\overline{A_1 B_1}$ en fonction de $\overline{AB}$ en utilisant la relation de grandissement de Newton pour la lentille L_1 : $\gamma_1 = \frac{f'_1}{\overline{F_1 A}}$. On obtient finalement :

$$\alpha' = \frac{\overline{AB}}{\overline{F_1 A}} \frac{f'_1}{f'_2}.$$

Il faut maintenant déterminer l'angle sous lequel AB est vu sans microscope. L'angle est montré sur la figure 5.36 (rappel : l'œil est placé en F'_2). α est négatif, très petit d'après les conditions de Gauss, et :

$$\alpha \simeq \tan \alpha = \frac{\overline{AB}}{\overline{AF'_2}}.$$

Or : $\overline{AF'_2} = \overline{AF_1} + \overline{F_1 F'_1} + \overline{F'_1 F_2} + \overline{F_2 F'_2} = \overline{AF_1} + 2f'_1 + \Delta + 2f'_2 = 310,1 \text{ mm}$. Le grossissement est donc :

$$G = \frac{\alpha'}{\alpha} = \frac{f'_1}{f'_2} \frac{\overline{AF'_2}}{\overline{F_1 A}} = -620. \tag{5.8}$$

On trouve un grandissement négatif car l'image sera vue renversée par rapport à l'objet ce qui est dû à l'objectif. On peut s'en rendre compte en déplaçant l'objet car on voit l'image se déplacer en sens inverse.

6.3 La lunette de Galilée

La lunette de Galilée est la lunette la plus simple qui permette d'observer des objets terrestres. On modélise l'objectif par une lentille convergente L_1 de distance focale $f'_1 = 60 \text{ cm}$ et l'oculaire par une lentille divergente L_2 de distance focale $f'_2 = -5 \text{ cm}$. Les valeurs ci-dessus sont des exemples, elles peuvent différer d'une lunette à l'autre.

Les objets observés sont à une distance grande devant la distance focale, on peut alors considérer qu'ils sont à l'infini : on dit que la lunette est réglée à l'infini. Comme pour le microscope, l'image finale doit être à l'infini. L'objet et l'image étant à l'infini, il s'agit d'un **système afocal**.

Puisque l'objet est à l'infini, l'image intermédiaire est dans le plan focal image de L_1 et comme l'image finale est à l'infini, l'image intermédiaire est dans le plan focal objet de L_2 :

$$\infty \xrightarrow{\text{objectif}} A_1 B_1 \xrightarrow{\text{oculaire}} \infty \text{ avec } A_1 = F_2 = F'_1.$$

Il faut donc disposer les lentilles de manière à ce que $F'_1 = F_2$ (voir figure 5.37).

Sur la figure, on a représenté un faisceau parallèle provenant de l'objet sous un angle α et émergeant de la lunette sous un angle α' . Comme cela a été fait, on a toujours intérêt à représenter l'image intermédiaire car cela peut faciliter les calculs ultérieurs.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

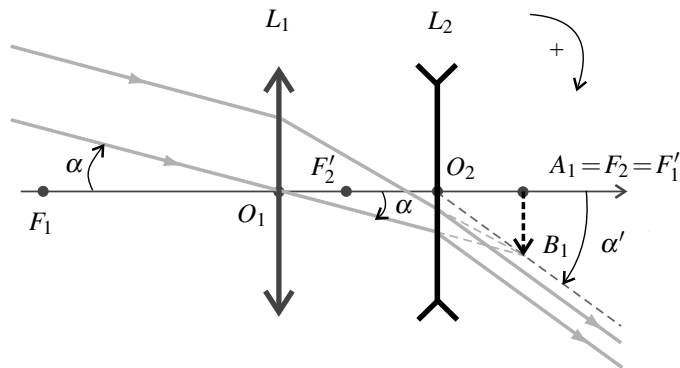


Figure 5.37 – Lunette de Galilée.

On cherche à calculer le grossissement :

$$G = \frac{\alpha'}{\alpha}.$$

On utilise pour cela l'image intermédiaire et les rayons passant par les centres optiques :

$$\alpha' \simeq \tan \alpha' = \frac{\overline{B_1 A_1}}{\overline{O_2 F_2}} \quad \text{et} \quad \alpha \simeq \tan \alpha = \frac{\overline{B_1 A_1}}{\overline{O_1 F_1'}},$$

avec le sens positif choisi pour les angles. Ainsi :

$$G = \frac{\alpha'}{\alpha} = \frac{\overline{O_1 F_1'}}{\overline{O_2 F_2}} = \frac{f_1'}{f_2} = 12.$$

Le grossissement est positif donc objet et image sont dans le même sens.

6.4 La lunette astronomique

La lunette astronomique est aussi une lunette permettant d'observer des objets à très grandes distances (à l'infini). On modélise l'objectif par une lentille convergente L_1 de distance focale $f_1' = 60$ cm et l'oculaire par une lentille convergente L_2 de distance focale $f_2' = 5$ cm. Les valeurs ci-dessus sont des exemples, elles peuvent différer d'une lunette à l'autre.

Comme la lunette de Galilée, la lunette astronomique est réglée à l'infini et $F_1' = F_2$ (voir figure 5.38) et c'est un système afocal. La méthode de calcul du grandissement est la même :

$$\alpha' \simeq \tan \alpha' = -\frac{\overline{A_1 B_1}}{\overline{F_2 O_2}} \quad \text{et} \quad \tan \alpha \simeq \alpha = \frac{\overline{A_1 B_1}}{\overline{O_1 F_1'}},$$

avec le sens positif choisi pour les angles. Ainsi :

$$G = \frac{\alpha'}{\alpha} = -\frac{\overline{O_1 F_1'}}{\overline{F_2 O_2}} = -\frac{f_1'}{f_2} = -12.$$

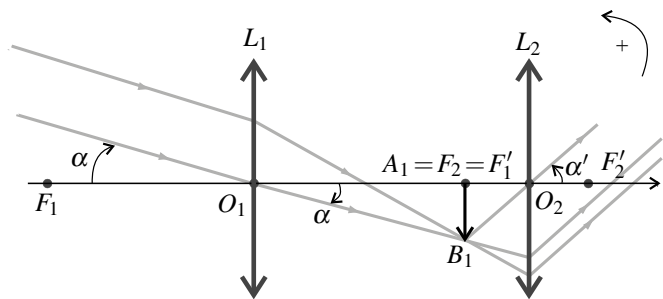


Figure 5.38 – Lunette astronomique.

Le grossissement est négatif, donc l’image est à l’envers. C’est pour cela qu’on ne se sert pas de cette lunette pour observer des objets terrestres.

7 L’œil

7.1 Description et modélisation

Une coupe de l’œil est présentée sur la figure 5.39.

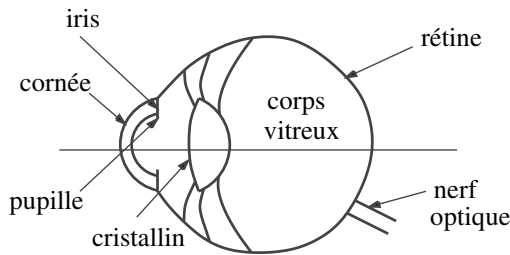


Figure 5.39 – Structure de l’œil.

- L’iris (partie colorée) est percé de la pupille dont le diamètre est variable (de 2 à 8 mm) et qui joue le rôle de diaphragme c’est-à-dire qui limite la puissance lumineuse pénétrant dans l’œil.
- Le cristallin est un muscle assimilable à une lentille mince biconvexe dont la distance focale est variable selon sa contraction. Il donne d’un objet une image renversée sur la rétine.
- La rétine est constituée de cellules sensibles à la lumière (les cônes et les bâtonnets).

On modélise l’œil par une lentille mince convergente de vergence variable - correspondant au cristallin - formant une image sur un écran fixe - correspondant à la rétine.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

7.2 Caractéristiques optiques

a) Limite de résolution angulaire

L'œil ne distingue deux détails différents de l'objet que si leur image se forme sur deux cellules différentes de la rétine.

Dans de bonnes conditions d'éclairement (ni trop sombre, ni trop lumineux), l'œil distingue des détails d'environ 1 minute d'arc, soit 3.10^{-4} rad.

Cette valeur constitue la **limite de résolution** ou **pouvoir séparateur** de l'œil. Elle n'est atteinte que dans des conditions d'éclairement et de contraste optimales.

b) Plage d'accommodation

L'œil ne peut voir une image nette que si elle se forme sur la rétine. Un œil au repos (cristallin non contracté) voit à une distance maximale. Le point situé à cette distance porte le nom de **Punctum Remotum** noté P.R.. Pour voir des objets plus proches, le cristallin doit se contracter pour être plus convergent (sa vergence augmente), on dit que l'œil **accommode**. Le point le plus proche que peut voir net un œil est appelé **Punctum Proximum** noté P.P.. Il correspond à la distance minimale de mise au point. La zone située entre le P.P. et le P.R. est appelé **champ en profondeur de l'œil**.

Pour un œil normal, le P.R. est à l'infini et le P.P. à environ 25 cm pour un adulte.

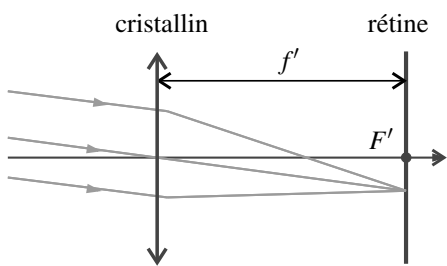


Figure 5.40 – Formation de l'image d'un objet au P.R.

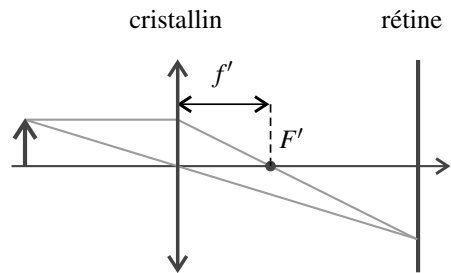


Figure 5.41 – Formation de l'image d'un objet au P.P.

Les figures 5.40 et 5.41 représentent la formation de l'image d'un objet au P.R. et au P.P.. La distance focale du cristallin varie mais la distance entre le cristallin et la rétine ne varie pas.

7.3 Défauts de l'œil

a) La myopie

Un œil **myope** possède un cristallin trop convergent. Par conséquent, le P.P. est plus proche que pour l'œil normal et le P.R. n'est plus à l'infini. On s'aperçoit sur la figure que l'image d'un objet ponctuel à l'infini se forme avant la rétine et l'observateur ne voit qu'une tache floue (le myope sans ses lunettes voit flou de loin). Il faut rendre l'œil moins convergent donc la correction se fait par l'utilisation d'une lentille divergente.

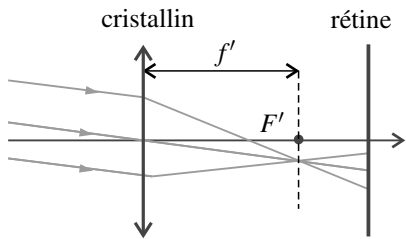


Figure 5.42 – Oeil myope sans accommodation.

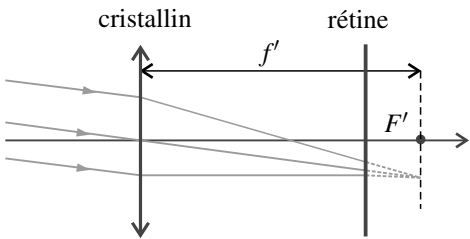


Figure 5.43 – Oeil hypermétrope sans accommodation.

b) L'hypermétropie

Un œil **hypermétrope** possède un cristallin insuffisamment convergent. L'œil doit donc accommoder pour voir à l'infini sinon l'image se forme derrière la rétine (voir figure 5.43) et l'observateur observe une tache floue. Le P.P. est plus éloigné que pour l'œil normal. Il faut rendre l'œil plus convergent donc la correction se fait par l'utilisation d'une lentille convergente.

8 Approche documentaire : influence des réglages sur l'image produite par un appareil photographique numérique

Dans ce paragraphe on va interpréter à l'aide des lois de l'optique l'influence sur l'image de trois réglages d'un appareil photographique :

- la durée d'exposition,
- l'ouverture du diaphragme,
- la distance focale.

8.1 Modélisation

Pour simplifier, on modélise l'objectif de l'appareil par une lentille convergente unique de distance focale f' , appelée dans la suite simplement **focale** (voir figure 5.44).

La lumière entrant dans l'objectif doit traverser le **diaphragme**, ouverture de forme proche d'un cercle de diamètre D réglable, que l'on modélisera par une ouverture circulaire placée juste devant la lentille mince précédente.

L'image est enregistrée à l'aide d'un capteur CCD situé dans un plan de front, à distance d de la lentille.

La lumière arrive sur le capteur pendant une durée contrôlée τ appelée **durée d'exposition**.

Pour obtenir une image correcte il faut que :

- l'image du sujet photographié se forme sur le capteur (à un point du sujet correspond un point du capteur),
- la quantité de lumière reçue par les éléments sensibles du capteur soit convenable.

Le première condition est réalisée par le réglage de mise au point qui joue sur la distance d entre l'objectif et le capteur. La valeur qu'il faut donner à d dépend de la distance du sujet à l'objectif, selon les formules de conjugaison de la lentille.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

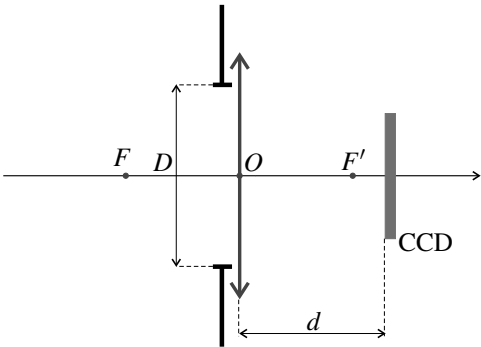


Figure 5.44 – Modélisation de l'appareil photographique.

La deuxième condition est réalisée par le réglage simultané de l'ouverture du diaphragme D et du temps d'exposition τ .

Dans la suite on va comparer des photos prises avec différents réglages pour voir leur influence sur l'image obtenue.

Les trois photos 5.45, 5.46 et 5.47 ont été prises avec une même focale $f' = 35$ mm, un diaphragme de diamètre $D = \frac{f'}{8}$ et des durées d'exposition τ différentes. La photo 5.45 est *sous-exposée* et la photo 5.47 est *surexposée*.

Pourquoi en est-il ainsi ? Le signal fourni par une cellule sensible du capteur dépend de l'énergie qu'elle reçoit, énergie qui est proportionnelle à la durée d'exposition τ (les autres paramètres étant fixés). L'échelle des gris, du noir au blanc, correspond à une énergie reçue variant entre $E_{\min}$ et $E_{\max}$; pour $E < E_{\min}$ (respectivement $E > E_{\max}$) le pixel sur l'image est noir (respectivement blanc).

Si trop de capteurs reçoivent une énergie inférieure à $E_{\min}$, l'image est trop noire (sous-exposée) et si trop de capteurs reçoivent une énergie supérieure à $E_{\max}$ l'image est trop blanche (surexposée).

8.2 Influence du diaphragme d'ouverture

Il est d'usage de donner l'ouverture par le *nombre d'ouverture* $N = \frac{f'}{D}$.

a) Exposition

Les photos 5.48, 5.46 et 5.49 ont été prises avec la même focale et des nombres d'ouvertures respectifs : $N = 2, 8$ et 22 . Les trois photos sont correctement exposées pour des durées d'exposition respectives : $\tau = \frac{1}{500}$ s, $\frac{1}{30}$ s et $\frac{1}{30}$ s. Quelle est la relation entre D et τ ?

L'onde lumineuse se caractérise par son éclairement $\mathcal{E}$ qui est une puissance surfacique.

Ainsi, la puissance entrant dans l'objectif est $\mathcal{E} \pi \frac{d^2}{4}$, qui augmente avec D . L'énergie reçue par un pixel est proportionnelle à la puissance qu'il reçoit (qui est proportionnelle à la



Figure 5.45 –
 $\tau = \frac{1}{4} \text{s}$



Figure 5.46 –
 $\tau = \frac{1}{30} \text{s}$

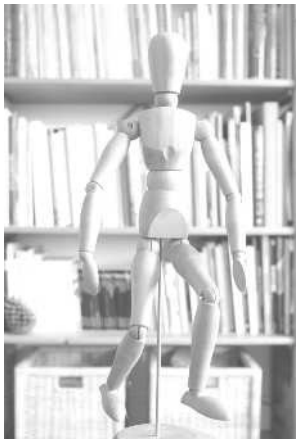


Figure 5.47 –
 $\tau = \frac{1}{125} \text{s}$



Figure 5.48 – $f' = 35 \text{ mm}$,
 $D = \frac{f}{2}, \tau = \frac{1}{500} \text{s}$.



Figure 5.49 – $f' = 35 \text{ mm}$,
 $D = \frac{f}{22}, \tau = \frac{1}{4} \text{s}$.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

puissance entrant dans l'objectif) multipliée par la durée d'exposition τ . Elle est donc proportionnelle au produit $D^2 \tau$ qui détermine la qualité de l'exposition.

Pour les photos 5.48, 5.46 et 5.49 le produit $\left(\frac{D}{f'}\right)^2 \tau$ vaut, en secondes :

$$\left(\frac{1}{2}\right)^2 \times \frac{1}{500} = 5,0 \cdot 10^{-4} \quad \left(\frac{1}{8}\right)^2 \times \frac{1}{30} = 5,2 \cdot 10^{-4}, \quad \text{et} \quad \left(\frac{1}{22}\right)^2 \times \frac{1}{4} = 5,2 \cdot 10^{-4}.$$

Il est quasiment constant : les trois photos sont correctement exposées.

b) Profondeur de champ

On remarque de plus que l'arrière plan est flou sur la photo 5.48, à peu près net sur la photo 5.46 et bien net sur la photo 5.49.

Comment des points en dehors du plan de mise au point peuvent-ils être nets ? Quel est le rôle du diamètre d'ouverture D ?

La figure 5.50 représente un point A sur lequel la mise au point est faite : son image A' se forme sur le capteur et un point B situé derrière A . L'image B' de B est située devant le capteur CCD. En effet d'après la formule de conjugaison de Newton :

$$\overline{F'B'} \overline{FB} = \overline{F'A'} \overline{FA} \quad \text{d'où} \quad F'B' = F'A' \frac{FA}{FB} < F'A'.$$

De ce fait, la lumière issue de B fait sur le capteur une tache qui a la forme du diaphragme d'ouverture. Pour un diaphragme circulaire de diamètre D , cette tache est circulaire de diamètre δ et d'après le théorème de Thalès :

$$\delta = D \frac{B'A'}{OB'}.$$

D'après cette formule : *plus le diamètre D du diaphragme est petit, plus le diamètre δ de la tache est petit.*

D'autre part, si B s'éloigne de A et de l'appareil, FB augmente, donc, d'après la relation de Newton, $F'B'$ diminue. Ainsi OB' diminue et $B'A'$ augmente, ce qui fait que δ augmente. *Plus le point B est éloigné derrière A , plus le diamètre de la tache est grand.*

Quelle est la condition pour que B soit net sur la photographie ? Si δ est inférieur à la taille d'une cellule élémentaire du capteur (soit un pixel de l'image finale), B sera aussi net que A sur l'image enregistrée. Mais, cette condition est inutilement contraignante parce qu'on ne distingue pas les pixels à l'unité sur une image numérique. Pour que B soit net, il suffit que δ soit plus petit qu'une distance $\delta_{\max}$ qui dépend des conditions de visualisation de la photographie et qui est en gros de l'ordre de 3 fois la taille d'un pixel.

Ainsi, pour les points de l'arrière plan, sur la photo 5.49 $\delta < \delta_{\max}$ alors que $\delta > \delta_{\max}$ sur la photo 5.48, ce qui provient du fait que D est 11 fois plus grand. Sur la photo 5.49 tout objet situé entre le mannequin et la bibliothèque serait net. Sur la photo 5.48 il y a quelque part entre le mannequin et la bibliothèque une limite à partir de laquelle les points ne sont plus nets.

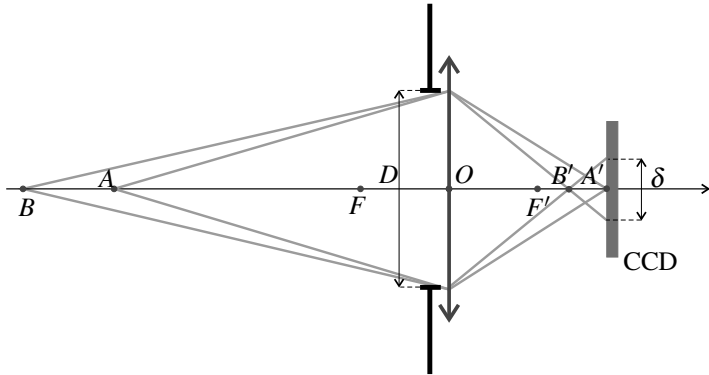


Figure 5.50 – Influence de l'ouverture du diaphragme sur la profondeur de champ.

De manière analogue, un point situé plus proche de l'objectif que le plan de mise au point serait net à condition de ne pas être trop éloigné de celui-ci.

On appelle *profondeur de champ* la distance entre le point net le plus proche de l'objectif et le point net le plus éloigné de l'objectif. *La profondeur de champ diminue si le diamètre D du diaphragme d'ouverture augmente.*

c) Diffraction

Le phénomène de diffraction limite-t-il la netteté de l'image pour les petits diamètres d'ouverture ?

La dispersion angulaire due à la diffraction est de l'ordre de $\frac{\lambda}{D}$ où λ est la longueur d'onde. Dans le cas d'une mise au point à l'infini, soit $d = f'$, la tache de diffraction sur le capteur a donc une dimension de l'ordre de $\delta' = \frac{f'\lambda}{D} = N\lambda$. Pour $N = 22$ et $\lambda = 0,5 \mu\text{m}$, $\delta' \simeq 10 \mu\text{m}$.

Il faut comparer δ' à la taille d'un pixel. Le capteur est un rectangle de $16 \text{ mm} \times 24 \text{ mm}$. Il comporte $15 \cdot 10^6$ pixels, carrés de côté $\sqrt{\frac{16 \cdot 10^{-3} \times 24 \cdot 10^{-3}}{15 \cdot 10^6}} \simeq 5 \cdot 10^{-6} \text{ m} = 5 \mu\text{m}$.

Ce calcul approché montre que, dans le cas du diamètre d'ouverture le plus petit, la tache de diffraction couvre quelques pixels, ce qui peut réduire le « piqué » de l'image.

8.3 Influence de la distance focale

Les photos 5.51 et 5.52 ont été prises toutes les deux du même endroit en faisant la mise au point sur l'infini, mais en utilisant deux objectifs différents, respectivement un objectif grand-angle de distance focale $f' = 24 \text{ mm}$ et un téléobjectif de distance focale $f' = 70 \text{ mm}$.

Le nombre d'ouverture $N = \frac{f}{D}$ est le même dans les deux cas.

On constate que :

- le champ angulaire (champ de vision de la photo) est nettement plus large avec l'objectif de courte focale,

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE



Figure 5.51 – $f' = 24 \text{ mm}$,
 $D = \frac{f}{4}$.



Figure 5.52 – $f' = 70 \text{ mm}$,
 $D = \frac{f}{4}$.

- le premier plan est net sur la photo 5.51 alors qu'il est flou sur l'autre : la profondeur de champ est plus grande avec l'objectif de courte focale.

a) Champ angulaire

On souhaite expliquer l'influence de la distance focale sur le champ angulaire.

La mise au point est faite sur l'infini : un point à l'infini a son image sur le capteur CCD qui est donc placé dans le plan focal image de l'objectif, à distance $d = f'$ de celui-ci.

Le capteur a la forme d'un rectangle de largeur $\ell = 16 \text{ mm}$ et longueur $L = 24 \text{ mm}$. On voit sur la figure 5.53 que l'angle d'ouverture α du champ de vision dans un plan horizontal est donné par :

$$\tan \frac{\alpha}{2} = \frac{\ell}{2f'}.$$

α diminue avec la focale f' . Numériquement : $\alpha = 2 \arctan \left(\frac{16}{2 \times 24} \right) = 37^\circ$ sur la photo 5.51,

et $\alpha = 2 \arctan \left(\frac{16}{2 \times 70} \right) = 13^\circ$ sur la photo 5.52.

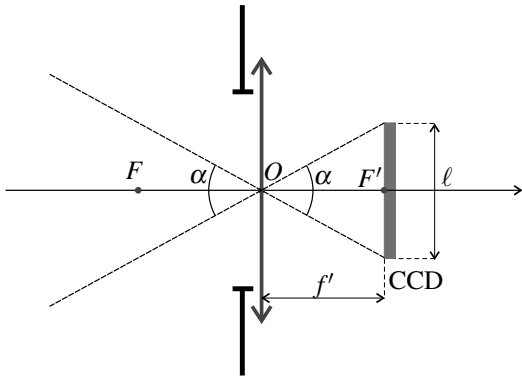


Figure 5.53 – Influence de la distance focale sur le champ angulaire.

b) Profondeur de champ

Pour expliquer l'influence de la distance focale sur la profondeur de champ, on cherche la position d'un point H de l'axe optique faisant sur le capteur une tache du diamètre maximum acceptable $\delta_{\max}$ (voir figure 5.54) .

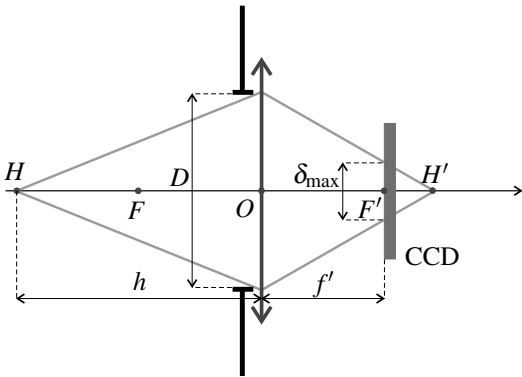


Figure 5.54 – Influence de la distance focale sur la profondeur de champ.

L'image H' de H est située derrière le capteur. D'après le théorème de Thalès :

$$\frac{\delta_{\max}}{D} = \frac{F'H'}{OH'}.$$

Pour exploiter cette relation on doit calculer les distances :

- $F'H'$ donnée par la formule de conjugaison de Newton : $F'H' = \overline{F'H'} = -\frac{f'^2}{\overline{FH}}$;
- OH' donnée par la formule de conjugaison de Descartes : $\frac{1}{OH'} = \frac{1}{OH} = \frac{1}{OH} + \frac{1}{f'}$.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

Il vient donc, en posant $h = OH$:

$$\frac{\delta_{\max}}{D} = -\frac{f'^2}{FH} \left(\frac{1}{OH} + \frac{1}{f'} \right) = -\frac{f'^2}{f' - h} \left(-\frac{1}{h} + \frac{1}{f'} \right) = \frac{f'}{h} \quad \text{soit} \quad h = \frac{f'D}{\delta_{\max}}.$$

h est appelée *distance hyperfocale*. Tous les points situés dans un plan de front à distance supérieure à h donnent des taches de diamètre $\delta < \delta_{\max}$, donc sont nets sur la photo. Plus h est petit, plus la profondeur de champ est grande.

Pour un nombre d'ouverture $N = \frac{f}{D}$ donné, la distance hyperfocale $h = \frac{f'^2}{N\delta_{\max}}$ est proportionnelle à f'^2 . Ceci explique pourquoi la profondeur de champ est plus grande sur la photo 5.51, prise avec l'objectif grand-angle, que sur la photo 5.52, prise avec le téléobjectif.

SYNTHÈSE

SAVOIRS

- approximation de l'optique géométrique
- lois de Descartes
- phénomène de réflexion totale
- image dans un miroir plan
- approximation de Gauss
- conditions de stigmatisme et aplanétisme approchés
- condition $D > 4f'$ pour la projection
- modèle de l'œil
- ordres de grandeur de la résolution angulaire de l'œil et de la plage d'accommodation

SAVOIR-FAIRE

- établir la condition de réflexion totale
- construire l'image d'un objet par un miroir
- construire l'image d'un objet par une lentille mince
- appliquer des formules de conjugaison fournies
- établir la condition $D > 4f'$ pour une projection
- modéliser un instrument d'optique à l'aide de plusieurs lentilles
- modéliser l'œil
- en comparant des photographies discuter l'influence de la durée d'exposition, du diaphragme et de la focale

MOTS-CLÉS

- | | | |
|------------------|------------------|--------------------------|
| • rayon lumineux | • rayon incident | • foyer |
| • miroir | • rayon émergent | • approximation de Gauss |
| • réflexion | • objet, image | • lentille convergente |
| • dioptrie | • réel, virtuel | • lentille divergente |
| • réfraction | • système centré | • distance focale |

S'ENTRAÎNER

5.1 Champ de vision dans un miroir (★)

- 1. Montrer que l'on voit un point A dans un miroir si et seulement si la droite $O'A$, où O' est le symétrique du point O où se trouve l'œil, coupe la surface du miroir.
- 2. En déduire quelle est la taille minimale d'un miroir dans lequel un homme de taille h peut se voir en entier. On assimilera le corps à un segment rectiligne vertical TP contenant la position O de l'œil. Comment ce miroir doit-il être placé ?

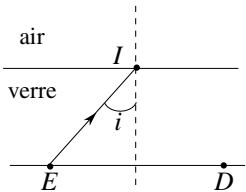
5.2 Miroir domestique (★)

Les miroirs domestiques sont des lames de verre dont la face arrière, recouverte d'un dépôt métallique d'argent, est une surface réfléchissante.

- 1. Représenter la trajectoire d'un rayon lumineux arrivant sur la lame de verre avec un angle d'incidence i . On note e l'épaisseur de la lame de verre et n son indice. Montrer que le rayon émergent du système est le même que si l'on avait uniquement une surface réfléchissante et exprimer la distance d entre cette surface et la face avant de la lame en fonction de e , i et r , angle de réfraction dans le verre correspondant à l'angle d'incidence i .
- 2. Montrer que dans les conditions de Gauss d ne dépend pas de i . Conclure.

5.3 Détection de pluie sur un pare-brise (★)

On modélise un pare-brise par une lame de verre à faces parallèles, d'épaisseur $e = 5$ mm, d'indice $n_v = 1,5$. Un fin pinceau lumineux issu d'un émetteur situé en E arrive de l'intérieur du verre sur le dioptre verre/air en I avec un angle d'incidence $i = 60^\circ$.



- 1. Montrer que le flux lumineux revient intégralement sur le détecteur situé en D et déterminer la distance ED .
- 2. Lorsqu'il pleut, une lame d'eau d'indice $n_e = 1,33$ et d'épaisseur $e' = 1$ mm se dépose sur le pare-brise. Représenter le rayon lumineux dans ce cas. À quelle distance du détecteur arrive-t-il ?

5.4 Erreur sur le positionnement d'un objet (★)

Un observateur de taille $t = 1,8$ m regarde un objet au fond d'un bassin depuis le bord. Le bassin de hauteur $h = 1,5$ m contient une épaisseur $e = 1,0$ m d'eau (indice $n = 1,33$). L'objet semble situé à une distance $d = 1$ m du bord. Déterminer la véritable position de l'objet.

5.5 Incidence de Brewster (★)

Un rayon lumineux arrive à l'interface plane séparant l'air d'un milieu d'indice n . Il se scinde en un rayon réfléchi et un rayon réfracté.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

1. Trouver l'angle d'incidence i_B , appelé angle de Brewster, pour lequel ces deux rayons sont perpendiculaires entre eux. Faire l'application numérique dans le cas de l'eau d'indice $n = 1,33$, puis d'un verre d'indice 1,5.

2. Lorsque l'angle d'incidence est i_B , la lumière réfléchie est polarisée rectilignement selon la direction perpendiculaire au plan d'incidence. Quelle application du polariseur pouvez-vous imaginer en photographie à partir de cette propriété ?

5.6 Image virtuelle avec une lentille (★)

1. Déterminer les positions des objets ayant une image virtuelle par une lentille convergente.
2. Déterminer les positions des objets ayant une image virtuelle par une lentille divergente.

5.7 Caractéristiques d'une lentille (★)

Une lentille mince sphérique donne d'un objet réel situé à 60 cm avant son centre une image droite réduite d'un facteur 5. Déterminer par le calcul et par une construction géométrique la position de l'image et les caractéristiques de la lentille.

5.8 Utilisation d'une lentille divergente (★)

Une lentille mince divergente a pour distance focale image $f' = -30$ cm.

1. Déterminer l'image d'un point A situé à 30 cm devant la lentille mince.
2. Si un objet AB dans le plan de front passant par A a pour taille 1 mm, quelle est la taille de son image ?

5.9 Doublet de Huygens (★★)

On appelle doublet un ensemble de deux lentilles minces de même axe optique. En appelant L_1 et L_2 les deux lentilles (la première lentille rencontrée par la lumière est L_1), on note O_1 et O_2 leurs centres optiques, F_1 et F_2 leurs foyers objets, F'_1 et F'_2 leurs foyers images. Un doublet est caractérisé par les distances focales image des deux lentilles f'_1 et f'_2 et son épaisseur $e = \overline{O_1 O_2}$.

Le doublet de Huygens est tel que : $f'_1 = 3a$, $e = 2a$ et $f'_2 = a$, où a est une longueur quelconque. On peut le désigner par le triplet (3, 2, 1).

1. Déterminer graphiquement la position du foyer image F' du doublet de Huygens (on pourra prendre l'échelle $a = 2$ cm).
2. Retrouver ce résultat par un calcul en déterminant l'expression de $\overline{F'_2 F'}$.
3. Déterminer graphiquement la position du foyer objet F du doublet de Huygens.
4. Retrouver ce résultat par un calcul en déterminant l'expression de $\overline{F_1 F}$.

5.10 Lentilles et système afocal (★)

1. On dispose un objet $\overrightarrow{A_0 B_0}$ orthogonalement à l'axe optique d'une lentille divergente L_1 de distance focale $f'_1 = -20$ cm. Où doit se trouver l'objet par rapport à la lentille pour que le grandissement transversal soit égal à 0,5 ?

2. Quelle est alors la position de l'image $\overrightarrow{A_1B_1}$?
3. On place après la lentille L_1 un viseur constitué d'une lentille convergente L_2 de même axe optique que L_1 et de distance focale $f'_2 = 40$ cm. On dispose également un écran perpendiculairement à l'axe optique à une distance $\overline{O_2E} = 80$ cm du centre du viseur. Calculer la distance $\overline{O_1O_2}$ entre les deux lentilles pour qu'on obtienne une image sur l'écran de l'objet initial.
4. On souhaite utiliser le dispositif précédent pour transformer un faisceau cylindrique de rayons parallèles à l'axe optique et de diamètre d en un faisceau cylindrique de rayons parallèles à l'axe optique et de diamètre D . Calculer la distance $\overline{O_1O_2}$ entre les deux lentilles pour obtenir un tel résultat.
5. Calculer dans ce cas le rapport entre les deux diamètres $\frac{D}{d}$.

APPROFONDIR

5.11 Dispersion de la lumière solaire lors d'un arc-en-ciel (★★)

Un milieu dispersif est un milieu dont l'indice optique dépend de la longueur d'onde de la lumière. Dans ce problème on s'intéresse à une conséquence du caractère dispersif de l'eau : le phénomène d'arc-en-ciel. L'indice optique de l'eau est donné par la loi de Cauchy : $n = a + \frac{b}{\lambda^2}$, où a et b sont des constantes positives caractéristiques du milieu et λ la longueur d'onde.

1. Questions préliminaires.

a. Rappeler les lois de Descartes pour la réfraction d'un rayon lumineux passant de l'air (milieu d'indice unité) vers un milieu d'indice n . On fera un schéma en notant i l'angle d'incidence et r l'angle de réfraction. Exprimer la dérivée $\frac{dr}{di}$ exclusivement en fonction de l'indice n et du sinus ($\sin i$) de l'angle d'incidence.

b. Exprimer, en fonction de i et de r , la valeur de la déviation du rayon lumineux, définie par l'angle entre la direction incidente et la direction émergente, orientées dans le sens de propagation.

c. Exprimer aussi, à l'appui d'un schéma, la déviation d'un rayon lumineux dans le cas d'une réflexion.

2. Lorsque le soleil illumine un rideau de pluie, on peut admettre que chaque goutte d'eau se comporte comme une sphère réceptionnant un faisceau de rayons parallèles entre eux.

Dans tout ce qui suit, on considérera que l'observation est faite par un oeil accommodant à l'infini, c'est-à-dire assimilable à une lentille convergente (cristallin) capable de focaliser sur un écran (rétine) tout faisceau de lumière parallèle issu d'une goutte d'eau.

On recherche, dans un premier temps, les conditions pour que la lumière émergente, issue d'une goutte d'eau, se présente sous forme d'un faisceau de lumière parallèle. Pour cela on fait intervenir l'angle de déviation D de la lumière à travers la goutte d'eau, mesuré entre le rayon émergent et le rayon incident. Cet angle de déviation D est une fonction de l'angle d'incidence i .

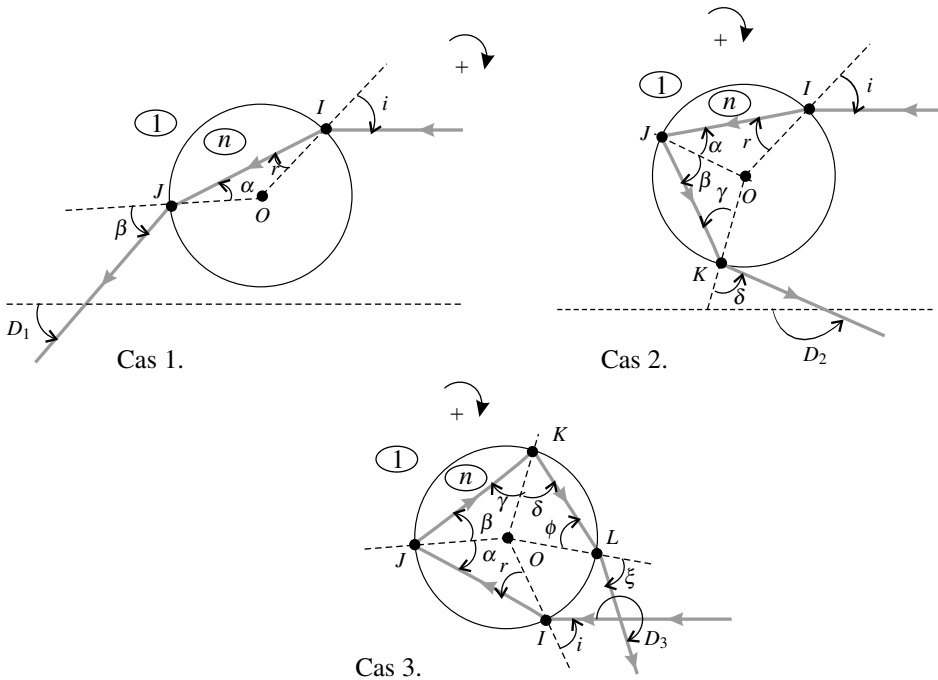
CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

Exprimer la condition de parallélisme des rayons émergents en la traduisant mathématiquement au moyen de la dérivée $\frac{dD}{di}$.

Une goutte d’eau quelconque, représentée par une sphère de centre O , est atteinte sous des incidences i variables comprises entre 0° et 90° . Son indice, pour une radiation donnée, sera noté n tandis que celui de l’air sera pris égal à l’unité.

On considère les trois cas représentés sur la figure ci-dessous :

1. lumière directement transmise,
2. lumière transmise après une réflexion partielle à l’intérieur de la goutte,
3. lumière transmise après deux réflexions partielles à l’intérieur de la goutte.



3. Pour le cas 1 :

- a. Exprimer en fonction de l’angle d’incidence i ou de l’angle de réfraction r , tous les angles marqués de lettres grecques.
- b. En déduire l’angle de déviation D propre à chaque cas, en fonction de i et de r .
- c. Rechercher ensuite, si elle existe, une condition d’émergence d’un faisceau parallèle, exprimée par une relation entre le sinus ($\sin i$) de l’angle d’incidence et l’indice n de l’eau.

4. Mêmes questions dans le cas 2.

5. Mêmes questions dans le cas 3.

6. Le soleil étant supposé très bas sur l’horizon, normal au dos d’un observateur, montrer que celui-ci ne pourra observer la lumière transmise que si la goutte d’eau se trouve sur

deux cônes d'axes confondus avec la direction solaire et de demi-angles au sommet $\theta_2 = 180^\circ + D_2 = 180^\circ - |D_2|$ (justification de l'arc primaire) et $\theta_3 = D_3 - 180^\circ$ (justification de l'arc secondaire).

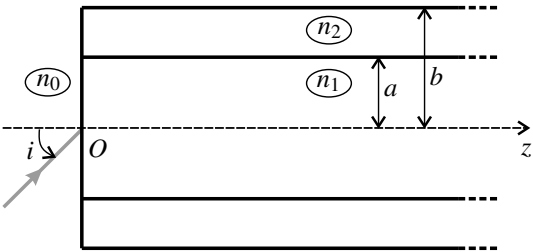
7. Les angles θ_2 et θ_3 dépendant de l'indice n de l'eau, on observe un phénomène d'irisation dû au fait que cet indice évolue en fonction de la longueur d'onde.

Calculer ces angles pour le rouge et le violet, sachant que pour le rouge l'indice vaut 1,3317 tandis que pour le violet il est égal à 1,3448.

En admettant que l'observateur se trouve face à un rideau de pluie, dessiner la figure qui apparaît dans son plan d'observation en notant la position respective des rouges et des violets.

5.12 Fibre optique à saut d'indice (★)

Le guidage de la lumière peut être assuré par des fibres optiques. Une fibre optique est constituée d'un cylindre de verre (ou de plastique) appelé cœur, entouré d'une gaine transparente d'indice de réfraction plus faible. La gaine contribue non seulement aux propriétés mécaniques de la fibre mais évite aussi les fuites de lumière vers d'autres fibres en cas de contact. Actuellement le diamètre du cœur d'une fibre varie de 3 à 200 μm selon ses propriétés et le diamètre extérieur de la gaine peut atteindre 400 μm .



On considère une fibre optique constituée d'un cœur cylindrique de rayon a et d'indice n_1 entouré d'une gaine d'indice n_2 inférieur à n_1 et de rayon b . Les faces d'entrée et de sortie sont perpendiculaires à l'axe du cylindre (Oz) formé par la fibre. L'ensemble, en particulier la face d'entrée, est en contact avec un milieu d'indice n_0 qui sera pris égal à l'indice de l'air pour les applications numériques.

1. Un rayon lumineux arrive en O . On appelle i l'angle d'incidence sur la surface d'entrée de la fibre. Déterminer en fonction de n_0 , n_1 et n_2 la condition que doit satisfaire i pour que le rayon réfracté ait une propagation guidée dans le cœur. On appelle angle d'acceptance i_a de la fibre la valeur maximale de i . Donner l'expression de i_a .
2. On appelle ouverture numérique ON de la fibre la quantité $ON = n_0 \sin i_a$. Exprimer ON en fonction de n_1 et n_2 . Application numérique : calculer la valeur de ON pour $n_1 = 1,456$ (silice) et $n_2 = 1,410$ (silicone).
3. On envoie dans la fibre un faisceau lumineux avec tous les angles d'incidence i compris entre 0 et i_a . Calculer la différence $\delta\tau$ entre la durée maximale et la durée minimale de propagation d'un bout à l'autre de cette fibre. On exprimera le résultat en fonction de la longueur L de la fibre, des indices n_1 et n_2 et de la vitesse de la lumière dans le vide $c = 3,00.10^8 \text{ m.s}^{-1}$. Application numérique : $L = 1,00 \text{ km}$, donner la valeur de $\delta\tau$.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

4. Le signal transporté par la fibre est constitué d'impulsions lumineuses d'une durée T_1 à intervalle régulier T . Quelle valeur minimale de T faut-il choisir pour que les impulsions soient distinctes à la sortie de la fibre ? Proposer une définition de la bande passante en bits (ou nombre d'impulsions) par seconde. Comparer la valeur de la bande passante obtenue ici avec celle d'un téléphone portable (64 bits par seconde) et celle de la télévision (100 Mbits par seconde).

5.13 Optique de l'œil (★)

Le cristallin de l'œil est assimilable à une lentille mince de centre optique O . On modélise l'œil par une lentille mince convergente de centre optique O , dont la vergence V est variable. L'image se forme sur la rétine, qui dans la réalité est à la distance $d_{\text{rétel}} = 15 \text{ mm}$ de O mais que l'on considérera égale à $d = 11 \text{ mm}$ pour compenser le fait qu'on néglige la présence du corps vitreux entre le cristallin et la rétine.

1. Un observateur doté d'une vision « normale » regarde un objet $\overrightarrow{AB}$ placé dans un plan de front à 1 m devant lui, et tel que $\overline{AB} = 10 \text{ cm}$.

a. Préciser si l'image formée par le cristallin est réelle ou virtuelle, droite ou renversée.

b. On note $\overrightarrow{A'B'}$ l'image de $\overrightarrow{AB}$ sur la rétine. Calculer le grandissement $\gamma = \frac{\overline{A'B'}}{\overline{AB}}$, et en déduire la taille de l'image $\overline{A'B'}$.

c. Calculer la vergence V du système.

2. L'observateur regarde maintenant un objet placé à 25 cm devant lui.

a. Préciser si l'image est réelle ou virtuelle, droite ou renversée.

b. Calculer la variation de la vergence par rapport à celle de la question (3) ainsi que la taille de l'image.

3. On s'intéresse maintenant à un sujet myope possédant donc un cristallin trop convergent. Lorsqu'il regarde à l'infini, l'image se forme à 0,5 mm en avant de la rétine (située à $d = 11 \text{ mm}$ de O). Pour corriger ce problème, cette personne est dotée de lunettes dont chaque verre est assimilé à une lentille mince de vergence V' constante et de centre optique O' , placé à $\ell = 2 \text{ cm}$ de O .

a. Calculer la vergence V' des verres de lunettes.

b. L'individu observe un objet situé à 1 m devant lui. Calculer la position de l'image intermédiaire ainsi que le grandissement de l'ensemble (lunette-cristallin).

5.14 Étude d'une lunette astronomique (★)

Une lunette astronomique est schématisée par deux lentilles minces convergentes de même axe optique Δ :

- l'une L_1 (objectif) de distance focale image $f'_1 = \overline{O_1F'_1}$;
- l'autre L_2 (oculaire) de distance focale image $f'_2 = \overline{O_2F'_2}$.

On rappelle qu'un œil normal voit un objet sans accommoder si celui-ci est placé à l'infini.

On souhaite observer la planète Mars qui est vue à l'œil nu sous un diamètre apparent α .

- Pour observer la planète avec la lunette, on forme un système afocal.
 - Que signifie l'adjectif afocal ? En déduire la position relative des deux lentilles.
 - Faire le schéma de la lunette pour $f'_1 = 5f'_2$. Dessiner sur ce schéma la marche à travers la lunette d'un faisceau lumineux (non parallèle à l'axe) formé de rayons issus de l'astre. On appelle $A'B'$ l'image intermédiaire.
 - On souhaite photographier cette planète. Où faut-il placer le capteur CCD ?
- On note α' l'angle que forment les rayons émergents extrêmes en sortie de la lunette.
 - L'image est-elle droite ou renversée ?
 - La lunette est caractérisée par son grossissement $G = \frac{\alpha'}{\alpha}$. Exprimer G en fonction de f'_1 et f'_2 .
- On veut augmenter le grossissement de cette lunette et redresser l'image. Pour cela, on interpose entre L_1 et L_2 une lentille convergente L_3 de distance focale image $f'_3 = \overline{O_3F'_3}$. L'oculaire L_2 est déplacé pour avoir de la planète une image nette à l'infini à travers le nouvel ensemble optique.
 - Quel couple de points doit conjuguer L_3 pour qu'il en soit ainsi ?
 - On appelle γ_3 le grandissement de la lentille L_3 . En déduire $\overline{O_3F'_1}$ en fonction de f'_3 et γ_3 .
 - Faire un schéma (on placera O_3 entre F'_1 et F_2 et on appellera $\overline{A'B'}$ la première image intermédiaire et $\overline{A''B''}$ la seconde image intermédiaire).
 - En déduire le nouveau grossissement G' en fonction de G et γ_3 . Comparer G' à G en signe et valeur absolue.

5.15 La loupe (★★)

Un observateur emmétrope, c'est-à-dire ayant un œil normal, peut voir distinctement de l'infini à une distance minimale d_m . On dit que l'observateur accommode si l'objet qu'il observe n'est pas à l'infini.

Cet observateur regarde à l'œil nu un tout petit objet plan que l'on assimilera à un segment AB de longueur ℓ , perpendiculaire à l'axe optique.

- Déterminer α_m , angle maximal sous lequel l'objet peut être vu.
- L'observateur regarde AB à travers une lentille mince convergente de distance focale f' et de centre O (loupe). Son œil est situé à une distance a de la loupe ($a < d_m$).
 - Déterminer les positions de l'objet rendant possible l'observation d'une image nette par l'observateur emmétrope. Faire une construction géométrique de l'image. L'image est-elle droite ou renversée ?
 - Pour quelle position de l'objet l'observation se fait-elle sans accommodation ? Exprimer l'angle α sous lequel l'œil voit l'image. Application numérique : que vaut le grossissement commercial de la loupe $G = \alpha/\alpha_m$? On donne $d_m = 0,25$ m et $f' = 50$ mm.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

5.16 Principe du viseur (★★)

Pour aborder cet exercice, il faut avoir fait l'exercice précédent.

Un viseur est constitué d'un objectif et d'un oculaire de même axe optique (Ox). On assimilera l'objectif à une lentille mince convergente L_1 de centre O_1 et de distance focale f'_1 et l'oculaire à une lentille mince convergente L_2 de centre O_2 et de distance focale f'_2 . On pose $\overline{O_1O} = D$ et $\overline{OO_2} = d$ (les distances D et d sont positives et réglables). Dans un plan transverse est disposé en O une croix appelée réticule, constituée de deux traits fins perpendiculaires. L'observateur place son œil à une distance a derrière l'oculaire ($a < d_m$).

1. Quel est l'intervalle des valeurs de d permettant à l'observateur de voir le réticule net à travers l'oculaire ? Où doit-on placer le réticule pour l'observer sans accommodation ? On effectuera les constructions dans les deux cas extrêmes de d .

2. Le réglage précédent est supposé réalisé. On souhaite observer à travers le viseur un point objet A situé sur l'axe optique à l'abscisse $x = \overline{OA}$. L'image intermédiaire doit se trouver dans le plan du réticule.

a. Montrer que $x \leq -4f'_1$

b. Déterminer l'expression de D en fonction de x .

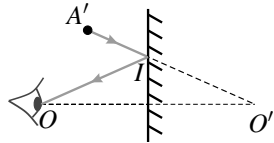
c. En déduire la plage de réglage de la distance D que le constructeur doit prévoir.

3. Un observateur myope souhaite utiliser le viseur sans ses verres correcteurs pour observer un objet A situé à l'infini, dans les conditions définies précédemment. Sachant que sa distance maximale de vision distincte est δ , calculer les valeurs des réglages qu'il doit effectuer.

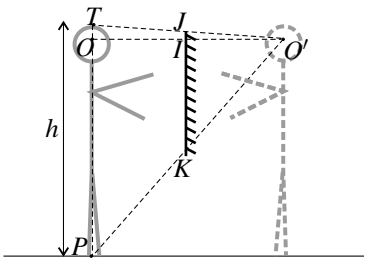
CORRIGÉS

5.1 Champ de vision dans un miroir

1. L'œil voit le point A « dans le miroir » s'il y a un rayon partant de A qui arrive en O après réflexion sur le miroir en I. Ce rayon peut aussi être parcouru dans l'autre sens et dans ce cas, après la réflexion, il semble provenir de O'. Ainsi la droite O'A coupe le miroir en I.

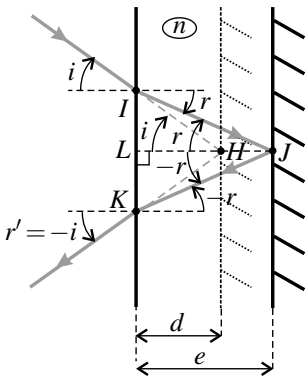


2. On applique le résultat de la question précédente. La figure montre le miroir le plus petit permettant à l'homme de se voir en entier. D'après le théorème de Thalès : $\frac{JK}{TP} = \frac{O'I}{O'O} = \frac{1}{2}$. La hauteur de ce miroir est donc $JK = \frac{1}{2}TP = \frac{h}{2}$. le bas du miroir doit être à une hauteur égale la moitié de la hauteur des yeux au dessus du sol.



5.2 Miroir domestique

1. Le rayon est réfracté à son entrée dans le verre d'indice n. L'angle de réfraction est r tel que : $n \sin r = \sin i$ (en prenant l'indice de l'air égal à 1. Le rayon réfracté rencontre avec un angle d'incidence égal à r (d'après la propriété des angles alterne-interne) la face arrière sur laquelle il est réfléchi. Il revient sur la surface du verre, avec un angle d'incidence égal à -r (propriété des angles alterne-interne) et est réfracté avec un angle de réfraction r' tel que $\sin r' = n \sin(-r) = -n \sin r = -\sin i$, soit $r' = -i$.



Le système donne le même rayon émergent qu'une surface réfléchissante plane parallèle aux faces de la lame et située à l'intersection des prolongements du rayon incident et du rayon émergent (représentée en pointillés sur la figure). La distance d entre ce miroir équivalent et la face avant de la lame est :

$$d = HL = IL \tan i = \left(\frac{LI}{\tan r} \right) \tan i = e \frac{\tan r}{\tan i}.$$

2. Dans les conditions de Gauss, $i \ll 1$ donc $\tan i \simeq i$. De plus $\sin r = \frac{\sin i}{n} \simeq \frac{i}{n} \ll 1$ donc $r \ll 1$ et $r \simeq \sin r \simeq \frac{i}{n}$; enfin $\tan r \simeq r \simeq \frac{i}{n}$. Il vient alors : $d \simeq \frac{e}{n}$, valeur indépendante de i.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

Ainsi, dans les conditions de Gauss, tous les rayons émergents semblent avoir été réfléchis par une surface réfléchissante plane située à $\frac{e}{n}$ derrière la face avant. Le miroir domestique a donc les propriétés de stigmatisme et d’aplanétisme approchées dans les conditions de Gauss.

5.3 Détection de pluie sur un pare-brise

1. Il peut y avoir réflexion totale en I , car le rayon passe du verre à l’air qui est moins réfringent. On calcule l’angle de réflexion limite en I : $R_{\text{lim},v} = \arcsin \frac{1}{n_v} = 41,8^\circ$ (on a pris 1 pour l’indice de l’air). Il y a réflexion totale en I si $i > R_{\text{lim},v}$, ce qui est le cas pour $i = 60^\circ$. Le flux lumineux est donc entièrement réfléchi.

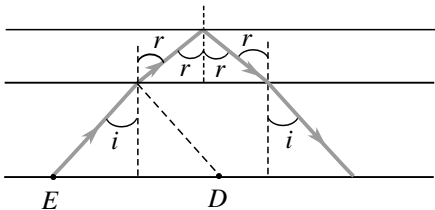
Pour trouver la position du détecteur, on applique la loi de réflexion en I : le rayon est réfléchi symétriquement par rapport à la normale : la distance entre l’émetteur E et le détecteur D doit donc être : $ED = 2e \tan i = 1,7 \text{ cm}$.

2. Les deux dioptrés sont maintenant verre-eau et eau-air. Dans les deux cas, il peut y avoir réflexion totale. On calcule les deux angles de réflexion totale :

- Pour le dioptré verre-eau :
 $n_v \sin R_{\text{lim},1} = n_e \sin \pi/2$ soit $R_{\text{lim},1} = 62,5^\circ$.
- Pour le dioptré eau-air :
 $n_e \sin R_{\text{lim},2} = \sin \pi/2$ soit $R_{\text{lim},2} = 48,75^\circ$.

L’angle i est inférieur à $R_{\text{lim},1}$, il y a donc réfraction sur le premier dioptré. On calcule l’angle r après réfraction : $n_v \sin i = n_e \sin r$ donc $r = 77,6^\circ$.

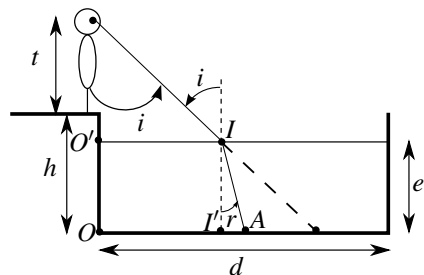
L’angle r est aussi l’angle d’incidence sur le deuxième dioptré, il y a donc réflexion totale sur ce dioptré puisque $r > R_{\text{lim},2}$. On obtient la marche du rayon lumineux représenté sur la figure. On constate que le rayon lumineux ne tombe plus sur le détecteur mais à distance $2e' \tan r = 0,9 \text{ cm}$ de celui-ci. Un système de commande relié au détecteur peut alors déclencher les essuie-glaces.



5.4 Erreur sur le positionnement d’un objet

L’observateur a l’impression que l’objet est dans la direction du rayon qu’il reçoit. On voit sur la figure que le rayon provenant de l’objet s’est éloigné de la normale au dioptré (air-eau) puisque $n > 1$. On note A la position réelle de l’objet et B sa position apparente.

- Dans le triangle TOB :
 $d = OB = (h + t) \tan i$ d’où $i = 16,9^\circ$.
- Dans le triangle $TO'I$: $O'I = (h - e + t) \tan i$.
- Dans le triangle $II'A$: $AI' = e \tan r$.
- Loi de Descartes en I : $\sin i = n \sin r$.



D'autre part, les angles sont inférieurs à $\frac{\pi}{2}$ donc : $0 \cos r = \sqrt{1 - \sin^2 r} = \sqrt{1 - \left(\frac{\sin i}{n}\right)^2}$.

On en déduit que : $OA = O'I + I'A = e \tan r + (h - e + t) \tan i$ donc :

$$OA = e \frac{\sin i}{n} \frac{1}{\sqrt{1 - \left(\frac{\sin i}{n}\right)^2}} + (h - e + t) \tan i = 0,92 \text{ m.}$$

5.5 Incidence de Brewster

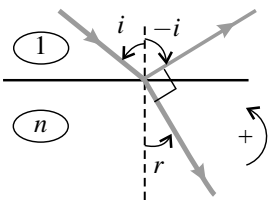
1. D'après la figure on doit avoir :

$$r + \frac{\pi}{2} + i = \pi \text{ soit } r = \frac{\pi}{2} - i.$$

La loi de la réflexion, $\sin i = n \sin r$, donne alors : $\sin i = n \cos i$ soit $\tan i = n$. L'angle d'incidence de Brewster est : $i_B = \arctan(n)$.

Application numérique : pour l'eau $i_{B,\text{eau}} = 53^\circ$ et pour le verre $i_{B,\text{verre}} = 56^\circ$.

2. En munissant l'objectif d'un polariseur, on peut réduire voire éliminer (si l'incidence est l'incidence de Brewster) les reflets sur l'eau ou sur une vitre en orientant convenablement le polariseur.



5.6 Image virtuelle avec une lentille

On note O le centre optique de la lentille, A l'objet, A' l'image et f' la distance focale image de la lentille.

Si l'image est virtuelle, cela signifie que $\overline{OA'} < 0$. Puisque dans l'exercice on impose une condition sur $\overline{OA}$, on a intérêt à utiliser la relation de conjugaison de Descartes (avec origine en O) : $\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{f'}$. La condition $\overline{OA'} < 0$ est équivalente à $\frac{1}{f'} + \frac{1}{\overline{OA}} < 0$ soit $\frac{1}{f'} < \frac{1}{\overline{AO}}$.

1. Pour la lentille convergente, comme f' est positive, on en déduit que $\overline{AO}$ est aussi positive donc A est réelle et $AO < f'$: A est entre F et O . Cette situation correspond à l'usage de la lentille convergente en loupe.

2. Pour la lentille divergente, f' est négative. Il y a donc deux possibilités :

- si $\overline{AO} > 0$ l'inégalité est toujours vérifiée : tout objet réel donne une image virtuelle.
- si $\overline{AO} < 0$ et (en faisant attention aux signes) $OA > -f'$: A est situé après la lentille et au-delà du foyer objet F (on rappelle que pour une lentille divergente le foyer objet est après la lentille.).

5.7 Caractéristique d'une lentille

Par le calcul, on utilise la formule de conjugaison avec origine au centre $\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{f'}$ et

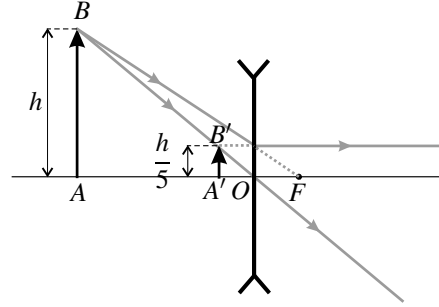
l'expression du grandissement avec origine au centre $\gamma = \frac{\overline{OA'}}{\overline{OA}}$. L'image se trouve en A' tel

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

que $\overline{OA'} = \gamma \overline{OA} = -12 \text{ cm}$ et la distance focale est $f' = \left(\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} \right) = \frac{\gamma}{1-\gamma} \overline{OA} = -15 \text{ cm}$.

La lentille est donc divergente.

Pour la construction, on place l'objet $\overline{AB}$ et la lentille (dont on ne connaît pas encore la nature) et on trace le rayon BO qui est non dévié. L'image étant droite et 5 fois plus petite, on place facilement B' sur ce rayon (on constate que $OA' = \frac{1}{5} OA$, ce qui provient du théorème de Thalès). On trace ensuite le rayon émergent parallèlement à l'axe semblant venir de B' : le prolongement de ce rayon avant la lentille passe par le foyer objet F que l'on détermine ainsi.



5.8 Utilisation d'une lentille divergente

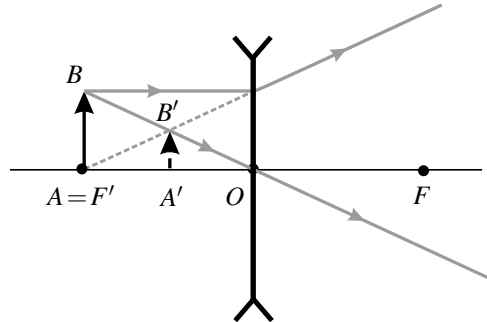
1. On remarque que l'objet A est confondu avec le foyer image de la lentille. Pour trouver l'image A' , on applique par exemple la relation de conjugaison de Newton :

$$\overline{F'A'} \cdot \overline{FA} = -f'^2.$$

Avec $\overline{FA} = 2f'$, on trouve :

$$\overline{F'A'} = -\frac{f'}{2} = \frac{1}{2} \overline{F'O}.$$

Ainsi A' est à mi-distance de F' et O .

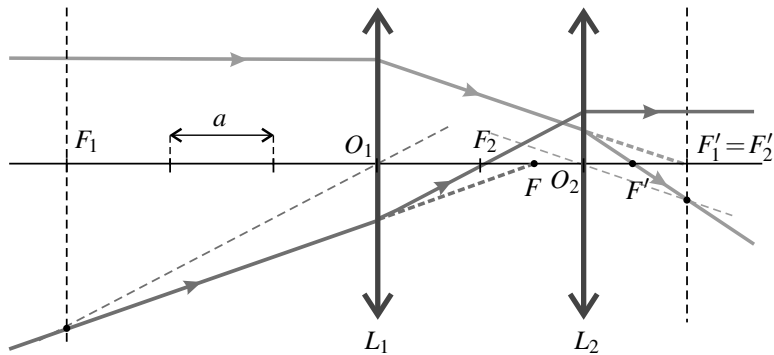


2. Pour répondre à cette question il faut calculer le grandissement entre les plans de front conjugués passant par A et A' . La formule avec l'origine au centre découle directement de la relation de Thalès dans le triangle OAB : $\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OA'}}{\overline{OA}} = \frac{1}{2}$. On en déduit que $\overline{A'B'} = 0,5 \text{ mm}$.

5.9 Doublet de Huygens

1. F' est l'image du point à l'infini dans la direction de l'axe optique. Tout rayon provenant de ce point est parallèle à l'axe. Un tel rayon, après avoir traversé L_1 passe par F'_1 . Il est ensuite dévié par L_2 mais le tracé n'est pas immédiat car ce n'est pas un rayon remarquable pour cette lentille. Pour construire le rayon émergent de L_2 , on utilise la règle numéro 4 de l'encadré du paragraphe 5.2.a). On trace un rayon auxiliaire qui est parallèle au rayon incident sur L_2 et passe par le centre optique O_2 de L_2 . Ce rayon n'est pas dévié par L_2 . Le rayon émergent cherché passe par l'intersection du rayon auxiliaire avec le plan focal image

de L_2 (foyer image secondaire). Finalement, l'intersection du rayon émergent du système avec l'axe optique est le foyer F' cherché.



2. Le point à l'infini sur l'axe a pour image F'_1 par L_1 qui a pour image F' par L_2 , ce que l'on peut résumer par le schéma : $\infty \xrightarrow{L_1} F'_1 \xrightarrow{L_2} F'$. D'après la relation de conjugaison de Newton : $\overline{F'_2 F'} = -\frac{(f'_2)^2}{\overline{F_2 F'_1}} = -\frac{a^2}{2a} = -\frac{a}{2}$.

3. Par définition, tout rayon incident sur le doublet passant par F donne un rayon émergent parallèle à l'axe. On trace un rayon émergent du système parallèle à l'axe quelconque ; ce rayon, avant la lentille L_2 , passait par son foyer objet F_2 , ce qui permet de le tracer entre les deux lentilles. Le tracé avant L_1 n'est pas immédiat car ce rayon n'est pas un rayon remarquable pour L_1 . On applique la règle 5 de l'encadré du paragraphe 5.2.a) : on trace un rayon auxiliaire parallèle au rayon, passant par le centre optique O_1 de L_1 . Ce rayon auxiliaire n'est pas dévié par L_1 . Le rayon incident sur L_1 cherché passe par l'intersection du rayon auxiliaire avec le plan focal objet de L_1 (foyer objet secondaire).

Finalement, l'intersection du rayon incident avec l'axe optique est le foyer F' cherché. Dans le cas présent il s'agit de l'intersection du prolongement de ce rayon avec l'axe : le foyer objet F du système est virtuel.

4. La succession des images est : $F \xrightarrow{L_1} F_2 \xrightarrow{L_2} \infty$. D'après la relation de conjugaison de Newton : $\overline{F_1 F} = -\frac{(f'_1)^2}{\overline{F_1 F_2}} = -\frac{9a^2}{-2a} = \frac{9a}{2}$.

Remarque

Les méthodes utilisées dans cet exercice s'appliquent à tous les doublets. On trouve dans le cas général :

$$\overline{F'_2 F'} = \frac{(f'_2)^2}{-f_1 + e + f'_2} \quad \text{et} \quad \overline{F_1 F} = \frac{-(f'_1)^2}{-f_1 + e + f'_2}.$$

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

5.10 Lentilles et système afocal

1. La formule de conjugaison s'écrit $\frac{1}{\overline{O_1A_1}} - \frac{1}{\overline{O_1A_0}} = \frac{1}{f'_1}$ et l'expression du grandissement est $\gamma = \frac{\overline{A_1B_1}}{\overline{A_0B_0}} = \frac{\overline{O_1A_1}}{\overline{O_1A_0}}$. On en déduit $\overline{O_1A_0} = \frac{(1-\gamma)f'_1}{\gamma} = -20 \text{ cm}$.

2. La position de l'image est donnée par $\overline{O_1A_1} = \gamma \overline{O_1A_0} = -10 \text{ cm}$.

3. Pour avoir une image sur l'écran de l'objet initial, il faut que L_2 donne une image de $\overrightarrow{A_1B_1}$ sur l'écran. En appliquant la formule de conjugaison, on obtient :

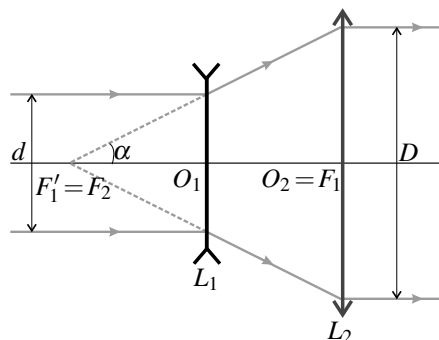
$$\overline{O_1O_2} = \frac{\overline{O_2E}f'_2}{\overline{O_2E} - f'_2} + \overline{O_1A_1} = 70 \text{ cm}.$$

4. Pour réaliser l'opération souhaitée, le système doit être afocal donc $F'_1 = F_2$. Par conséquent, on en déduit que $\overline{O_1O_2} = 20 \text{ cm}$.

5. On détermine l'expression de $\tan \alpha$, en notant α l'angle que fait le rayon issu de F'_1 avec l'axe optique entre les deux lentilles, soit :

$$\tan \alpha = \frac{d}{-f'_2} = \frac{D}{f'_1}. \text{ On en déduit :}$$

$$\frac{D}{d} = -\frac{f'_2}{f'_1} = 2.$$



5.11 Dispersion de la lumière solaire lors d'un arc-en-ciel

1. a. Loi de Descartes pour la réfraction : le rayon réfracté appartient au plan d'incidence et $\sin i = n \sin r$. On dérive cette loi par rapport à i sachant que n est indépendant de i :

$$\cos i = n \cos r \frac{dr}{di} \Rightarrow \frac{dr}{di} = \frac{1}{n} \frac{\cos i}{\cos r}.$$

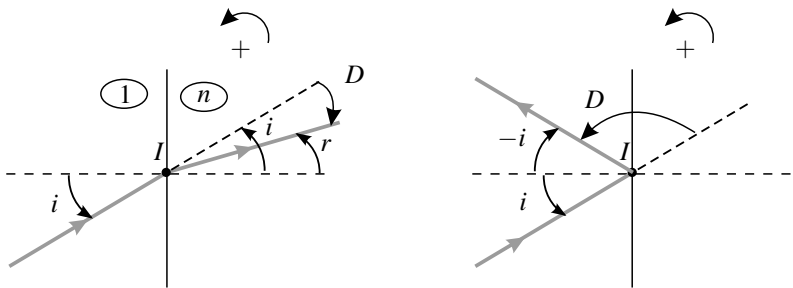
Les angles sont dans l'intervalle $\left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$, donc le cosinus est positif ; on peut utiliser la

formule $\cos = \sqrt{1 - \sin^2}$ pour trouver :

$$\frac{dr}{di} = \frac{1}{n} \frac{\sqrt{1 - \sin^2 i}}{\sqrt{1 - \sin^2 r}} = \frac{1}{n} \frac{\sqrt{1 - \sin^2 i}}{\sqrt{1 - \left(\frac{\sin i}{n}\right)^2}}.$$

b. La déviation est l'angle dont a tourné le rayon lumineux par rapport à la direction incidente. Dans le cas de la réfraction (figure ci-dessous, à gauche), la déviation D est égale à $r - i$.

c. Dans le cas de la réflexion (figure ci-dessous à droite), la déviation vaut $D = \pi - 2i$.



2. Les rayons émergents sont parallèles quel que soit l'angle d'incidence i si la déviation D est indépendante de i soit $\frac{dD}{di} = 0$.

3. Cas 1

a. Le triangle OIJ est isocèle, donc $\alpha = -r$. Dans ce cas, la loi de Descartes en J s'écrit $\sin \beta = n \sin(-r) = -n \sin r = -\sin i$ donc $\beta = -i$.

b. La déviation en I est $D = r - i$, et celle en J est $D' = \beta - \alpha = r - i$, soit globalement :

$$D_1 = D + D' = 2(r - i). \tag{5.9}$$

c. On calcule la dérivée de D_1 par rapport à i : $\frac{dD_1}{di} = 2 \frac{dr}{di} - 2 = 2 \frac{1}{n} \frac{\sqrt{1 - \sin^2 i}}{\sqrt{1 - (\frac{\sin i}{n})^2}} - 2$.

La dérivée est nulle si $\sqrt{1 - \sin^2 i} = n \sqrt{1 - (\frac{\sin i}{n})^2}$. En élevant au carré et en simplifiant, on trouve $n = 1$, ce qui n'a pas d'intérêt. Dans ce cas, il n'est pas possible d'avoir des rayons émergents parallèles.

4. Cas 2

a. L'application de la loi de Descartes pour la réflexion en J donne $\beta = -\alpha$. D'autre part, les triangles OIJ et OJK sont isocèles donc $r = \beta = -\alpha = -\gamma$. Alors, l'application de la loi de Descartes en K donne $\sin \delta = n \sin \gamma = -n \sin r = -\sin i$ soit $\delta = -i$.

b. Pour la déviation, on a la somme de trois termes : la déviation en I , $D = r - i$; la déviation en J , $D' = \pi - 2\alpha = \pi + 2r$; la déviation en K , $D'' = \delta - \gamma = -i + r$. Donc :

$$D_2 = D + D' + D'' = \pi + 4r - 2i. \tag{5.10}$$

c. On dérive D_2 : $\frac{dD_2}{di} = 4 \frac{dr}{di} - 2 = 4 \frac{1}{n} \frac{\sqrt{1 - \sin^2 i}}{\sqrt{1 - (\frac{\sin i}{n})^2}} - 2$. Cette dérivée s'annule

si :

$$\sin^2 i = \frac{4 - n^2}{3}. \tag{5.11}$$

5. Cas 3

a. De manière analogue : $r = \beta = \delta = -\alpha = -\gamma = -\phi$ et $\xi = -i$.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

b. Il vient : $D_3 = (r - i) + (\pi - 2\alpha) + (\pi - 2\gamma) + \xi - \phi = (r - i) + (\pi + 2r) + (\pi + 2r) + (-i + r)$ soit :

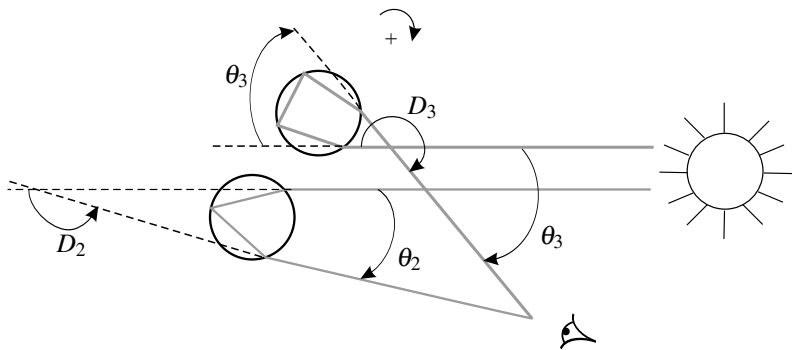
$$D_3 = 2\pi + 6r - 2i. \tag{5.12}$$

c. On dérive D_3 : $\frac{dD_3}{di} = 6 \frac{dr}{di} - 2 = 6 \frac{1}{n} \frac{\sqrt{1 - \sin^2 i}}{\sqrt{1 - \left(\frac{\sin i}{n}\right)^2}} - 2$. Cette dérivée s'annule si :

$$\sin^2 i = \frac{9 - n^2}{8}. \tag{5.13}$$

6. Le soleil étant très bas sur l'horizon, on peut supposer que les rayons incidents sont horizontaux. La figure ci-dessous représente les deux cas possibles. Les angles θ_2 et θ_3 entre les rayons incidents pour l'œil et les rayons provenant du soleil valent :

$$\begin{cases} \theta_2 = \pi + D_2 = \pi - |D_2| \\ \theta_3 = D_3 - \pi \end{cases}$$

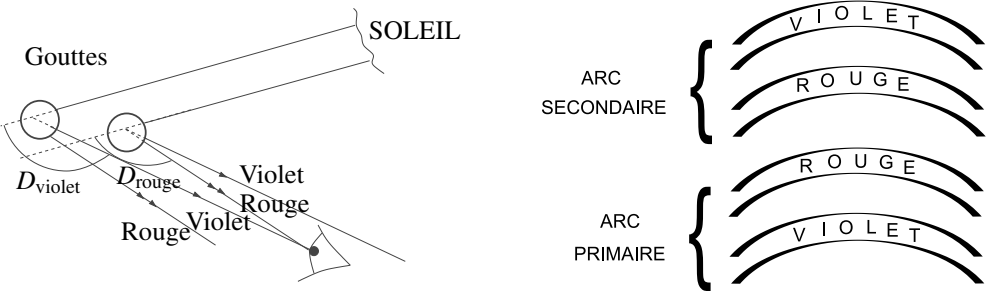


7. Pour l'arc primaire, on calcule i avec la relation (5.11), on en déduit r avec la loi de Descartes puis D_2 avec la relation (5.10). On procède de même pour D_3 avec les relations (5.13) et (5.12).

couleur	i	r	$ D_2 $	$\theta_2(\text{rad})$	$\theta_2(^{\circ})$
rouge	1.0382	0.7035	2.4039	0.7377	42.2675
violet	1.0250	0.6887	2.4366	.7050	40.3927
couleur	i	r	D_3	θ_3	$\theta_3(^{\circ})$
rouge	1.2546	0.7948	4.0238	0.8823	50.5494
violet	1.2473	0.7825	4.0830	0.9414	53.9363

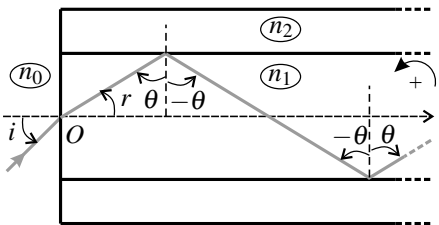
L'angle θ_2 est plus grand pour le rouge que pour le violet. Les rayons rouges observés par l'œil proviennent donc de gouttes de plus haute altitude, car dans ce cas les rayons violets « passent au-dessus de l'œil », et ce sont les rayons violets des gouttes de plus basse altitude qui entrent dans l'œil. On voit donc le violet en dessous du rouge pour l'arc primaire.

En ce qui concerne l'arc secondaire, il n'est pas toujours visible car beaucoup moins intense (il y a deux réflexions dans la goutte). On remarque que $\theta_3 > \theta_2$, donc les gouttes donnant l'arc secondaire visible par l'œil seront à plus haute altitude que celles donnant l'arc primaire. L'arc secondaire est donc vu au-dessus de l'arc primaire. D'autre part, les couleurs sont inversées car θ_3 est plus grand pour le violet que pour le rouge. Les figures suivantes récapitulent l'ordre des couleurs.



5.12 Fibre optique

1. Le rayon est guidé le long de la fibre s'il y a réflexion totale sur la surface de séparation cœur - gaine, ce qui est possible puisque $n_2 < n_1$. Pour cela il faut que θ , angle d'incidence du rayon sur ce dioptre, vérifie $|\sin \theta| > \frac{n_2}{n_1}$ (attention : θ est négatif sur la figure). Or $\theta = r - \frac{\pi}{2}$ où r est l'angle de réfraction à l'entrée du rayon dans le cœur, donné par la loi de Descartes : $n_1 \sin r = n_0 \sin i$.



On a donc : $|\sin \theta| = \cos r = \sqrt{1 - \sin^2 r} = \sqrt{1 - \left(\frac{n_0 \sin i}{n_1}\right)^2}$. La condition de réflexion

totale devient : $\frac{n_0^2}{n_1^2} \sin^2 i < 1 - \frac{n_2^2}{n_1^2}$, soit : $-i_a < i < i_a$ avec $i_a = \text{Arcsin} \frac{\sqrt{n_1^2 - n_2^2}}{n_0}$.

2. D'après la définition, $ON = \sqrt{n_1^2 - n_2^2}$. L'application numérique donne $ON = 0,363$.
3. Le temps le plus court est celui obtenu en se déplaçant sur l'axe. On obtient comme valeur $\frac{Ln_1}{c}$. Le temps le plus long correspond à l'angle d'acceptance i_a soit $\frac{Ln_1}{c \cos \theta}$. On en déduit $\delta \tau = \frac{Ln_1}{c} \left(\frac{1}{\cos \theta} - 1 \right)$. L'application numérique donne $\delta \tau = 0,158 \mu s$.

4. Il faut que $T > \delta \tau$ soit $f < \frac{1}{\delta \tau} = \Delta f = 6,3 \cdot 10^6$ bits par seconde. La fibre étudiée convient pour le téléphone mais pas pour la télévision.

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

5.13 Optique de l'œil

1. a. L'objet est réel et l'image se formant sur la rétine (qui joue le rôle d'écran) est donc réelle elle-aussi. La seule possibilité d'objet réel donnant une image réelle correspond au cas d'une lentille convergente (ici le cristallin de l'œil), avec l'objet avant F et l'image après F' et dans ce cas l'image est renversée.

b. On connaît les distances de l'objet ($\overline{OA} = -D$) et de l'image au centre optique, donc on utilise la formule de grandissement avec origine en O :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OA'}}{\overline{OA}} = \frac{d}{D} = 11 \cdot 10^{-3}.$$

On en déduit la taille de l'image : $\overline{A'B'} = \gamma \overline{AB} = 1,1 \text{ mm}$.

c. Pour calculer la vergence du système, on applique la relation de Descartes :

$$V = \frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{d} - \frac{1}{-D} = 92 \delta.$$

2. a. On ne peut se servir des résultats des questions précédentes. Le fait que la position de la rétine (position de l'image) soit inchangée alors que la position de l'objet a changé implique que la vergence V de la lentille a évolué. Le raisonnement est le même qu'à la question (1) : image réelle renversée puisqu'elle se forme sur un écran (la rétine) et que l'objet est réel.

b. On applique de nouveau la loi de Descartes :

$$V = \frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{11 \cdot 10^{-3}} - \frac{1}{-25 \cdot 10^{-2}} = 491 \delta.$$

On trouve que le cristallin s'est contracté (la vergence a augmenté) pour voir plus près. La taille de l'image est obtenue par l'expression du grandissement :

$$\overline{A'B'} = \gamma \overline{AB} = \frac{\overline{OA'}}{\overline{OA}} \overline{AB} = \frac{0,11}{-0,25} 0,1 = -4,4 \text{ mm}.$$

3. a. On calcule tout d'abord la vergence V du cristallin de l'œil myope. L'image d'un objet à l'infini se forme dans le plan focal image de la lentille donc à une distance f' de O soit $f' = 11 - 0,5 = 10,5 \text{ mm}$ soit $V = 95,2 \delta$. On appelle L_D le verre de lunette et R le point d'intersection de l'axe optique avec la rétine. Pour que le myope voit net l'objet à l'infini, on doit avoir la suite de conjugaisons suivante : $\infty \xrightarrow{L_D} F'_D \xrightarrow{\text{cristallin}} R$. Il s'agit donc de trouver la position de l'antécédent de R par le cristallin. Avec $\overline{OR} = d$, la relation de Descartes donne :

$$V = \frac{1}{d} - \frac{1}{\overline{OF'_D}} \Rightarrow \overline{OF'_D} = \frac{d}{1 - Vd} = -23,3 \text{ cm}.$$

On remarque que F'_D est à 23,3 cm avant O donc à 21,3 cm avant L_D ce qui correspond bien à une lentille divergente de distance focale image $f'_D = -21,3 \text{ cm}$, donc de vergence $V' = 1/f'_D = -4,7 \delta$.

b. Dans ce cas, la vergence du cristallin a changé par rapport à la question précédente, puisque l'œil n'observe plus à l'infini. On sait cependant que l'image finale A' est sur la rétine ($\overline{OA'} = d$). On calcule la position de l'image intermédiaire A_i par la formule Descartes appliquée à L_D :

$$\frac{1}{\overline{O'A_i}} - \frac{1}{\overline{O'A}} = V' \Rightarrow \overline{O'A_i} = \frac{\overline{O'A}}{1 + \overline{O'A} V'} = -17,5 \text{ cm.}$$

On en déduit $\overline{OA_i} = -19,5 \text{ cm.}$

Le grandissement de l'ensemble est le produit des grandissements :

$$\gamma = \frac{\overline{O'A_i} \overline{OA'}}{\overline{O'A} \overline{OA_i}} = -0,01.$$

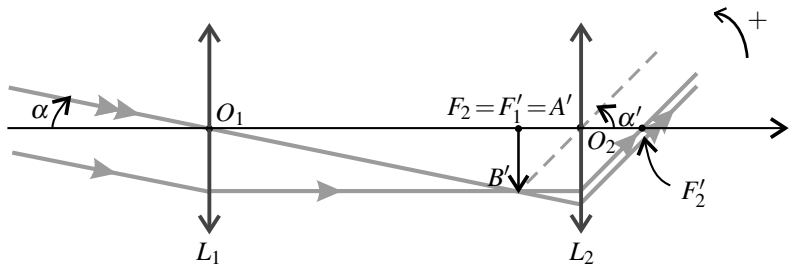
5.14 Étude d'une lunette astronomique

1. a. Un système est dit afocal s'il conjugue l'infini avec l'infini, c'est-à-dire si un objet à l'infini a une image à l'infini.

L'image d'un objet à l'infini par L_1 est dans le plan focal image de L_1 . Cette image est un objet pour L_2 , qui en donne une image à l'infini, donc elle doit être dans le plan focal objet de L_2 , d'où la série des conjugaisons :

$$\infty \xrightarrow{L_1} F'_1 = F_2 \xrightarrow{L_2} \infty.$$

b. Le schéma est sur la figure ci-dessous. Pour le tracer on utilise la règle 5 du paragraphe 5.2.a) d'après laquelle les rayons émergeant de L_2 et provenant de B' sont tous parallèles entre eux et donc parallèles au rayon non dévié $B'O_2$.



c. La lentille oculaire L_2 est nécessaire pour observer directement à l'œil (qui observe à l'infini pour ne pas accommoder). Si l'on souhaite prendre une photo, on la supprime et on place le capteur CCD dans le plan focal image de l'objectif L_1 .

2. a. Sur la figure ci-dessus on remarque que pour un faisceau lumineux incident provenant du haut, le faisceau émergent semble provenir du bas : l'image est donc renversée, ce qui est confirmé par le sens différent des angles α et α' .

CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

b. Dans les triangles $O_1A'B'$ et $O_2A'B'$ (on rappelle que $A' = F_2 = F'_1$), on écrit :

$$\tan \alpha = \frac{\overline{A'B'}}{\overline{O_1A'}} = \frac{\overline{A'B'}}{f'_1} \quad \text{et} \quad \tan \alpha' = -\frac{\overline{A'B'}}{\overline{A'O_2}} = -\frac{\overline{A'B'}}{f'_2}.$$

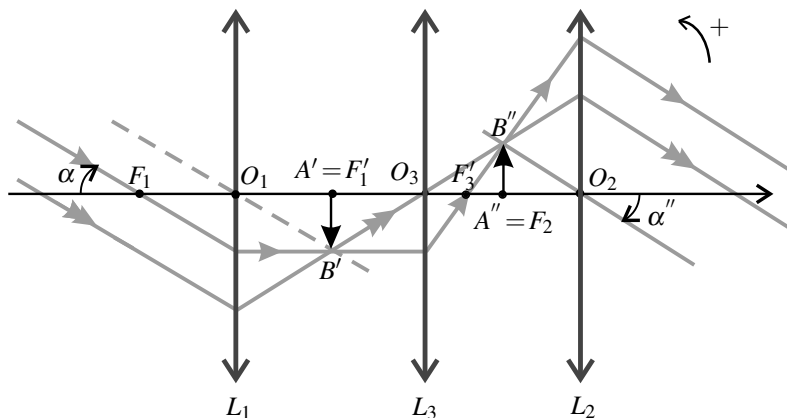
On note que $\alpha < 0$, $\alpha' > 0$ et $\overline{A'B'} < 0$. Dans l'approximation de Gauss, $\tan \alpha \simeq \alpha$ et $\tan \alpha' \simeq \alpha'$, soit

$$G = \frac{\alpha'}{\alpha} = -\frac{f'_1}{f'_2} = -5.$$

3. a. Par la lunette, on a toujours un objet à l'infini dont l'image par l'objectif L_1 est dans le plan focal image de L_1 et une image finale à l'infini, dont l'antécédent par l'oculaire L_2 est dans le plan focal objet de L_2 . Il faut donc que la lentille intermédiaire L_3 conjugue les points F'_1 et F_2 , soit : $F'_1 \xrightarrow{L_3} F_2$.

b. La formule du grandissement avec origine au centre optique de L_3 pour F'_1 et F_2 s'écrit : $\gamma_3 = \frac{\overline{O_3F_2}}{\overline{O_3F'_1}}$ d'où $\overline{O_3F_2} = \gamma_3 \overline{O_3F'_1}$. On reporte dans la relation de conjugaison avec origine au centre pour L_3 : $\frac{1}{\overline{O_3F_2}} - \frac{1}{\overline{O_3F'_1}} = \frac{1}{f'_3}$ soit $\frac{1}{\overline{O_3F'_1}} \left(\frac{1}{\gamma_3} - 1 \right) = \frac{1}{f'_3}$, d'où : $\overline{O_3F'_1} = f'_3 \left(\frac{1 - \gamma_3}{\gamma_3} \right)$.

c. Schéma avec les trois lentilles :



d. En utilisant les notations de la figure, on obtient pour L_1 , $\tan \alpha \simeq \alpha = \frac{\overline{A'B'}}{f'_1}$; puis pour L_3 , $\gamma_3 = \frac{\overline{A''B''}}{\overline{A'B'}}$; puis pour L_2 , $\tan \alpha'' \simeq \alpha'' = -\frac{\overline{A''B''}}{f'_2}$.

On vérifie les signes des expressions : les deux angles sont négatifs ainsi que $\overline{A'B'}$ et γ_3 , les autres grandeurs étant positives. En combinant ces équations, on obtient :

$$G = \frac{\alpha''}{\alpha} = \gamma_3 \left(-\frac{f'_1}{f'_2} \right).$$

Puisque γ_3 est négatif (voir figure), G' est bien positif et l'image est à l'endroit. En ce qui concerne la valeur absolue de G' cela dépend de celle de γ_3 . Si $|\gamma_3| > 1$ alors $G' > |G|$.

5.15 La loupe

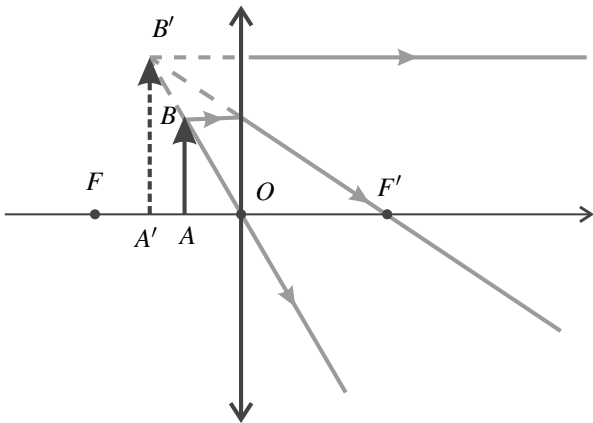
1. À l'œil nu, un objet à une distance d_m et de taille ℓ est vu sous l'angle α_m tel que : $\tan \alpha_m = \frac{\ell}{d_m}$ avec $\tan \alpha_m \simeq \alpha_m$ dans l'approximation de Gauss. Alors : $\alpha_m = \frac{\ell}{d_m}$.

a. La lentille est utilisée en loupe, donc l'objet est placé entre le foyer objet F et le centre optique O . On note $x = \overline{OA}$ et $y = \overline{OA'}$, A' étant l'image de A . La relation de conjugaison de Descartes donne : $\frac{1}{y} - \frac{1}{x} = \frac{1}{f'}$ d'où $y = \frac{xf'}{f' + x}$.

On note O_b la position de l'œil de l'observateur. Pour que l'image soit vue nette il faut que : $O_b A' \geq d_m$, soit en valeur algébrique : $-a + y \leq -d_m$.

Avant de chercher l'intervalle des positions x de l'image, il faut connaître les variations de y en fonction de x . On calcule la dérivée : $\frac{dy}{dx} = \frac{1}{(f' + x)^2} (f'(f' + x) - xf') = \frac{(f')^2}{(f' + x)^2}$. Elle est positive donc y est une fonction croissante de x : à l'objet le plus proche correspond l'image la plus proche, et à l'objet le plus lointain ($A = F$) correspond l'image la plus lointaine (à l'infini).

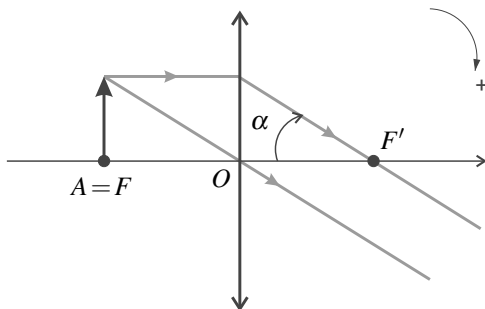
L'objet le plus proche vu net par l'œil est tel que $y = a - d_m$, soit : $\frac{1}{a - d_m} - \frac{1}{x} = \frac{1}{f'}$ d'où $x = \frac{f'(a - d_m)}{f' - a + d_m}$. En conclusion, la zone possible pour les objets donnant une vision nette est l'intervalle $\left[-f' ; \frac{f'(a - d_m)}{f' - a + d_m}\right]$.



CHAPITRE 5 – OPTIQUE GÉOMÉTRIQUE

b. L'œil n'accommode pas si l'image A' est à l'infini, donc l'objet doit être dans le plan focal de la lentille : $A = F$. D'après la figure ci-contre : $\tan \alpha \simeq \alpha = \frac{\ell}{f'}$. On déduit de ce qui précède le grossissement commercial de la loupe :

$$G = \frac{\alpha}{\alpha_m} = \frac{\ell}{f'} \frac{d_m}{\ell} \Rightarrow G = \frac{d_m}{f'} = 5.$$



5.16 Principe du viseur

1. L'oculaire joue le rôle de la loupe étudiée dans le précédent problème et le réticule le rôle de l'objet A , mais $d = -x$ donc d varie dans l'intervalle : $\left[\frac{f'(a - d_m)}{f' - a + d_m} ; f' \right]$.

Pour qu'il n'y ait pas d'accommodation, l'image du réticule doit être à l'infini, donc le réticule doit être placé au foyer objet F_2 de L_2 . Les constructions correspondent aux deux figures de l'exercice précédent.

2. a. D'après l'énoncé $A \xrightarrow{L_1} O$. On sait que si un objet réel a une image réelle par une lentille convergente, la distance entre l'objet et l'image est au minimum égale à 4 fois la distance focale image de la lentille. Cette condition s'écrit ici : $x \leq -4f'_1$.

b. D'après la formule de conjugaison de Descartes $\frac{1}{O_1O} - \frac{1}{O_1A} = \frac{1}{f'_1}$ soit $\frac{1}{D} - \frac{1}{x - D} = \frac{1}{f'_1}$ d'où : $D^2 + xD - xf'_1 = 0$. On choisit pour D la valeur plus petite pour minimiser l'encombrement de l'appareil, soit : $D = -\frac{x}{2} - \frac{1}{2}\sqrt{x^2 + 4xf'_1}$ (remarquer que $x < 0$).

c. Lorsque x varie entre $-\infty$ et $-4f'_1$, D varie entre f'_1 et $2f'_1$. La plage de réglage que le constructeur doit prévoir pour D est : $D \in [f'_1 ; 2f'_1]$.

3. Quand l'observateur change, seul le réglage de l'oculaire est à modifier. On part d'un viseur réglé pour un œil emmétrope à l'infini. Dans ce cas, le réticule est dans le plan focal objet de l'oculaire donc l'image finale est à l'infini. On souhaite augmenter la distance d , de manière à ce que l'image finale soit au *Punctum Remotum* de l'observateur. En appliquant la relation de conjugaison de Descartes : $\frac{1}{O_2A'} - \frac{1}{O_2A} = \frac{1}{f'_2}$. Or $\overline{O_2A'} = a - \delta$, ce qui donne :

$$d = \overline{OO_2} = \frac{(\delta - a)f'_2}{\delta - a + f'_2}$$

Introduction au monde quantique

6

La mécanique quantique a été élaborée dans la première moitié du $XX^{\text{ème}}$ siècle par de nombreux physiciens dont N. Bohr, L. de Broglie, P. Dirac, A. Einstein, W.K. Heisenberg, M. Planck et E. Schrödinger. Cette théorie est nécessaire pour décrire la matière à l'échelle atomique et au-dessous. Ses prédictions n'ont, à l'heure actuelle, jamais été mises en défaut. Elles ont permis la réalisation d'inventions aussi importantes que le laser, le microprocesseur, l'horloge atomique, l'imagerie médicale par résonance magnétique nucléaire...

Les phénomènes quantiques qui apparaissent à l'échelle microscopique sont parfois difficiles à appréhender car il ne correspondent pas à notre intuition naturelle fondée sur notre expérience du monde macroscopique. À la base de leur compréhension se trouve l'idée de dualité onde-particule et la notion d'onde de matière. Il en découle l'inégalité de Heisenberg et la quantification de l'énergie.

1 La dualité onde-particule de la lumière

1.1 Introduction

On a présenté dans les chapitres précédents l'onde lumineuse et les phénomènes d'interférences et de diffraction. Ces phénomènes caractéristiques des ondes ont été étudiés au $XIX^{\text{ème}}$ siècle. Ils montrent l'**aspect ondulatoire** de la lumière.

D'autre part, des phénomènes découverts au $XX^{\text{ème}}$ siècle ne peuvent être interprétés que si l'on admet l'existence de « particules de lumière », que l'on appelle les **photons**. La lumière a aussi un **aspect corpusculaire**.

Ce double comportement est appelé **dualité onde-particule**. Il est impossible à théoriser dans le cadre de la mécanique classique, mais trouve son interprétation dans le cadre la mécanique quantique.

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

1.2 Historique de la découverte du photon

a) Le rayonnement thermique

En 1900, pour réussir à expliquer les propriétés de l'émission thermique du rayonnement électromagnétique d'un corps chauffé (ce rayonnement est pour l'essentiel dans le domaine infrarouge), Max Planck utilisa l'hypothèse que l'énergie s'échange entre la matière et le rayonnement par multiples d'une valeur minimale, le **quantum d'énergie**, dont l'expression est :

$$E = h\nu$$

où ν est la fréquence du rayonnement et h une constante. Cette constante appelée *Hilfskonstante* par Planck, soit en français *constante auxiliaire*, est devenue depuis l'une des constantes fondamentales de la physique. Sa valeur actuellement admise est :

$$h = 6,636\,176.10^{-34} \text{ J.s.}$$

b) L'effet photoélectrique

L'**effet photoélectrique** est l'émission d'électrons par un métal lorsqu'il est éclairé par un rayonnement du domaine visible ou ultraviolet (figure 6.1). Le phénomène n'existe que si la fréquence du rayonnement est supérieure à une fréquence seuil ν_S qui dépend de la nature du métal. Si la fréquence est plus petite que ν_S il n'y a pas d'effet photoélectrique, même si le faisceau est très intense.

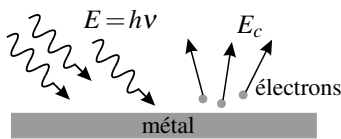


Figure 6.1 – Effet photoélectrique. Chaque électron de masse m_e et vitesse v_e emporte une énergie cinétique : $E_c \leq h\nu - W$.

En 1905, Albert Einstein proposa une interprétation théorique de l'effet photoélectrique en reprenant l'idée de Planck. Il supposa que le rayonnement lui-même est constitué de « quanta de lumière », sortes de grains de lumière contenant l'énergie $E = h\nu$.

L'hypothèse de base de la théorie d'Einstein est qu'un électron du métal peut absorber *un seul* quantum de lumière. Il est alors arraché au métal si l'énergie $E = h\nu$ est supérieure à une valeur minimale dépendant du métal et appelée travail d'extraction W . La condition $E > W$ se traduit par :

$$\nu > \nu_S = \frac{W}{h}.$$

La théorie d'Einstein explique ainsi l'existence de la fréquence seuil. De plus elle prédit que l'énergie cinétique maximale emportée par l'électron est :

$$E_{c,\max} = E - W = h(\nu - \nu_S).$$

Des expériences, menées par Robert Andrews Millikan entre 1905 et 1915, confirmèrent cette formule et donnèrent une valeur de la constante de Planck en bon accord avec la valeur provenant des expériences sur le rayonnement thermique. Einstein reçut le prix Nobel de physique pour son travail sur l'effet photoélectrique en 1922.

c) La diffusion Compton

Le phénomène qui finit de convaincre la communauté scientifique de l'existence de particules associées au rayonnement électromagnétique est la **diffusion Compton** découverte par Arthur Compton en 1922.

La diffusion d'un rayonnement est le phénomène selon lequel un échantillon de matière, recevant un rayonnement incident, renvoie dans tout l'espace un rayonnement de même nature appelé onde diffusée.

Classiquement on explique la diffusion des ondes électromagnétiques par le fait que les électrons contenus dans la matière sont mis en mouvement sous l'action du champ électromagnétique de l'onde incidente. L'onde diffusée est créée par les électrons en mouvement qui oscillent à la fréquence de l'onde incidente (parce que le principe fondamental de la dynamique est une équation linéaire, comme on le verra dans la partie *Mécanique* de ce livre). L'onde diffusée a donc la même fréquence que l'onde incidente (parce que les équations de l'électromagnétisme sont aussi linéaires, voir cours de deuxième année).

En envoyant des rayons X de longueur d'onde $\lambda = 0,071 \text{ nm}$ sur une cible de carbone, A. Compton observa un rayonnement diffusé de longueur d'onde *différente* de la longueur d'onde incidente, en contradiction avec la théorie classique. Il put interpréter les résultats expérimentaux en faisant hypothèse d'une collision entre les électrons contenus dans l'échantillon et des particules arrivant avec le rayonnement incident (voir figure 6.2), particules dotées de l'énergie $E = h\nu = \frac{hc}{\lambda}$ et de la quantité de mouvement :

$$p = \frac{h\nu}{c} = \frac{h}{\lambda}.$$

Au cours de cette collision la particule du rayonnement perd une partie de son énergie, ce qui explique qu'elle reparte avec une énergie $E' < E$, donc une fréquence $\nu' < \nu$ et une longueur d'onde $\lambda' > \lambda$. Le calcul fondé sur les lois de la mécanique relativiste donne des résultats quantitatifs en parfait accord avec l'expérience.

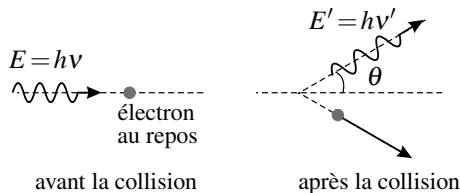


Figure 6.2 – Effet Compton : collision entre un photon et un électron. Le calcul relativiste donne : $\lambda' - \lambda = \frac{h}{m_e c}(1 - \cos \theta)$ où m_e est la masse de l'électron.

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

1.3 Le photon

a) Définition

La particule associée à la lumière s'appelle le **photon**. Ses propriétés sont les suivantes :

- Le photon a une masse nulle.
- Le photon se déplace à la vitesse de la lumière, égale à $c = 3,00 \cdot 10^8 \text{ m} \cdot \text{s}^{-1}$ dans le vide.
- Le photon associé à une lumière de fréquence ν et de longueur d'onde $\lambda = \frac{c}{\nu}$ possède l'énergie :

$$E = h\nu = \frac{hc}{\lambda} \tag{6.1}$$

où h est la constante de Planck.

- Le photon associé à une lumière de fréquence ν se propageant dans la direction du vecteur unitaire $\vec{u}$ possède la quantité de mouvement :

$$\vec{p} = \frac{E}{c} \vec{u} = \frac{h\nu}{c} \vec{u} = \frac{h}{\lambda} \vec{u}. \tag{6.2}$$



Dans le cas du photon, la formule classique de la quantité de mouvement donnée dans la partie *Mécanique*, $\vec{p} = m\vec{v}$, ne s'applique pas. Il faut utiliser les équations relativistes, qui sont en dehors du programme.

b) Ordre de grandeur de l'énergie d'un photon

L'énergie du photon correspondant à une lumière de longueur d'onde $\lambda = 500 \text{ nm}$ est :

$$E = \frac{6,64 \cdot 10^{-34} \times 3,00 \cdot 10^8}{500 \cdot 10^{-9}} = 3,98 \cdot 10^{-19} \text{ J}.$$

On constate que le joule n'est pas une unité adaptée pour exprimer cette énergie. Il est courant d'exprimer l'énergie du photon en électron-volt, une unité d'énergie telle que $1 \text{ eV} = 1,6 \cdot 10^{-19} \text{ J}$. L'énergie du photon précédent est égale à $\frac{3,98 \cdot 10^{-19}}{1,60 \cdot 10^{-19}} = 2,48 \text{ eV}$.

Pour un photon de longueur d'onde quelconque λ , sachant que l'énergie est inversement proportionnelle à la longueur d'onde, on peut écrire : $E = 2,48(\text{eV}) \times \frac{500 \text{ (nm)}}{\lambda \text{ (nm)}}$. Il est ainsi pratique de retenir que l'énergie du photon est donnée par :

$$E \text{ (eV)} = \frac{1240}{\lambda \text{ (nm)}}. \tag{6.3}$$

La notion de photon s'étend aux ondes électromagnétiques autres que la lumière. Le tableau 6.1 donne les gammes d'énergie des photons pour les différents domaines des ondes électromagnétiques. L'aspect corpusculaire ressort d'autant plus que l'énergie du photon est élevée.

On peut donc prévoir qu’il sera très marqué pour les rayonnements γ et X, mais pratiquement imperceptible pour les ondes hertziennes. Inversement l’aspect ondulatoire du rayonnement ressort d’autant plus que la longueur d’onde est élevée.

rayonnement	λ (m)	ν (Hz)	E (eV)
rayons gamma	$< 2.10^{-11}$	$> 1,5.10^{19}$	$> 60\,000$
rayons X	$2.10^{-11} - 1.10^{-8}$	$3.10^{16} - 1,5.10^{19}$	$125 - 60\,000$
ultra-violets	$1.10^{-8} - 4.10^{-7}$	$7,5.10^{14} - 3.10^{16}$	$3 - 125$
visible	$4.10^{-7} - 7,5.10^{-7}$	$4.10^{14} - 7,5.10^{14}$	$1,5 - 3$
infrarouges	$7,5.10^{-7} - 3.10^{-4}$	$1.10^{12} - 4.10^{14}$	$4.10^{-3} - 1,5$
ondes hertziennes	$> 3.10^{-4}$	$< 1.10^{12}$	$< 4.10^{-3}$

Tableau 6.1 – Longueur d’onde, fréquence et énergie des différents types de photons.

c) Nombre de photons

Dans les expériences d’optique, l’aspect granulaire de la lumière n’apparaît généralement pas. L’explication est que ces expériences mettent en jeu un très grand nombre de photons à la seconde.

On peut calculer, par exemple, le nombre de photons émis chaque seconde par un laser hélium-néon de longueur d’onde $\lambda = 633\text{ nm}$ et de puissance $P = 1,0\text{ mW}$. En $\Delta t = 1\text{ s}$, l’énergie lumineuse délivrée par le laser est $E_{\text{laser}} = P\Delta t = 1,0.10^{-3}\text{ J}$. L’énergie d’un photon est $E = \frac{hc}{\lambda} = \frac{6,64.10^{-34} \times 3.10^8}{633.10^{-9}} = 3,1.10^{-19}\text{ J}$. Le nombre de photons est donc :

$$N = \frac{E_{\text{laser}}}{E} \simeq 3,2.10^{15}.$$

C’est un nombre élevé mais, à titre de comparaison, bien plus petit que le nombre d’Avogadro qui est l’ordre de grandeur du nombre d’atomes ou de molécules dans une expérience de chimie.

Pour voir les photons un à un, on peut imaginer réduire la puissance du faisceau. En réduisant d’un facteur 10^9 la puissance, on obtient $3,2.10^6$ photons par seconde, soit un photon toutes les $0,31.10^{-6}\text{ s}$ en moyenne. Les photons se déplaçant à la vitesse $c = 3,0.10^8\text{ m}\cdot\text{s}^{-1}$, deux photons successifs sont alors séparés en moyenne de $3,0.10^8 \times 0,31.10^{-6} = 93\text{ m}$. Avec un détecteur de résolution temporelle de l’ordre de 10^{-7} s , on peut espérer recevoir des photons un à un. En fait ce système simple ne fonctionne pas correctement : à cause du caractère aléatoire de l’émission du laser, les photons sont émis irrégulièrement et plusieurs photons peuvent être détectés en même temps.

On sait cependant produire des photons un à un avec des systèmes beaucoup plus élaborés appelés **sources de photons uniques**. Avec de telles sources on réalise des expériences qui illustrent remarquablement le double aspect, corpusculaire et ondulatoire, de la lumière.

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

1.4 Une expérience avec des photons uniques

a) La lame semi-réfléchissante

Définition La **lame semi-réfléchissante** est un dispositif optique qui divise un faisceau lumineux incident en un faisceau lumineux réfléchi et un faisceau lumineux transmis d'égales puissances.

Interprétation ondulatoire Dans une vision ondulatoire de la lumière, l'onde incidente donne deux ondes : une onde réfléchie et une onde transmise (voir figure 6.3). La puissance transportée par chacune de ces deux ondes est égale à la moitié de la puissance transportée par l'onde incidente :

$$\mathcal{P}_{\text{transmise}} = \mathcal{P}_{\text{réfléchie}} = \frac{1}{2} \mathcal{P}_{\text{incidente}}.$$

Or on a vu que la puissance $\mathcal{P}$ transportée par l'onde lumineuse est proportionnelle au carré de son amplitude A :

$$\mathcal{P} = KA^2$$

où K est une constante. Les amplitudes des ondes transmise et réfléchie, $A_{\text{réfléchie}}$ et $A_{\text{transmise}}$ sont donc reliées à l'amplitude $A_{\text{incidente}}$ de l'onde incidente par :

$$A_{\text{transmise}} = A_{\text{réfléchie}} = \frac{1}{\sqrt{2}} A_{\text{incidente}}.$$

Interprétation corpusculaire Dans l'image où la lumière est constituée photons indivisibles, la division du faisceau incident en deux faisceaux s'interprète ainsi : chaque photon peut être soit transmis, il traverse la lame, soit réfléchi, il « rebondit » sur la lame (voir figure 6.4).

Il y a autant de photons transmis que de photons réfléchis puisque les deux faisceaux ont la même puissance.

On a besoin des deux interprétations, ondulatoire et corpusculaire, pour comprendre les expériences : elles sont complémentaires. La mécanique quantique permet d'unifier les deux points de vue.

b) Une preuve du comportement corpusculaire de la lumière

On ne peut plus considérer à l'heure actuelle que l'effet photoélectrique et l'effet Compton prouvent le caractère corpusculaire de la lumière car il est possible de les interpréter de ma-

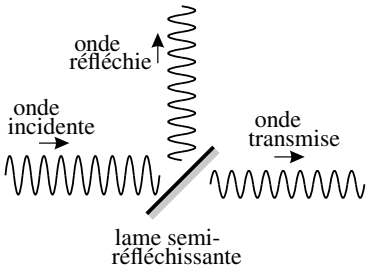


Figure 6.3 – Lame semi-réfléchissante dans une vision ondulatoire.

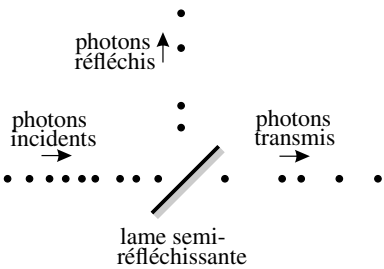


Figure 6.4 – Lame semi-réfléchissante dans une vision corpusculaire.

nière satisfaisante en traitant le rayonnement comme un champ électromagnétique classique et l'électron comme un système quantique.

Une expérience considérée à l'heure actuelle comme une preuve du comportement corpusculaire consiste à envoyer un faisceau de photons uniques sur une lame semi-réfléchissante. Des détecteurs très sensibles recueillent les faisceaux transmis et réfléchis (voir figure 6.5).

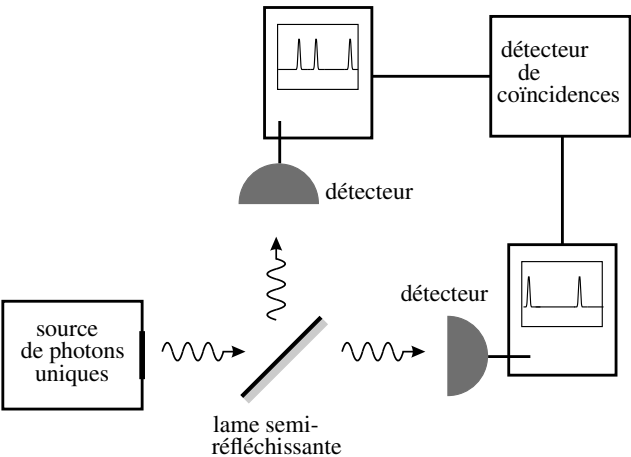


Figure 6.5 – Expérience démontrant le caractère corpusculaire de la lumière.

On constate que :

- les détecteurs fournissent des signaux constitués de pics de durée très brève à des instants aléatoires,
- les pics des deux détecteurs n'arrivent *jamais* simultanément (pas de coïncidence détectée).

Le premier point évoque l'arrivée de photons successifs mais n'est pas incompatible avec une image ondulatoire : en effet, l'onde pourrait être une suite d'impulsions très courtes arrivant aléatoirement. Mais dans ce cas, les pics devraient être détectés simultanément sur les deux détecteurs ce qui n'est pas le cas d'après la deuxième observation. En conclusion, l'expérience montre sans équivoque une manifestation corpusculaire de la lumière.

De plus, d'après la première observation, les photons sont *aléatoirement* réfléchis ou transmis. Quelles sont les probabilités de ces deux événements ? La lame semi-réfléchissante partageant le faisceau incident en deux faisceaux transportant la même puissance, on a :

$$Prob(\text{réflexion}) = \frac{1}{2} \quad \text{et} \quad Prob(\text{transmission}) = \frac{1}{2}.$$

1.5 Franges d'interférences et photons

a) Expérience des fentes de Young

L'expérience des fentes de Young est une expérience classique d'interférences lumineuses qui sera étudiée de manière approfondie dans le cours de deuxième année. On se contentera ici d'en donner le principe.

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

Une source ponctuelle, quasiment monochromatique, éclaire un écran opaque percé de deux fentes rectilignes identiques, très fines, derrière lesquelles on place un écran parallèle au plan des fentes. Sur l'écran on observe des **franges rectilignes** parallèles aux fentes.

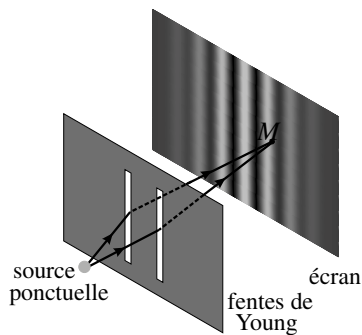


Figure 6.6 – Principe de l'expérience des fentes de Young.

L'explication qualitative du phénomène est la suivante : en tout point M de l'écran parviennent deux ondes qui sont diffractées par les deux fentes de Young. Les longueurs des deux rayons lumineux entre la source et M sont différentes, donc les deux ondes sont décalées dans le temps ce qui se traduit par un déphasage qui dépend de la position de M sur l'écran.

On observe ainsi des franges brillantes, au centre desquelles le déphasage est un multiple de 2π (condition d'interférence constructive) et des franges sombres au centre desquelles le déphasage est un multiple impair de π (condition d'interférence destructive).

b) Construction d'une figure d'interférences photon par photon

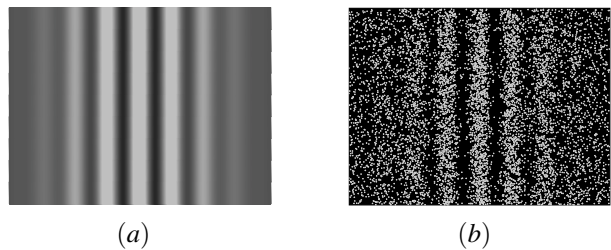


Figure 6.7 – Franges d'interférences : (a) d'après le calcul de l'optique ondulatoire, (b) avec 7500 photons environ (simulations numériques de l'auteur).

Que voit-on si l'on réalise l'expérience des fentes de Young (ou une autre expérience d'interférences lumineuses) avec une source envoyant très peu de photons ? L'expérience peut être réalisée avec une source classique atténuée ou avec une source de photons uniques. On enregistre la figure sur un capteur CCD très sensible. Le temps de pose détermine le nombre

de photons reçus. On peut aussi superposer différents clichés à temps de pose égal pour avoir un nombre croissant de photons. Les observations sont les suivantes :

- Chaque photon donne un point lumineux sur l'image, se comportant donc comme un corpuscule localisé, mais en un point aléatoire.
- Les impacts des photons se répartissent inégalement sur l'écran, dessinant, au fur et à mesure que le nombre de photons augmente, les franges d'interférences prédites par la théorie des interférences de deux ondes lumineuses.

La figure 6.7 montre le résultat d'une simulation de figure d'interférences pour un nombre macroscopique de photons et pour un petit nombre de photons. La simulation est fondée sur l'hypothèse que la probabilité de recevoir un photon en un point est proportionnelle au carré de l'amplitude de l'onde lumineuse (calculée par la théorie ondulatoire) en ce point. L'aspect de la figure simulée est très semblable au résultat de l'expérience réelle. S'il souhaite le vérifier, le lecteur curieux pourra se rendre sur la page http://www.physique.ens-cachan.fr/old/franges_photon/interference.htm et y télécharger le film d'une expérience réelle, montrant une figure d'interférences analogue se construisant photon par photon.

c) Conclusion

La lumière présente un double aspect corpusculaire et ondulatoire, tous les deux perceptibles dans les expériences d'interférences lumineuses avec peu de photons. Les photons sont détectés de manière aléatoire. La probabilité de détecter un photon en un point est proportionnelle au carré de l'amplitude de l'onde lumineuse en ce point.

2 La dualité onde-particule de la matière

2.1 La longueur d'onde de de Broglie

a) L'onde « de matière » de de Broglie

On vient de voir le double aspect ondulatoire et corpusculaire de la lumière et plus généralement de tout rayonnement électromagnétique.

L'idée symétrique que les particules de matière (électrons, protons, neutrons...), par définition de nature corpusculaire, puissent avoir aussi un comportement ondulatoire, a été envisagée pour la première fois en 1923 par Louis de Broglie. Il postula l'existence pour toute particule d'une **onde de matière** qui lui est associée et établit sur des arguments théoriques une expression pour la longueur d'onde λ_{DB} de cette onde :

$$\lambda_{DB} = \frac{h}{p}, \tag{6.4}$$

où p est la quantité de mouvement (ou impulsion) de la particule et h la constante de Planck. La relation fondamentale précédente est appelée **relation de de Broglie** et λ_{DB} est appelée **longueur d'onde de de Broglie** de la particule.

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

Remarque

On a vu plus haut une relation identique à la relation (6.4) pour le photon.

Pour calculer la quantité de mouvement de la particule on peut utiliser l'expression de la mécanique classique $p = mv$ où v est la vitesse de la particule (voir partie *Mécanique* de cet ouvrage) à condition que v soit très inférieure à la vitesse de la lumière $c \simeq 3.10^8 \text{ m}\cdot\text{s}^{-1}$. La relation de de Broglie devient dans ce cas :

$$\lambda_{\text{DB}} = \frac{h}{mv}. \tag{6.5}$$

Dans le cas contraire il faut utiliser la formule relativiste pour la quantité de mouvement. Il est courant d'admettre que la formule classique est valable si $v < \frac{c}{10}$ parce que l'erreur relative introduite par l'usage de la formule classique, égale à $\frac{1}{2} \left(\frac{v}{c} \right)^2$, est inférieure à 1% dans ce cas.

b) Première vérification expérimentale

En 1927, les américains C.J. Davisson et L. Germer apportent la première preuve expérimentale de l'idée de de Broglie, en observant la diffraction d'un faisceau d'électrons par un monocristal de nickel.

La diffraction est un phénomène caractéristique des ondes. Si un faisceau d'électrons est diffracté, c'est la preuve qu'il possède un caractère ondulatoire.

Pourquoi un solide cristallin diffracte-t-il certaines ondes ? Les solides cristallins présentent au niveau atomique un arrangement parfaitement ordonné (voir cours de chimie), périodique dans trois directions de l'espace avec des périodes de l'ordre de $10^{-10} \text{ m} = 0,1 \text{ nm}$. Or les objets périodiques ont la propriété de diffracter d'une manière caractéristique une onde dont la longueur d'onde est proche de leur période. Ainsi, la découverte en 1912 par Max von Laue de la diffraction des rayons X par les cristaux (voir figure 6.8) avait prouvé à la fois la nature ondulatoire de ces rayons et la structure périodique des cristaux. Elle avait aussi permis de mesurer la longueur d'onde des rayons X.



Figure 6.8 – Cliché de diffraction des rayons X par l'austénite.

De la même manière, la diffraction d'un faisceau d'électrons par un cristal de nickel montra que ces particules peuvent avoir un comportement ondulatoire et permit de vérifier quantitativement la formule de de Broglie.

Exemple

Quelle était la longueur d'onde de de Broglie dans l'expérience de Davisson et Germer ?
Les électrons, accélérés par une différence de potentiel de 54 V (voir le chapitre sur les

mouvements de charges dans la partie *Mécanique*), avaient une énergie cinétique

$$E_c = 54 \text{ eV} = 8,6 \cdot 10^{-18} \text{ J}.$$

La masse d'un électron est $m_e = 9,11 \cdot 10^{-31} \text{ kg}$. Si l'on suppose que la mécanique classique s'applique, l'énergie cinétique est donnée par la formule $E_c = \frac{1}{2} m_e v^2$, de sorte que la vitesse d'un électron ayant cette énergie est :

$$v = \sqrt{\frac{2E_c}{m_e}} = \sqrt{2 \times \frac{8,6 \cdot 10^{-18}}{9,11 \cdot 10^{-31}}} = 4,35 \cdot 10^6 \text{ m} \cdot \text{s}^{-1}.$$

Cette vitesse étant inférieure à $\frac{c}{10}$ l'emploi de la formule classique est validé. On peut ensuite calculer :

$$\lambda_{\text{DB}} = \frac{h}{m_e v} = \frac{6,64 \cdot 10^{-34}}{9,11 \cdot 10^{-31} \times 4,35 \cdot 10^6} = 1,68 \cdot 10^{-10} \text{ m}.$$

Cette longueur d'onde est bien de l'ordre de 0,1 nm.

c) Applications actuelles

Microscopie électronique Un microscope optique ne peut révéler des détails plus petits que l'ordre de grandeur de la longueur d'onde de la lumière, soit plus petits qu'un micromètre. La longueur d'onde de de Broglie d'électrons suffisamment accélérés est bien plus courte. Elle est voisine de 0,1 nm pour des électrons d'énergie cinétique égale à 100 eV. Dans certains appareils l'énergie cinétique des électrons atteint 100 keV et la longueur d'onde est de l'ordre de 1 pm = $1 \cdot 10^{-12} \text{ m}$ (pour appliquer la formule de de Broglie dans ce cas il est nécessaire d'utiliser la formule relativiste de la quantité de mouvement). Cependant il faut que l'échantillon supporte une telle énergie.

Exemple

Quelle est la longueur d'onde de de Broglie d'un électron ayant une énergie cinétique $E_c = 100 \text{ keV}$? En appliquant la formule classique de l'énergie cinétique on trouve :

$$v = \sqrt{\frac{2E_c}{m_e}} = \sqrt{2 \times \frac{100 \cdot 10^3 \times 1,60 \cdot 10^{-19}}{9,11 \cdot 10^{-31}}} = 1,9 \cdot 10^8 \text{ m} \cdot \text{s}^{-1} \simeq 0,6c.$$

Cette vitesse est supérieure 0,1c. Les formules classiques ne sont donc pas valides. Il faut donc appliquer les formules relativistes :

$$E_c = (\gamma - 1) m_e c^2 \quad \text{et} \quad p = \gamma m_e v, \quad \text{avec} \quad \gamma = \frac{1}{\sqrt{1 - \frac{v^2}{c^2}}}.$$

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

Ces formules *ne sont pas exigibles* dans le programme des classes préparatoires. En revanche, le candidat aux concours peut avoir à les appliquer si elles lui sont fournies.

On en déduit :

$$\gamma = 1 + \frac{E_c}{m_e c^2} = 1 + \frac{100 \cdot 10^3 \times 1,60 \cdot 10^{-19}}{9,11 \cdot 10^{-31} \times (3,00 \cdot 10^8)^2} \simeq 1,20,$$

puis :

$$v = c \sqrt{1 - \frac{1}{\gamma^2}} = 3,00 \cdot 10^8 \times \sqrt{1 - \frac{1}{(1,20)^2}} \simeq 1,64 \cdot 10^8 \text{ m} \cdot \text{s}^{-1},$$

puis :

$$p = \gamma m_e v = 1,20 \times 9,11 \cdot 10^{-31} \times 1,64 \cdot 10^8 \simeq 1,80 \cdot 10^{-22} \text{ kg} \cdot \text{m} \cdot \text{s}^{-1},$$

et finalement :

$$\lambda_{\text{DB}} = \frac{h}{p} = \frac{6,64 \cdot 10^{-34}}{1,80 \cdot 10^{-22}} = 6,75 \cdot 10^{-12} \text{ m} = 3,7 \text{ pm}.$$

Diffraction des neutrons Les neutrons dits « thermiques », appelés ainsi parce qu'ils sont à une température de l'ordre de la température ambiante (de l'ordre de 300 K) ont une longueur d'onde de de Broglie de l'ordre de la taille d'un atome. Ainsi, ces neutrons sont utilisés comme sonde pour explorer la matière à l'échelle atomique.

Exemple

Quelle est la longueur d'onde de de Broglie de neutrons de masse $m_n = 1,67 \cdot 10^{-27} \text{ kg}$ et de température $T = 300 \text{ K}$?

L'énergie cinétique moyenne d'un neutron sous l'effet de l'agitation thermique (voir la partie *Thermodynamique*) est :

$$E_c = \frac{3}{2} k_B T,$$

où $k_B \simeq 1,38 \cdot 10^{-23} \text{ J} \cdot \text{K}^{-1}$ est la constante de Boltzmann. Or $E_c = \frac{1}{2} m_n v^2$, le neutron a donc une vitesse égale en moyenne à :

$$v = \sqrt{\frac{3 k_B T}{m_n}} = \sqrt{\frac{3 \times 1,38 \cdot 10^{-23}}{1,67 \cdot 10^{-27}}} = 2,73 \cdot 10^3 \text{ m} \cdot \text{s}^{-1}.$$

Cette vitesse très inférieure à c justifie le calcul par une formule classique. La longueur d'onde de de Broglie est donc :

$$\lambda_{\text{DB}} = \frac{h}{m_n v} = \frac{h}{\sqrt{3 m_n k_B T}} \simeq \frac{6,64 \cdot 10^{-34}}{\sqrt{3 \times 1,67 \cdot 10^{-27} \times 1,38 \cdot 10^{-23} \times 300}} \simeq 1,5 \cdot 10^{-10} \text{ m}.$$

2.2 Expériences d'interférences de particules

La théorie de de Broglie est illustrée de manière marquante par des expériences d'interférences de particules réalisées au cours des trois dernières décennies.

a) Interférences d'électrons

La figure 6.9 montre le résultat d'une expérience d'interférences d'électrons réalisée en 1989 par les chercheurs japonais A. Tonomura, J. Endo, T. Matsuda, T. Kawasaki et H. Ezawa. Il s'agit d'une expérience analogue à l'expérience des fentes de Young. Les électrons, après leur passage à travers les « fentes », tombent sur un film fluorescent jouant le rôle d'« écran ». Chaque électron arrivant sur le film provoque l'émission d'environ 500 photons, collectés par un dispositif d'imagerie. Sur le document présenté sur la figure 6.9 chaque point lumineux témoigne de l'arrivée d'un électron.

Dans cette expérience les électrons arrivent un à un sur le détecteur. En effet, le flux d'électrons à travers l'appareil est maintenu à une valeur très faible, de l'ordre de 10^3 électrons par seconde soit un électron toutes les 10^{-3} s. Ces électrons ont une énergie cinétique moyenne de 50 keV, donc une vitesse (relativiste) de $1,25 \cdot 10^8 \text{ m} \cdot \text{s}^{-1}$. Ainsi la distance entre deux électrons consécutifs serait de l'ordre de $1,25 \cdot 10^8 \times 1 \cdot 10^{-3} = 1,25 \cdot 10^5 \text{ m} = 125 \text{ km}$ soit bien plus que la taille de l'appareil. Il y a donc un seul électron à la fois dans l'appareil.

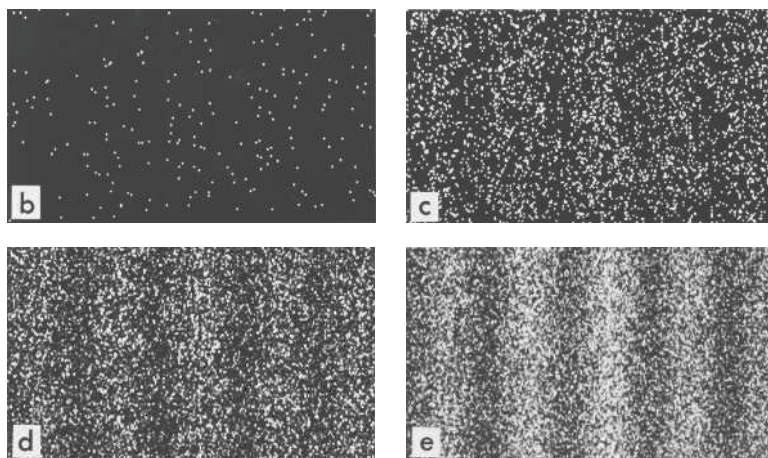


Figure 6.9 – Construction progressive d'un figure d'interférences d'électrons. Chaque point lumineux correspond à l'arrivée d'un électron. Nombres d'électrons respectifs sur les images (b), (c), (d) et (e) : 100, 3000, 20000 et 70000.

L'observation du document de la figure 6.9 appelle les remarques suivantes :

- Les électrons ne se comportent pas comme on s'y attendrait d'après les lois de la mécanique classique. La mécanique classique (c'est-à-dire non quantique) est déterministe : pour des conditions initiales données (position et vitesse initiale de l'électron) on trouve

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

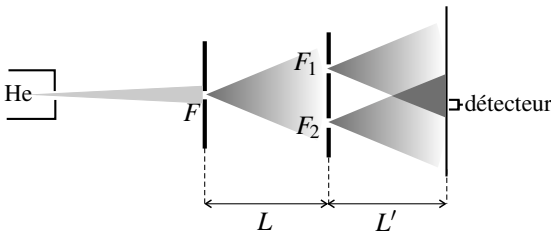


Figure 6.10 – Expérience d'interférences d'atomes d'hélium. O. Carnal, J. Mlynek *Phys. Rev. Lett.*, 1991.

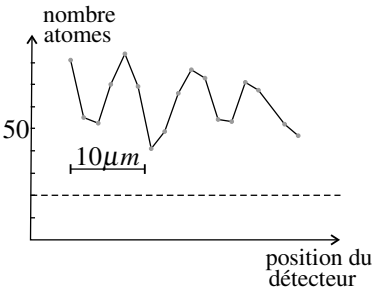


Figure 6.11 – Nombres d'atomes détectés durant 10 minutes en fonction de la position du détecteur.

une trajectoire bien définie¹. Dans ce cas on devrait avoir en théorie *un* point d'impact sur l'écran et en pratique des points regroupés. Au lieu de cela, on observe, notamment sur la photo (b), une répartition aléatoire des points d'impacts. Ainsi : *chaque électron est détecté en un point aléatoire*.

- De la photographie b à la photographie e, au fur et à mesure que le nombre d'électrons détectés croît, apparaît une modulation régulière du nombre d'impacts enregistrés. Des franges d'interférences tout à fait analogues à celles que l'on obtient avec des photons se dessinent peu à peu.

b) Interférences d'atomes

En 1991, O. Carnal et J. Mlynek ont réalisé une véritable expérience de fentes de Young avec des atomes d'hélium. La très grande difficulté de cette expérience tient, entre autres, à la très faible valeur de la longueur d'onde de de Broglie, qui dans ce cas vaut $\lambda_{\text{DB}} = 0,103\text{ nm}$. Les fentes de Young ont une largeur de $s_2 = 1\text{ }\mu\text{m}$ et leurs centres sont séparés de $a = 8\text{ }\mu\text{m}$. Dans le cas d'une expérience d'interférences lumineuses, la longueur d'onde est 10^3 à 10^4 fois plus grande et l'écartement des fentes 10^3 fois plus grand. Les dimensions de l'expérience sont adaptées à la longueur d'onde mise en jeu.

Le schéma de l'expérience est représenté sur la figure 6.10 : le faisceau d'atomes d'hélium est diffracté par la fente F de largeur $s_1 = 2\text{ }\mu\text{m}$, puis par les fentes F_1 et F_2 placées à distance $L = 64\text{ cm}$ de la première fente. On détecte les atomes d'hélium dans le plan situé à distance $L' = 64\text{ cm}$ des fentes, avec un détecteur de $2\text{ }\mu\text{m}$ de largeur, dans la zone de recouvrement des faisceaux issus de F_1 et F_2 .

Les atomes d'hélium ont, comme les électrons de l'expérience précédente, un comportement quantique : un atome donné a une certaine probabilité d'être vu par le détecteur. Le résultat du comptage d'atomes par le détecteur est donc aléatoire et il ne traduit cette loi de probabilité que si un nombre suffisamment grand d'atomes est passé par le dispositif.

Le nombre d'atomes détectés durant 10 minutes est représenté en fonction de la position du

1. Ceci est encore vrai en mécanique relativiste.

détecteur sur la figure 6.11. Le trait en pointillé représente le bruit de fond du détecteur que l'on mesure en occultant le faisceau à l'entrée du dispositif. Les points expérimentaux sont en gris et la ligne noire est seulement un guide pour l'œil. On observe ici aussi une variation alternée du nombre d'atomes détectés pouvant correspondre à des franges d'interférences. L'interfrange, distance entre deux maxima, est estimée expérimentalement à $8,4 \mu\text{m} \pm 0,8 \mu\text{m}$.

Cette interfrange peut être évaluée théoriquement en utilisant la formule 3.3 du chapitre 3 avec la longueur d'onde de de Broglie puisque $L' \gg a$ et $a \gg \lambda_{\text{DB}}$. On trouve :

$$\ell_i = \frac{\lambda_{\text{DB}} L'}{a} = \frac{1,03 \cdot 10^{-10} \times 64 \cdot 10^{-2}}{8 \cdot 10^{-6}} = 8,2 \cdot 10^{-6} \text{ m} = 8,2 \mu\text{m},$$

résultat en bon accord avec la valeur expérimentale.

c) Interférences de molécules

A une température donnée, la vitesse d'agitation thermique d'une particule de masse m est (voir la partie *Thermodynamique*) : $v = \sqrt{\frac{3k_B T}{m}}$ où k_B est la constante de Boltzmann. La

longueur d'onde de de Broglie de cette particule est donc : $\lambda_{\text{DB}} = \frac{h}{mv} = \frac{h}{\sqrt{3mk_B T}}$. Pour une molécule de masse moléculaire M , donc de masse $m = M \times 1,66 \cdot 10^{-27} \text{ kg}$ ($1,66 \cdot 10^{-27} \text{ kg}$ est la masse du proton), à la température $T = 300 \text{ K}$, cette longueur d'onde s'écrit :

$$\lambda_{\text{DB}} = \frac{6,64 \cdot 10^{-34}}{\sqrt{M \times 3 \times 1,66 \cdot 10^{-27} \times 1,38 \cdot 10^{-23} \times 300}} = \frac{1,46 \cdot 10^{-10} \text{ m}}{\sqrt{M}}.$$

Plus la molécule est lourde, plus cette longueur d'onde est petite et donc plus il est difficile de mettre en évidence un comportement ondulatoire.

Des chercheurs de l'université de Vienne ont pu réaliser des interférences avec des molécules de très grande taille : en 1999, le fullerène C_{60} , molécule en forme de ballon de football de masse moléculaire 720, et récemment des molécules dérivées de tétraphénylporphyrine de masse moléculaire dépassant 5000.

3 Fonction d'onde et probabilités

3.1 Analyse d'une expérience d'interférences quantiques

L'expérience de fentes de Young, représentée sur la partie gauche de la figure 6.12, peut être résumée ainsi : une particule (neutron, électron, atome...) peut parvenir en M sur l'écran soit en passant soit par la fente F_1 soit par la fente F_2 . Dans le cas où la particule a un comportement quantique la réponse à la question « la particule est-elle détectée par un détecteur placé au point M ? » est aléatoire. La réponse est « oui » avec une probabilité que l'on note

$$p(M) = \text{Prob}(\text{détecté en } M).$$

La situation peut être considérée comme la superposition des deux situations :

- situation 1 : il n'y a que la fente F_1 , la particule peut parvenir en M en passant par la fente F_1 (situation représentée sur la partie droite de la figure 6.12) ;
- situation 2 : il n'y a que la fente F_2 , la particule peut parvenir en M en passant par la fente F_2 .

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

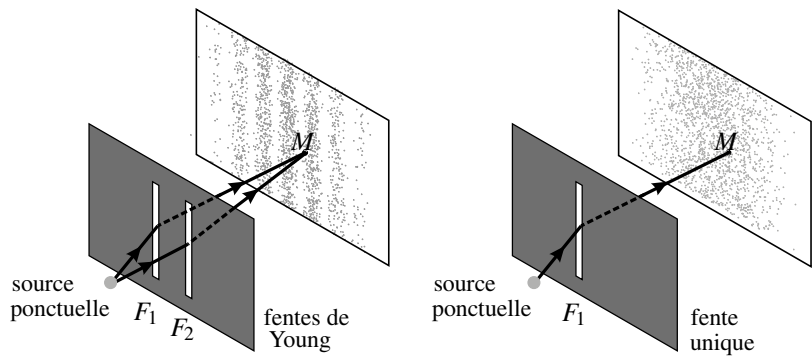


Figure 6.12 – A gauche l’expérience des fentes de Young, à droite la « situation 1 » où la fente F_2 est occultée.

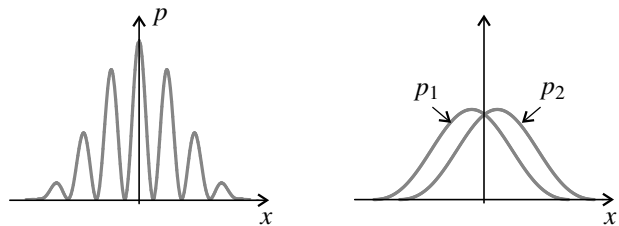


Figure 6.13 – Probabilités de détection de la particule : dans l’expérience d’interférence (à gauche) et dans les « situations 1 et 2 » (à droite).

On appelle $p_1(M)$ et $p_2(M)$ les probabilités de détection de la particule en M dans ces deux situations.

Des expériences comme celles qui ont été décrites ci-dessus permettent, par le comptage des particules arrivant sur le détecteur, de déterminer les probabilités $p(M)$, $p_1(M)$ et $p_2(M)$. Un résultat typique est visible sur la figure 6.12 où chaque point gris sur l’écran représente une particule détectée. Les allures des courbes donnant $p(M)$, $p_1(M)$ et $p_2(M)$ en fonction de la coordonnée x selon un axe perpendiculaire aux fentes sont représentées sur la figure 6.13.

Ces résultats sont surprenants. En effet, en s’appuyant sur l’image classique selon laquelle la particule passe soit par une fente, soit par l’autre, on s’attend à ce que les probabilités $p_1(M)$ et $p_2(M)$ associées aux deux chemins, s’ajoutent lorsque la particule peut emprunter les deux chemins. Or il est manifeste que :

$$p(M) \neq p_1(M) + p_2(M).$$

Plus étonnant encore, il y a des points où $p(M)$ est pratiquement nulle alors que $p_1(M)$ et $p_2(M)$ ne sont pas nulles. En ces points, la particule peut parvenir si l’une des fentes est ouverte, mais ne le peut plus si les deux fentes sont ouvertes.

Ce résultat paradoxal s’interprète par l’*interférence* entre les deux ondes associées à la particule dans les situations 1 et 2.

3.2 Notion de fonction d'onde et probabilité de détection

La grandeur qui s'additionne lorsqu'on superpose les situations 1 et 2 n'est pas la probabilité de détection en M mais la *fonction d'onde* en M . Comme une onde classique, la fonction d'onde est une fonction du point M de l'espace et du temps t . Cependant, elle prend ses valeurs dans l'ensemble des nombres complexes $\mathbb{C}$.

On peut associer à une particule quantique une **fonction d'onde**, $\psi(M,t)$, fonction du point M et du temps, à valeurs complexes. Cette fonction caractérise l'état de la particule.
En particulier la *probabilité* de trouver la particule au point M et à l'instant t est proportionnelle au *module au carré de la fonction d'onde*, $|\psi(M,t)|^2$.

3.3 Interprétation de l'expérience des fentes de Young

Si on note $\psi_1(M,t)$ la fonction d'onde de la particule dans la situation 1 et $\psi_2(M,t)$ la fonction d'onde de la particule dans la situation 2, dans l'expérience des fentes de Young qui est la superposition des situations 1 et 2, la fonction d'onde est :

$$\psi(M,t) = \psi_1(M,t) + \psi_2(M,t).$$

La probabilité de détection de la particule est proportionnelle au module au carré de la fonction d'onde. Or, on peut écrire (en omettant la mention de M et de t pour alléger les écritures) :

$$\begin{aligned} |\psi_1 + \psi_2|^2 &= (\psi_1 + \psi_2)(\psi_1^* + \psi_2^*) \\ &= \psi_1\psi_1^* + \psi_2\psi_2^* + \psi_1\psi_2^* + \psi_1^*\psi_2 \\ &= |\psi_1|^2 + |\psi_2|^2 + \underbrace{\psi_1\psi_2^* + \psi_1^*\psi_2}_{\text{terme d'interférences}}. \end{aligned}$$

Ainsi : $|\psi(M,t)|^2 \neq |\psi_1(M,t)|^2 + |\psi_2(M,t)|^2$ donc $p(M) \neq p_1(M) + p_2(M)$. La différence est due au terme d'interférences.

3.4 Complémentarité

Si l'on considère les particules comme des corpuscules indivisibles, on est amené à se demander à travers laquelle des deux fentes elles passent. Or, dès que l'on tente de répondre à cette question, on provoque inmanquablement la disparition des franges d'interférences.

Une manière d'essayer de déterminer la fente traversée serait d'éclairer les particules passant par une des deux fentes avec un laser. Mais on ne peut résoudre le dilemme suivant : la longueur d'onde du laser doit être suffisamment courte pour différencier les deux fentes et suffisamment grande pour que la quantité de mouvement des photons soit assez faible pour ne pas trop perturber la particule.

Cette impossibilité est fondamentale et provient du **principe de complémentarité** énoncé par Niels Bohr : une particule quantique ne peut se comporter en même temps comme une onde et comme un corpuscule.

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

4 L'inégalité de Heisenberg (PTSI)

4.1 Indétermination quantique

Dans les expériences présentées ci-dessus, on a vu que des particules quantiques arrivent en des points aléatoires du plan de détection. Ainsi, quand on mesure la coordonnée x ou y d'une particule dans le plan de détection, on obtient un résultat aléatoire.

Il s'agit d'une propriété générale des systèmes quantiques :

La mesure d'une grandeur X sur un système quantique donne *a priori* un résultat aléatoire.

Soit $\langle X \rangle$ la moyenne des résultats possibles pondérés par leur probabilité. On définit en mathématiques l'écart quadratique moyen :

$$\Delta X = \sqrt{\langle (X - \langle X \rangle)^2 \rangle} = \sqrt{\langle X^2 \rangle - \langle X \rangle^2}.$$

L'écart quadratique moyen ΔX nous renseigne sur la dispersion des résultats possibles pour la mesure de la grandeur X . On l'appelle **indétermination quantique** de X .

Remarque

Il ne faut pas confondre l'indétermination quantique ΔX avec l'incertitude sur la mesure. ΔX est totalement indépendante de l'instrument de mesure et du protocole suivi. Par exemple sur le document de la figure 6.9, l'incertitude sur la coordonnée x est donnée par la taille d'un point correspondant à un électron, alors que l'indétermination quantique Δx est au moins aussi grande que la largeur du champ visible.

4.2 Exemple : diffraction d'une particule par une fente

On considère une particule de quantité de mouvement $\vec{p} = p\vec{u}_z$ arrivant orthogonalement sur une fente de largeur a parallèle à (Oy) (voir figure 6.14). L'onde associée à la particule est diffractée par la fente exactement comme une onde lumineuse l'est.

Le demi-angle d'ouverture θ du faisceau diffracté est donné par :

$$\sin \theta = \frac{\lambda_{DB}}{a},$$

où λ_{DB} est la longueur d'onde de de Broglie de la particule. Si l'on fait une mesure de la coordonnée x de la particule immédiatement après la fente, le résultat de la mesure est compris entre $-\frac{a}{2}$ et $\frac{a}{2}$ donc l'indétermination quantique est du même ordre de grandeur que a :

$$\Delta x \sim a.$$

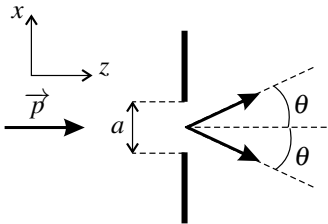


Figure 6.14 – Diffraction d'une particule par une fente.

Si l'on fait la mesure de sa quantité de mouvement p_x selon (Ox) immédiatement après la fente, la dispersion des résultats possibles est due au phénomène de diffraction. Les valeurs extrêmes de p_x correspondent aux cas représentés sur la figure où la particule part dans la direction faisant l'angle θ ou $-\theta$ avec l'axe (Oz) . Les valeurs de p_x possibles sont donc comprises entre $-p \sin \theta$ et $p \sin \theta$. Ainsi, en utilisant $\lambda_{DB} = \frac{h}{p}$ on trouve :

$$\Delta p_x \sim p \sin \theta = p \frac{\lambda_{DB}}{a} = \frac{h}{a}.$$

Il vient donc :

$$\Delta x \Delta p_x \sim h.$$

Ainsi, en localisant la particule avec précision à l'aide d'une fente étroite, on introduit une dispersion sur sa quantité de mouvement.

4.3 L'indétermination position-quantité de mouvement

Le résultat précédent est un cas particulier du **principe d'indétermination de Heisenberg** qui est aussi connu sous le nom, moins approprié, de principe d'incertitude de Heisenberg.

Principe d'indétermination de Heisenberg :

Les mesures de position x et de quantité de mouvement p_x selon un même axe (Ox) sont affectées d'indéterminations quantiques Δx et Δp_x dont le produit est *au moins de l'ordre de grandeur* de la constante de Planck réduite $\hbar = \frac{h}{2\pi}$, ce que l'on écrit :

$$\Delta x \Delta p_x \gtrsim \hbar.$$



La valeur précise de la constante de Planck réduite est $\hbar = 1,056\,180 \cdot 10^{-34} \text{ J}\cdot\text{s}$. Il est facile de retenir que : $\hbar \simeq 1 \cdot 10^{-34} \text{ J}\cdot\text{s}$.

Cette inégalité représente une contrainte fondamentale. Si la position d'une particule est connue avec une grande précision, alors une mesure de la quantité de mouvement sera affectée d'une grande indétermination. C'est par exemple le cas pour une particule enfermée dans une « boîte » de petites dimensions. De même, si l'on a préparé la particule de manière à bien connaître sa vitesse, alors on ne peut avoir que très peu d'information sur sa position.

Ainsi, l'inégalité de Heisenberg rend caduque, à l'échelle microscopique, la notion de trajectoire (pour laquelle position et vitesse sont parfaitement connues) sur laquelle la mécanique classique est fondée. En revanche pour un système macroscopique, la petitesse de la valeur de $\hbar$ fait que la limitation posée par l'inégalité de Heisenberg n'est pas perceptible, donc la mécanique classique s'applique (on verra un exemple d'application numérique au

5 Quantification de l'énergie d'une particule confinée

5.1 Notion de quantification

Un des résultats les plus remarquables de la mécanique quantique est la quantification de certaines grandeurs physiques.

Une grandeur physique est **quantifiée** lorsqu'elle ne peut prendre qu'une suite de valeurs discrètes. On parle de **quantification**.

L'énergie mécanique d'une particule confinée dans une région limitée de l'espace est quantifiée. Le cas le plus simple est celui d'une particule dans un puits infini à une dimension.

5.2 Particule dans un puits infini à 1 dimension

a) Puits infini à 1 dimension

Le **puits infini à une dimension** est le nom donné au modèle théorique d'énergie potentielle suivant :

$$E_p(x) = \begin{cases} \infty & \text{si } x < 0, \\ 0 & \text{si } 0 < x < \ell, \\ \infty & \text{si } x > \ell. \end{cases}$$

Concrètement cela décrit une particule se déplaçant librement entre deux « murs » infranchissables perpendiculaires à (Ox) et distants de ℓ (voir figure 6.15). La particule est confinée dans l'intervalle $0 < x < \ell$.

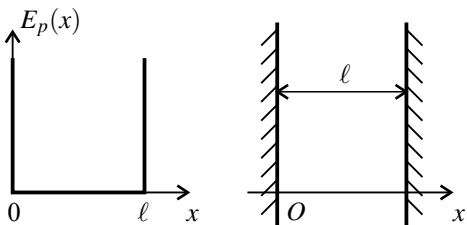


Figure 6.15 – Modèle du puits infini à 1 dimension.

b) Longueurs d'onde possibles pour la particule

Quelles sont les longueurs d'ondes possibles pour une particule placée dans le puits infini ? L'onde de matière est limitée à l'intervalle $0 < x < \ell$. On a vu dans un chapitre précédent qu'une onde dans un espace confiné est nécessairement une onde stationnaire. Il en est obligatoirement de même pour l'onde de de Broglie associée à la particule. *L'onde associée à la particule est une onde stationnaire.*

D'autre part, la probabilité de détecter la particule en un point est proportionnelle au carré de l'amplitude de la fonction d'onde en ce point. On sait que cette probabilité est nulle pour $x < 0$ et pour $x > \ell$. Par continuité, elle est aussi nulle pour $x = 0$ et pour $x = \ell$. Ainsi l'amplitude de l'onde associée à la particule s'annule en $x = 0$ et $x = \ell$. *L'onde associée à la particule présente un nœud de vibration en $x = 0$ et en $x = \ell$.*

L'étude des ondes a montré que la distance entre deux nœuds de vibration consécutifs est égale à la moitié de la longueur d'onde. Ainsi *la largeur du puits est nécessairement un multiple entier de la demi-longueur d'onde*. Ceci s'écrit : $\ell = n \frac{\lambda_{DB}}{2}$, où n est un entier.

Finalement, l'onde associée à une particule placée dans le puits de potentiel a nécessairement une longueur d'onde de de Broglie prenant l'une des valeurs :

$$\lambda_{DB,n} = \frac{2\ell}{n} \text{ où } n \text{ est un entier.} \tag{6.6}$$

c) Niveaux d'énergie

L'énergie mécanique de la particule à l'intérieur du puits, $E = E_c + E_p$, se résume à son énergie cinétique $E_c = \frac{p^2}{2m}$ puisque l'énergie potentielle est nulle. En remplaçant p par $\frac{h}{\lambda_{DB}}$ selon la relation de de Broglie, on trouve :

$$E = \frac{p^2}{2m} = \frac{h^2}{2m\lambda_{DB}^2}.$$

Lorsque la longueur d'onde de de Broglie est égale à $\lambda_{DB,n} = \frac{2\ell}{n}$, alors : $E = \frac{h^2 n^2}{2m(2\ell)^2}$.

Ainsi, l'énergie ne peut prendre que l'une des valeurs suivantes :

$$E_n = n^2 \frac{h^2}{8m\ell^2} \text{ où } n \text{ est un entier.} \tag{6.7}$$

L'énergie de la particule dans le puits est *quantifiée*. Les premiers **niveaux d'énergie** sont représentés sur la figure 6.16. L'écart entre deux niveaux consécutifs :

$$E_{n+1} - E_n = (2n + 1) \frac{h^2}{8m\ell^2}$$

augmente avec n . Sa valeur minimale est $E_1 = \frac{h^2}{8m\ell^2}$.

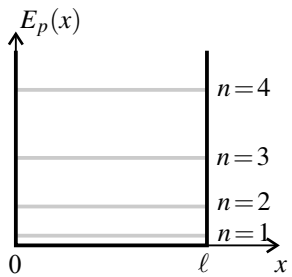


Figure 6.16 – Les 4 premiers niveaux d'énergie de la particule dans le puits infini.

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

Remarque

En appliquant la mécanique classique on trouverait une énergie minimale nulle et l'énergie pourrait prendre toutes les valeurs positives.

d) Transitions entre niveaux d'énergie

La particule doit être dans l'un des niveaux d'énergie précédemment déterminés. Elle peut passer d'un niveau E_n à un niveau plus bas $E_{n'}$ ($n' < n$) en émettant un photon dont la fréquence est donnée par la loi de conservation de l'énergie :

$$h\nu = E_n - E_{n'} = (n^2 - n'^2) \frac{h^2}{8m\ell^2}.$$

Elle peut aussi passer du niveau $E_{n'}$ au niveau E_n en absorbant un photon ayant cette même fréquence.

La plus petite fréquence pouvant être émise ou absorbée correspond à $n = 2$ et $n' = 1$ et son expression est :

$$\nu_0 = \frac{E_2 - E_1}{h} = \frac{3h}{8m\ell^2}.$$

Exemple

On sait réaliser des puits quantiques constitués de couches de semi-conducteurs : par exemple une couche de GaAs en « sandwich » entre deux couches de GaAlAs. On appelle ceci une hétérostructure. Le puits de potentiel correspond à la couche centrale dont la largeur ℓ est de quelques nanomètres. Ces puits quantiques sont à la base des lasers qui équipent la plupart des lecteurs de CD et de DVD de salon.

Dans l'approximation du puits de profondeur infinie, on peut calculer la fréquence ν_0 en utilisant pour les électrons à l'intérieur du puits une masse effective $m_e^* = 0,067m_e$ à la place de leur véritable masse m_e pour tenir compte du fait qu'ils sont à l'intérieur d'un semi-conducteur (la valeur 0,067 correspond au semi-conducteur AsGa). On calcule ainsi pour $\ell = 3$ nm, $\nu_0 = 4,5 \cdot 10^{14}$ Hz qui correspond à une longueur d'onde $\lambda_0 = \frac{c}{\nu} = 662$ nm, soit une lumière visible de couleur rouge.

Remarque

Si ℓ tend vers l'infini, ν_0 tend vers 0 et la différence entre deux niveaux d'énergie consécutifs, $E_{n+1} - E_n$, tend vers 0 aussi. Alors les niveaux sont tellement serrés qu'ils forment quasiment un continuum : l'énergie n'est plus quantifiée. Ainsi la quantification de l'énergie provient du confinement dans une zone limitée de l'espace.

5.3 Généralisation : lien entre confinement spatial et quantification

L'exemple du puits infini montre que la notion d'onde de matière conduit, pour une particule confinée, à la quantification de l'énergie. Ceci est un fait qualitatif général :

Une particule quantique confinée dans une région de l'espace de taille finie a son énergie quantifiée.

En appliquant le modèle du puits infini, fortement idéalisé, on peut chercher une estimation de l'ordre de grandeur de l'écart entre deux niveaux d'énergie, soit $\frac{3h^2}{8m\ell^2}$ pour une particule de masse m confinée dans un espace de largeur ℓ .

Par exemple :

- pour un électron de masse $m_e \sim 10^{-30}\text{kg}$ dans un puits de la taille d'un atome, soit $\ell \sim 10^{-10}\text{m}$, on trouve $\frac{3h^2}{8m_e\ell^2} \sim 10^{-17}\text{J} \sim 100\text{eV}$.
- pour un nucléon (c'est-à-dire un proton ou un neutron) de masse $m \sim 10^{-27}\text{kg}$ dans un puits de la taille d'un noyau atomique, soit $\ell \sim 10^{-15}\text{m}$, on trouve $\frac{3h^2}{8m\ell^2} \sim 10^{-11}\text{J} \sim 100\text{MeV}$.

Ces valeurs sont dans les deux cas supérieures d'un facteur 10 à 100 aux valeurs expérimentales (spectres d'émission en énergie) mais le calcul interprète bien le fait que les énergies mises en jeu dans le noyau sont 10^5 à 10^6 fois plus importantes que celles qui sont mises en jeu dans l'atome.

SYNTHÈSE

SAVOIRS

- relation de Planck
- relation de de Broglie
- notion d'indétermination quantique
- inégalité de Heisenberg (PTSI)

SAVOIR-FAIRE

- évaluer des ordres de grandeur typiques intervenant dans des phénomènes quantiques
- décrire une expérience mettant en évidence la notion de photon
- décrire une expérience illustrant la notion d'onde de matière
- décrire une expérience d'interférences particule par particule
- expliquer qualitativement la nécessité d'une amplitude de probabilité dont le module est proportionnel à la probabilité
- obtenir les niveaux d'énergie du puits quantique infini
- établir le lien qualitatif entre confinement spatial et quantification

MOTS-CLÉS

- | | | |
|---------------------------|-------------------|--------------------|
| – dualité onde-corpuscule | – indétermination | – puits quantique |
| – photon | – probabilité | – niveau d'énergie |
| – onde de matière | – quantification | |

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

S'ENTRAÎNER

6.1 « Photons hertziens » (★)

Un émetteur radio émet un signal de fréquence 105,5 MHz et de puissance 100 kW. Évaluer le nombre de photons qu'il émet par seconde.

6.2 Couleur d'un laser (★)

La lumière d'un faisceau laser est émise par des atomes effectuant une transition entre deux niveaux d'énergie distants de 2,28 eV. Quelle est la couleur de ce laser ?

6.3 Longueur d'onde de de Broglie (★)

1. Calculer la longueur d'onde de de Broglie d'un homme de 75 kg marchant à $5,0 \text{ km} \cdot \text{h}^{-1}$. Comparer à la largeur de la porte de votre chambre et conclure.
2. Quelle énergie, en électronvolts, doit-on communiquer à des électrons, de masse $m_e = 9,11 \cdot 10^{-31} \text{ kg}$, pour que leur longueur d'onde de de Broglie soit égale à $0,1 \text{ nm}$?
3. Calculer les longueurs d'onde de Broglie pour un électron et un proton, de masse $m_p = 1,67 \cdot 10^{-27} \text{ kg}$, dont les énergies cinétiques valent toutes 100 eV.

6.4 Nombres de photons (★)

1. Le flux solaire au niveau du sol terrestre vaut, par beau temps, environ $\Phi_S = 1000 \text{ W} \cdot \text{m}^{-2}$. En prenant pour les photons solaires une longueur d'onde moyenne $\lambda_m = 0,5 \mu\text{m}$, trouver l'ordre de grandeur du nombre de photons reçus par un capteur solaire de surface $S = 1 \text{ m}^2$ pendant $\Delta t = 1 \text{ s}$.
2. Il est possible de voir à l'œil nu une étoile de magnitude 6,5. La magnitude m est reliée au flux d'énergie Φ provenant de l'étoile par la relation $m - m_0 = 2,5 \log_{10} \frac{\Phi_0}{\Phi}$ où m_0 et Φ_0 correspondent à une étoile de référence. La magnitude du Soleil est égale à $-26,8$. Trouver l'ordre de grandeur du nombre de photons provenant d'une étoile de magnitude 6,5 entrant pendant 1 s dans un œil dont la pupille, ouverte au maximum, a un diamètre de 7 mm. On prendra pour longueur d'onde moyenne des photons de l'étoile la même valeur que pour les photons solaires (hypothèse valide si la température de l'étoile est proche de celle du Soleil).

6.5 Longueur d'onde Compton de l'électron (★)

Trouver l'expression de la longueur d'onde λ_C telle que l'énergie d'un photon vaut deux fois l'énergie cinétique d'un électron s'ils ont tous les deux cette longueur d'onde. Calculer numériquement λ_C sachant que la masse de l'électron est $m_e = 9,11 \cdot 10^{-31} \text{ kg}$. À quel domaine du spectre électromagnétique appartient-elle ?

6.6 Vitesse de propagation de l'onde de de Broglie (★)

Une particule de masse m se déplace à la vitesse v très inférieure à la vitesse de la lumière. Elle n'est soumise à aucune force donc son énergie E se réduit à son énergie cinétique. On souhaite trouver la vitesse de propagation de l'onde de de Broglie.

1. Exprimer le vecteur d'onde k_{DB} de l'onde de de Broglie en fonction de m et v .
2. En admettant que la formule reliant l'énergie du photon à sa pulsation est aussi valable pour la particule, trouver une expression ω_{DB} en fonction de m et v .
3. En déduire la vitesse de propagation de l'onde de de Broglie.

6.7 Inégalité de Heisenberg (PTSI) (★)

1. Quelle est l'indétermination quantique minimale sur la vitesse d'un adénovirus dont la masse est égale à $2,4 \cdot 10^{-16}$ g et dont la position est connue à 10 nm près (soit un dixième de sa taille) ?
2. Un radar autoroutier « flashe » une voiture de masse $m = 1,3$ t roulant à une vitesse de $150 \text{ km} \cdot \text{h}^{-1}$. L'éclair du flash dure 0,01 s. Quelle est l'indétermination sur la position de la voiture ? En déduire une minoration de l'indétermination quantique de la vitesse. Conclure.

6.8 Fonction d'onde d'une particule dans un puits infini (★★)

Une particule quantique est confinée dans la zone comprise entre les plans $x = 0$ et $x = \ell$ dans un puits infini. On admet que sa fonction d'onde est de la forme :

$$\psi(x, t) = A \sin(kx) \exp(-i\omega t)$$

où A , k et ω sont des constantes réelles positives.

1. Déterminer les valeurs possibles de k en fonction de ℓ et d'un entier n positif quelconque.
2. La probabilité de trouver la particule dans l'intervalle $[x, x + dx]$ est $|\psi(x, t)|^2 dx$. Justifier la condition de normalisation $\int_0^\ell |\psi(x, t)|^2 dx = 1$. L'utiliser pour trouver l'expression de A en fonction de ℓ .
3. Tracer $|\psi(x, t)|^2$ en fonction de x dans les cas $n = 1$ et $n = 2$. Commenter. Comparer aussi au cas d'une particule classique.

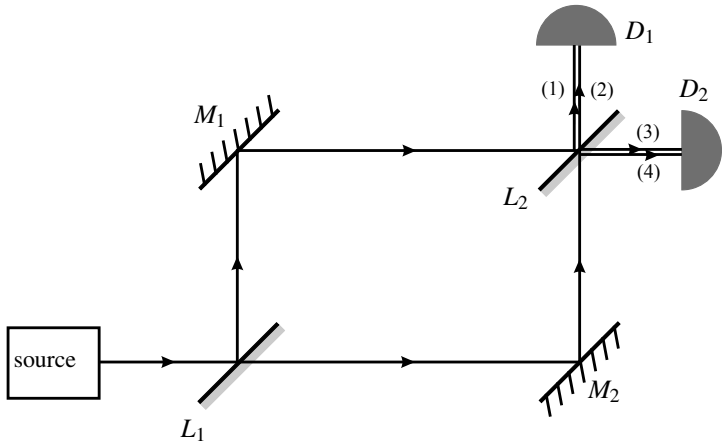
APPROFONDIR

6.9 lame semi-réfléchissante, interféromètre de Mach - Zender (★★)

On considère une lame semi-réfléchissante non équilibrée qui divise un faisceau incident de puissance $\mathcal{P}$ en un faisceau réfléchi de puissance $\mathcal{P}_{\text{réfléchi}} = R\mathcal{P}$ et un faisceau transmis de puissance $\mathcal{P}_{\text{transmis}} = T\mathcal{P}$ avec $R + T = 1$.

1. Traduire ces renseignements en termes de probabilité de réflexion et de transmission du photon.
2. On réalise un interféromètre de Mach-Zender avec deux lames semi-réfléchissantes de ce type. La puissance du faisceau entrant dans l'interféromètre est $\mathcal{P}_0$.

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE



- a. Quelles sont les puissances des quatre faisceaux sortant de l'interféromètre ?
- b. Quelle serait, dans une image corpusculaire, les puissances reçues par les détecteurs D_1 et D_2 ? Expérimentalement la puissance reçue par D_2 est nulle. Expliquer l'échec du modèle corpusculaire à prévoir cette observation.
- c. Dans l'image ondulatoire, en appelant A_0 l'amplitude de l'onde entrant dans l'appareil, écrire les amplitudes des quatre ondes sortant de l'interféromètre.
- d. L'appareil est réglé de manière à ce que les deux trajets de la lumière soient géométriquement rigoureusement identiques. On admet que le faisceau qui se réfléchit sur l'arrière de L_2 subit un déphasage de π , ce qui n'arrive pas au faisceau se réfléchissant sur l'avant de cette lame. Quelles sont les amplitudes des ondes reçues par chacun des détecteurs ? En déduire les puissances qu'ils reçoivent en fonction de la puissance $\mathcal{P}_0$ entrant dans l'appareil. Conclure.
- e. On fait varier le déphasage φ entre les deux ondes arrivant sur le détecteur D_1 . Exprimer les puissances reçues par chacun des détecteurs en fonction de $\mathcal{P}_0$, R , T , φ . Comparer le cas où les lames sont parfaitement équilibrées ($R = T = \frac{1}{2}$) et le cas où elles ne le sont pas.

6.10 Interférences d'atomes d'hélium (Carnal et Mlynek, 1991) (★★)

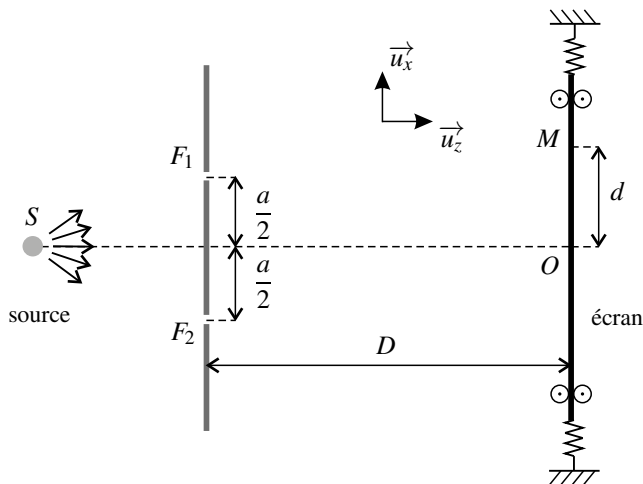
Les informations nécessaires pour répondre aux questions suivantes seront prises dans le paragraphe 2.2.b) de ce chapitre.

- 1. La masse d'un atome d'hélium est $m_{\text{He}} = 6,70 \cdot 10^{-27}$ kg. Quelle est la vitesse des atomes d'hélium de cette expérience ? Estimer le temps mis par un atome pour aller de la fente F au détecteur.
- 2. Calculer l'angle θ de diffraction de l'onde de matière par la fente F . Vérifier que les fentes F_1 et F_2 reçoivent bien cette onde.
- 3. Calculer l'angle θ' de diffraction de l'onde de matière par la fente F_1 ou F_2 . Calculer la largeur de la zone de recouvrement des deux ondes diffractées dans le plan de détection.

4. Combien d'atomes détecte-t-on en moyenne pendant 10 minutes ? Trouver l'ordre de grandeur du nombre d'atomes traversant l'appareil et vus par le détecteur pendant 10 minutes. Conclure en comparant au résultat de la question 1.
5. Il y a des points du plan de détection où la probabilité de détection s'annule par interférence destructrice. Comment se fait-il que, sur la figure 6.11, le nombre d'atomes détectés ne soit jamais nul ?

6.11 Principe d'indétermination et complémentarité (PTSI) (★★)

On réalise l'expérience des fentes de Young avec des photons de longueur d'onde λ que l'on envoie un à un. La distance entre les fentes est a , la distance entre le plan des fentes et l'écran de détection est D .



1. Dans cette question l'écran de détection est fixe. M est un point de l'écran tel que $\overrightarrow{OM} = d\vec{u}_x$.
- a. Décrire ce que l'on observe sur l'écran au fur et à mesure de l'arrivée des photons.
 - b. Les fentes sont à égale distance de la source et on admet que : $F_2M - F_1M \simeq \frac{ad}{D}$.
Exprimer le déphasage en M des ondes passant par les deux fentes.
 - c. Donner l'ordre de grandeur d'une distance caractéristique du phénomène observé sur l'écran de détection.
2. Dans cette question l'écran est monté sur un dispositif mobile de manière qu'il puisse se déplacer en translation dans son plan, selon l'axe (Ox) . On se place dans le modèle corpusculaire. Quand un photon arrive sur l'écran, l'écran l'absorbe et gagne sa quantité de mouvement selon (Ox) . En mesurant la quantité de mouvement de l'écran juste après la détection du photon, on doit pouvoir savoir de quelle fente il provient. On note p la norme de la quantité de mouvement du photon.

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

a. Exprimer p_{1x} , quantité de mouvement selon (Ox) d'un photon parvenant en M et passant par la fente F_1 en fonction de p , d , a et D . Exprimer de même p_{2x} , quantité de mouvement selon (Ox) d'un photon provenant de la fente F_2 .

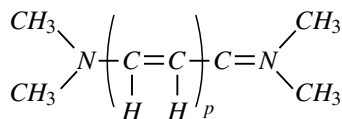
b. En déduire que l'on sait de quelle fente provient le photon seulement si l'indétermination quantique sur la quantité de mouvement de l'écran est très inférieure à une valeur que l'on exprimera en fonction de p , a et D .

c. En se plaçant dans cette hypothèse, comparer l'indétermination quantique sur la position de l'écran et la distance caractéristique définie à la première question. Conclure.

6.12 Molécules de colorant modélisées par un puits quantique (d'après Pour la Science, mai 2011) (★★★)

Les cyanines sont des colorants organiques répandus.

On s'intéresse aux streptocyanines, molécules représentées ci-contre, dans lesquelles le motif du polyacétylène est répété p fois. Ces molécules ont une bande d'absorption dans le visible qui se déplace lorsque p varie, passant du bleu (pour $p = 2$) au rouge (pour $p = 5$). La théorie du puits quantique permet d'interpréter ces propriétés. Il y a en effet sur la molécule $p + 1$ doublets d'électrons « π » délocalisés (les électrons des doubles liaisons) qui sont libres de se déplacer entre les deux atomes d'azote.



1. En admettant que toutes les liaisons $C - C$, ainsi que les deux liaisons $C - N$ concernées ont la même longueur $a = 0,139$ nm, exprimer en fonction de p et a la longueur L sur laquelle les électrons délocalisés sont libres de se déplacer.

2. On suppose que ces électrons sont dans un puits quantique infini de largeur L . Exprimer les énergies E_n possibles pour un électron de masse m_e dans ce puits.

3. Les électrons remplissent ces niveaux d'énergie en respectant le principe de Pauli et la règle de Hund (voir cours de chimie). Montrer que le plus haut niveau d'énergie peuplé est le niveau E_{p+1} .

4. Lorsque la molécule absorbe un photon, un électron passe du niveau d'énergie où il se trouve à un niveau supérieur non occupé. La différence entre les deux niveaux d'énergie est égale à l'énergie du photon. En ne tenant compte que des électrons du puits, montrer que la plus grande longueur d'onde que la molécule absorbe est donnée par la formule :

$$\lambda = \frac{8ma^2c}{h} \frac{(2p+2)^2}{2p+3}.$$

5. Expérimentalement on mesure les longueurs d'onde d'absorption suivantes $p = 2, 3, 4$ et 5 : 416 nm, 519 nm, 625 nm, et 735 nm. Commenter.

6. Quelle est la couleur de la molécule pour $p = 1$?

CORRIGÉS

6.1 « Photons hertziens »

L'énergie d'un photon est : $E_{\text{photon}} = h\nu = 6,63 \cdot 10^{-34} \times 105,5 \cdot 10^6 \simeq 7,0 \cdot 10^{-26}$ J. L'énergie émise par antenne en 1 seconde est : $E_{\text{émise}} = 100 \cdot 10^3 \times 1 = 1,00 \cdot 10^5$ J. Le nombre de photons émis en 1 seconde est donc : $\frac{E_{\text{émise}}}{E_{\text{photon}}} = \frac{1,00 \cdot 10^5}{7,0 \cdot 10^{-26}} = 1,4 \cdot 10^{30}$ soit 2,4 millions de moles de photons. L'aspect corpusculaire ne se manifeste pas dans le cas d'une onde hertzienne.

6.2 Couleur d'un laser

Chaque atome se désexcitant émet un photon d'énergie $E = 2,28 \text{ eV} = 3,65 \cdot 10^{-19}$ J. D'après la formule de Planck : $E = h\nu = \frac{hc}{\lambda}$ donc : $\lambda = \frac{hc}{E} = \frac{6,63 \cdot 10^{-34} \times 3,00 \cdot 10^8}{3,65 \cdot 10^{-19}} = 5,46 \cdot 10^{-7} \text{ m}$ soit $\lambda = 546 \text{ nm}$. C'est une radiation de couleur verte.

6.3 Longueur d'onde de de Broglie

1. D'après la forme de de Broglie : $\lambda_{\text{DB}} = \frac{h}{mv} = \frac{6,63 \cdot 10^{-34}}{75 \times 5,0 \cdot 10^3 / 3600} = 6,4 \cdot 10^{-36} \text{ m}$. La largeur d'une porte de l'ordre de 1 m est bien plus grande. L'homme ne subit pas de diffraction quand il franchit une porte.

2. D'après la formule de de Broglie, la quantité de mouvement de l'électron est $p = \frac{h}{\lambda_{\text{DB}}}$ et son énergie cinétique :

$$E_c = \frac{p^2}{2m_e} = \frac{h^2}{2m_e \lambda_{\text{DB}}^2} = \frac{(6,63 \cdot 10^{-34})^2}{2 \times 9,11 \cdot 10^{-31} \times (0,1 \cdot 10^{-9})^2} = 2,4 \cdot 10^{-17} \text{ J} \simeq 150 \text{ eV}.$$

3. $E_c = \frac{p^2}{2m}$ donc $p = \sqrt{2mE_c}$ et, d'après la relation de de Broglie : $\lambda_{\text{DB}} = \frac{h}{p} = \frac{h}{\sqrt{2mE_c}}$.

Pour l'électron : $\lambda_{\text{DB,électron}} = \frac{6,63 \cdot 10^{-34}}{\sqrt{2 \times 9,11 \cdot 10^{-31} \times 100 \times 1,6 \cdot 10^{-19}}} = 1,24 \cdot 10^{-10} \text{ m} = 0,124 \text{ nm}$.

Et pour le proton :

$$\lambda_{\text{DB,proton}} = \lambda_{\text{DB,électron}} \sqrt{\frac{m_e}{m_p}} = 1,24 \cdot 10^{-10} \times \sqrt{\frac{9,11 \cdot 10^{-31}}{1,67 \cdot 10^{-27}}} = 2,87 \cdot 10^{-12} \text{ m} = 2,89 \text{ pm}.$$

6.4 Nombre de photons

1. L'énergie reçue par le capteur solaire pendant Δt est : $E = \Phi_S S \Delta t$. L'énergie moyenne d'un photon du rayonnement solaire est : $E_{\text{photon}} = \frac{hc}{\lambda_m}$. Le nombre de photons est donc :

$$N_S = \frac{E}{E_{\text{photon}}} = \frac{\Phi_S S \Delta t \lambda_m}{hc} = \frac{1000 \times 1 \times 1 \times 0,5 \cdot 10^{-6}}{6,63 \cdot 10^{-34} \times 3,00 \cdot 10^8} \sim 10^{21}.$$

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

2. Le flux d'énergie Φ provenant de l'étoile de magnitude m est donné par :

$$2,5 \log_{10} \frac{\Phi_S}{\Phi} = m - m_S = 6,5 + 26,8 = 32,3, \quad \text{soit} \quad \Phi = 10^{-\frac{32,3}{2,5}} \Phi_S = 1,2 \cdot 10^{-13} \Phi_S.$$

Le nombre de photons de l'étoile reçus par une surface S pendant Δt est : $N = 1,2 \cdot 10^{-13} N_S$.
Le nombre de photons de l'étoile traversant la pupille de surface S_{pupille} en Δt est :

$$N_{\text{pupille}} = N \frac{S_{\text{pupille}}}{S} = N \frac{\pi \times (7 \cdot 10^{-3})^2 / 4}{1} = 3,8 \cdot 10^{-5} N = 3,8 \cdot 10^{-5} \times 1,2 \cdot 10^{-13} N_S \sim 5 \cdot 10^3.$$

6.5 Longueur d'onde Compton de l'électron

La relation définissant λ_C est : $\frac{hc}{\lambda_C} = 2 \frac{p^2}{2m} = \frac{1}{m_e} \left(\frac{h}{\lambda_C} \right)^2$, en utilisant la relation de de Broglie. On a donc : $\lambda_C = \frac{h}{m_e c} = \frac{6,63 \cdot 10^{-34}}{9,11 \cdot 10^{-31} \times 3,00 \cdot 10^8} = 2,43 \cdot 10^{-12} \text{ m} = 2,43 \text{ pm}$. C'est une longueur d'onde de rayon X.

6.6 Vitesse de propagation de l'onde de de Broglie

1. $k_{\text{DB}} = \frac{2\pi}{\lambda_{\text{DB}}} = \frac{2\pi p}{h} = \frac{2\pi mv}{h}.$

2. L'énergie de la particule est $E = E_c = \frac{1}{2}mv^2$. En admettant qu'on peut appliquer la formule de Planck, $E = h\nu_{\text{DB}} = \frac{h}{2\pi} \omega_{\text{DB}}$, il vient alors : $\omega_{\text{DB}} = \frac{\pi mv^2}{h}.$

3. La vitesse de phase de l'onde de matière est donc : $v = \frac{\omega_{\text{DB}}}{k_{\text{DB}}} = \frac{v}{2}$. Ce n'est pas la vitesse de la particule. En fait c'est une autre vitesse associée à l'onde, la vitesse de groupe qui sera vue en deuxième année, qui est égale à la vitesse de la particule.

6.7 Inégalité de Heisenberg (PTSI)

1. L'indétermination quantique de la position de la particule vérifie : $\Delta x < 10 \text{ nm}$. D'après l'inégalité de Heisenberg, l'indétermination quantique de la quantité de mouvement Δp_x est telle que : $\Delta p_x \gtrsim \frac{\hbar}{\Delta x} > \frac{1 \cdot 10^{-34}}{10 \cdot 10^{-9}} = 1,10^{-25} \text{ kg} \cdot \text{m} \cdot \text{s}^{-1}$. Puisque $p_x = mv_x$, l'indétermination quantique Δv_x sur sa vitesse vérifie : $\Delta v_x = \frac{\Delta p_x}{m} \gtrsim \frac{1 \cdot 10^{-25}}{2,4 \cdot 10^{-19}} = 4,2 \cdot 10^{-7} \text{ m} \cdot \text{s}^{-1}$. Si on veut prendre une image de ce virus au microscope électronique avec une résolution au moins égale à 10 nm il faut que le temps de pose soit au plus de l'ordre de grandeur de $\frac{10 \cdot 10^{-9}}{4,2 \cdot 10^{-7}} = 2,4 \cdot 10^{-2} \text{ s}$.

2. L'indétermination quantique sur la position de la voiture est telle que : $\Delta x < \frac{150 \cdot 10^3}{3600} \times 0,01 = 0,42 \text{ m}$. L'indétermination quantique sur sa quantité de mouvement vérifie donc :

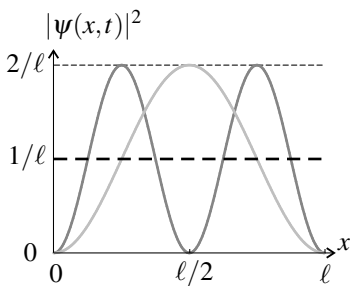
$\Delta p_x \gtrsim \frac{\hbar}{\Delta x} > \frac{1.10^{-34}}{0,42} = 2,4.10^{-34} \text{ kg}\cdot\text{m}\cdot\text{s}^{-1}$. Il en résulte que l'indétermination quantique sur la vitesse vérifié : $\Delta v_x = \frac{\Delta p_x}{m} \gtrsim \frac{2,4.10^{-34}}{1,3.10^3} = 1,8.10^{-37} \text{ m}\cdot\text{s}^{-1}$. L'inégalité d'Heisenberg ne permettra pas à cet automobiliste de contester la contravention !

6.8 Fonction d'onde d'une particule dans un puits infini

1. La fonction d'onde doit s'annuler en $x = 0$ et en $x = \ell$. Il faut et il suffit pour cela que $A \sin(k\ell) = 0$. A ne doit pas être nulle (sinon il n'y a plus de fonction d'onde) donc $\sin(k\ell) = 0$ soit $k\ell = n\pi$ où n est un entier. Les valeurs possibles pour le vecteur d'onde sont : $k = \frac{n\pi}{\ell}$.

2. L'intégrale $\int_0^\ell |\psi(x,t)|^2 dx$ représente la probabilité de trouver la particule quelque part (n'importe où) entre $x = 0$ et $x = \ell$. Comme il est certain que la particule se trouve dans cet intervalle, cette probabilité est égale à 1.

Or : $\psi(x,t) = A \sin\left(\frac{n\pi x}{\ell}\right)$ et $|\psi(x,t)|^2 = A^2 \sin^2\left(\frac{n\pi x}{\ell}\right)$. La condition de normalisation s'écrit : $A^2 \int_0^\ell \sin^2\left(\frac{n\pi x}{\ell}\right) dx = A^2 \frac{\ell}{2} = 1$. On en tire : $A = \sqrt{\frac{2}{\ell}}$.



3. Les courbes sont représentées sur la figure : cas $n = 1$, en gris clair, cas $n = 2$ en gris foncé. On constate que la probabilité de trouver la particule est maximale en $x = \frac{\ell}{2}$ pour $n = 1$. Il en est de même pour toutes les valeurs impaires de n . Pour $n = 2$ et toutes les valeurs paires de n , la probabilité de présence est nulle au centre du puits. Dans le cas classique, la particule fait des allers-retours à vitesse constante d'une extrémité à l'autre du puits, la probabilité de la trouver en un point à un instant quelconque est indépendante de ce point (courbe en tireté sur la figure).

6.9 lame semi-réfléchissante, interféromètre de Mach-Zender

1. Un photon arrivant sur la lame est aléatoirement transmis ou réfléchi. La probabilité qu'il soit réfléchi est R , la probabilité qu'il soit transmis est $1 - R$. Elle ne sont pas égales *a priori*.
2. a. Le faisceau (1) se réfléchit sur L_1 et sur L_2 , sa puissance est : $\mathcal{P}_1 = R^2 \mathcal{P}_0$. Le faisceau (2) est transmis par L_1 et par L_2 , sa puissance est : $\mathcal{P}_2 = T^2 \mathcal{P}_0$. Le faisceau (3) se réfléchit

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

sur L_1 et est transmis par L_2 , sa puissance est : $\mathcal{P}_3 = TR\mathcal{P}_0$. Le faisceau (4) est transmis par L_1 et se réfléchit sur L_2 , sa puissance est : $\mathcal{P}_4 = RT\mathcal{P}_0$.

On a bien : $\mathcal{P}_1 + \mathcal{P}_2 + \mathcal{P}_3 + \mathcal{P}_4 = (R^2 + T^2 + 2RT)\mathcal{P}_0 = (R + T)^2\mathcal{P}_0 = \mathcal{P}_0$.

b. Dans une image corpusculaire, les photons reçus par D_1 sont ceux du faisceau (1) auxquels s'ajoutent ceux du faisceau (2) donc la puissance reçue par D_1 est : $\mathcal{P}_1 + \mathcal{P}_2 = (R^2 + T^2)\mathcal{P}_0$. De même la puissance reçue par D_2 est : $\mathcal{P}_3 + \mathcal{P}_4 = 2RT\mathcal{P}_0$.

Ce n'est pas ce que donne l'expérience parce qu'il y a des interférences.

c. La puissance d'un faisceau est proportionnelle au carré de l'amplitude de l'onde lumineuse. Ainsi on peut déduire de la première question que les amplitudes des faisceaux (1), (2), (3) et (4) sont respectivement : $A_1 = RA_0$, $A_2 = TA_0$ et $A_3 = A_4 = \sqrt{RT}A_0$.

d. Les ondes (1) et (2) arrivant sur D_1 sont en phase (elles ont suivi des chemins de longueurs strictement égales et ont été réfléchies sur les même faces des lames). Elles ont donc une interférence constructive et l'amplitude de l'onde résultante reçue par D_1 est : $A_{D_1} = A_1 + A_2 = (R + T)A_0 = A_0$.

Les ondes (3) et (4) arrivant sur D_2 sont en opposition de phase (elles ont suivi des chemins de longueurs strictement égales mais se sont réfléchies sur des faces différentes des lames). Elles ont donc une interférence destructive et l'amplitude de l'onde résultante reçue par D_1 est : $A_{D_2} = A_3 - A_4 = 0$.

Ainsi le détecteur D_1 reçoit l'intégralité de la puissance $\mathcal{P}_0$ et le détecteur D_2 reçoit une puissance nulle. Ces résultats sont indépendants de R et T .

e. On applique la formule des interférences avec un déphasage φ entre les ondes (1) et (2) : $A_{D_1}^2 = A_1^2 + A_2^2 + 2A_1A_2\cos\varphi = A_0^2(R^2 + T^2 + 2RT\cos\varphi)$ donc la puissance reçue par D_1 est : $\mathcal{P}_{D_1} = \mathcal{P}_0(R^2 + T^2 + 2RT\cos\varphi)$.

On applique la formule des interférences avec un déphasage $\varphi + \pi$ entre les ondes (3) et (4) : $A_{D_2}^2 = A_3^2 + A_4^2 - 2A_3A_4\cos\varphi = 2RTA_0^2(1 - \cos\varphi)$ donc la puissance reçue par D_2 est : $\mathcal{P}_{D_2} = 2RT\mathcal{P}_0(1 - \cos\varphi)$.

Dans le cas où les lames sont parfaitement équilibrées ($R = T = \frac{1}{2}$) on a : $\mathcal{P}_{D_1} = \frac{1}{2}\mathcal{P}_0(1 + \cos\varphi)$

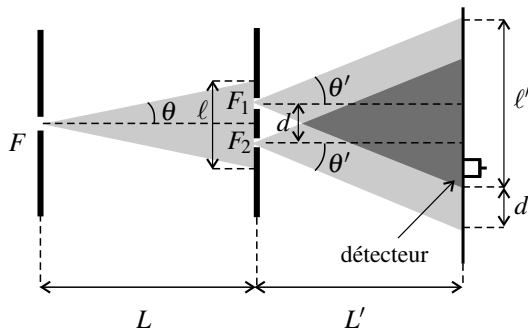
et $\mathcal{P}_{D_2} = \frac{1}{2}\mathcal{P}_0(1 - \cos\varphi)$. Si la lame n'est pas équilibrée :

- l'amplitude d'oscillation des puissances est plus faible car $2RT = 2R(1 - R) < \frac{1}{2}$,
- la puissance arrivant sur D_1 ne s'annule jamais.

6.10 Interférences d'atomes d'hélium (Carnal et Mlynek, 1991)

1. Le texte donne la longueur d'onde des atomes $\lambda_{DB} = 0,103 \text{ nm}$. D'après la relation de de Broglie : $\lambda_{DB} = \frac{h}{p} = \frac{h}{m_{He}v}$, donc : $v = \frac{h}{m\lambda_{DB}} = \frac{6,62 \cdot 10^{-34}}{6,70 \cdot 10^{-27} \times 0,103 \cdot 10^{-9}} = 961 \text{ m} \cdot \text{s}^{-1}$.

La distance de la fente F au détecteur est environ $L + L' = 1,20 \text{ m}$. Le « temps de vol » de l'atome est environ : $\Delta t = \frac{L + L'}{v} = \frac{1,20}{961} \simeq 1,2 \cdot 10^{-3} \text{ s}$.



2. La fente F ayant une largeur $s_1 = 2 \mu\text{m}$, l'onde de matière diffractée a une demi-ouverture angulaire $\theta = \frac{\lambda_{\text{DB}}}{s_1} = \frac{0,103 \cdot 10^{-9}}{2 \cdot 10^{-6}} = 5 \cdot 10^{-5} \text{ rad}$ (voir figure). Sur l'écran percé des fentes F_1 et F_2 cette onde s'étend sur une largeur $\ell = 2L \tan \theta \simeq 2L\theta = 2 \times 0,60 \times 5 \cdot 10^{-5} = 6 \cdot 10^{-5} \text{ m} = 60 \mu\text{m}$. L'écartement des fentes, $d = 8 \cdot 10^{-6} \text{ m}$, est plus petit que ℓ donc les deux fentes reçoivent bien cette onde.

3. Les fentes F_1 et F_2 ont la même largeur $s_2 = 1 \mu\text{m}$ donc $\theta' = \frac{\lambda_{\text{DB}}}{s_2} = \frac{0,103 \cdot 10^{-9}}{1 \cdot 10^{-6}} = 1 \cdot 10^{-4} \text{ rad}$. Chacune des deux ondes diffractées par les fentes donne sur le plan de détection un faisceau de largeur $\ell' = 2L'\theta = 2 \times 0,60 \times 1 \cdot 10^{-4} = 1,2 \cdot 10^{-4} \text{ m} = 120 \mu\text{m}$.

Ces deux faisceaux sont écartés de $d = 8 \mu\text{m}$ donc leur partie commune a une largeur d'environ $\ell' - d = 120 - 8 \simeq 110 \mu\text{m}$.

4. D'après la figure 6.11 du cours, le nombre d'atomes reçus par le détecteur en 10 minutes est en gros de 65, mais il y a un bruit de fond correspondant au comptage de 20 atomes. Ainsi, le détecteur reçoit en 10 minutes en moyenne $65 - 20 = 45$ atomes provenant des fentes.

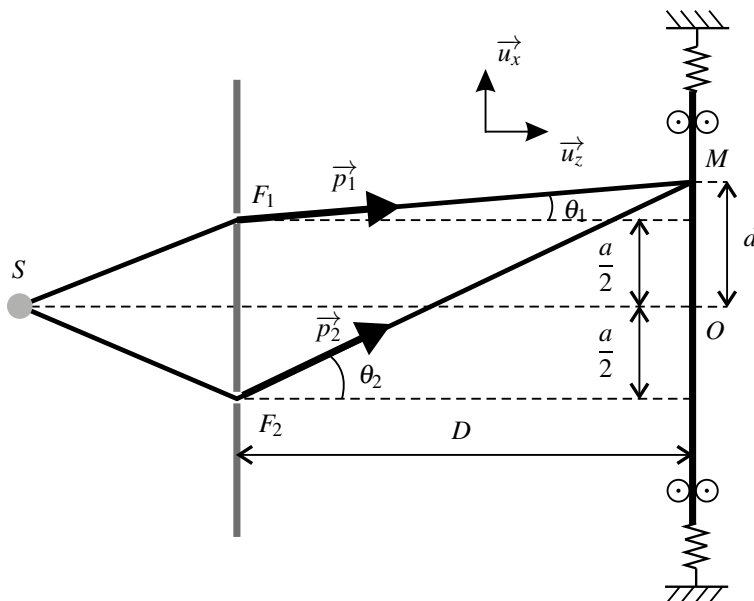
Le détecteur a une largeur de $2 \mu\text{m}$. Sur toute la largeur de la zone d'interférences arrivent donc en moyenne, en 10 minutes : $45 \times \frac{110}{2} \simeq 2500$ atomes.

Le temps moyen entre l'arrivée des atomes est donc $\frac{10 \times 60}{2500} \simeq 0,2 \text{ s}$. Ce temps est bien supérieur au "temps de vol" trouvé à la première question. Il n'y a jamais plus d'un atome détecté à la fois dans l'appareil.

5. Ceci est dû à la largeur du détecteur de $2 \mu\text{m}$ qui n'est pas négligeable devant l'interfrange de $8 \mu\text{m}$.

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

6.11 Principe d'indétermination et complémentarité (PTSI)



1. a. Les photons arrivent en des points aléatoires de l'écran et les points d'impact dessinent peu à peu des franges rectilignes, perpendiculaires aux fentes et régulièrement espacées.

b. La différence des temps de propagation entre S et M est : $\Delta t = \frac{F2M}{c} - \frac{F1M}{c} \simeq \frac{ad}{cD}$. Le déphasage qui lui correspond est : $\Delta\phi = \omega\Delta t = 2\pi\nu\frac{ad}{cD} = 2\pi\frac{ad}{\lambda D}$.

c. Les impacts s'accumulent là où l'amplitude de l'onde est la plus forte, c'est-à-dire là où il y a interférence constructive. La condition d'interférence constructive est $\Delta\phi = 2n\pi$ où n est un entier. Elle s'écrit : $2\pi\frac{ad}{\lambda D} = 2n\pi$ soit $d = n\frac{\lambda D}{a}$. Ainsi, les centres de deux franges sur lesquelles les impacts sont les plus denses (franges brillantes) sont espacés de $i = \frac{\lambda D}{a}$ qui est l'interfrange.

2. a. La quantité de mouvement suivant (Ox) du photon arrivant de la fente F_1 est : $p_{1x} = p \sin \theta_1 \simeq p \theta_1$. En observant la figure ci-dessus on trouve : $\theta_1 \simeq \tan \theta_1 = \frac{d - \frac{a}{2}}{D}$. Finalement : $p_{1x} \simeq p \frac{2d - a}{2D}$. De manière analogue on obtient : $p_{2x} \simeq p \tan \theta_2 = p \frac{2d + a}{2D}$.

b. $p_{2x} - p_{1x} = p \frac{a}{D}$. Pour savoir de quelle fente le photon provient il faut mesurer la quantité de mouvement de la plaque avec une incertitude expérimentale très inférieure à $p \frac{a}{D}$. Or l'incertitude expérimentale est supérieure à l'indétermination quantique, on a donc :

$$\Delta p_x \ll p \frac{a}{D}.$$

c. La plaque est soumise aux lois de la mécanique quantique. D'après l'inégalité de Heisenberg l'indétermination quantique Δx sur sa position vérifie : $\Delta x \gtrsim \frac{\hbar}{\Delta p_x} \gg \frac{\hbar D}{pa} = \frac{\lambda D}{2\pi a}$, soit $\Delta x \gg i$. Une telle indétermination quantique empêche l'observation des franges d'interférences. Le principe de complémentarité est bien respecté : on ne peut à la fois savoir par quel chemin le photon passe et observer des interférences entre les deux chemins.

6.12 Molécules de colorant modélisées par un puits quantique (d'après Pour la Science, mai 2011)

1. Les électrons π peuvent se déplacer le long de $2p$ liaisons CC et de deux liaisons CN donc : $L = 2(p+1)a$.

2. D'après le cours, les niveaux d'énergie possibles pour ces électrons sont : $E_n = n^2 \frac{h^2}{8m_e L^2}$.

3. D'après le principe de Pauli, il y a au plus 2 électrons dans un niveau d'énergie donné. D'après la règle de Hund, la somme des énergies des électrons doit être minimale.

Il y a $p+1$ doubles liaisons donc $2(p+1)$ électrons π , qui occupent donc les niveaux : $E_1, E_2, \dots, E_{p+1}$.

4. La plus grande longueur d'onde absorbée correspond à la plus petite différence d'énergie entre les deux niveaux (puisque $E = \frac{hc}{\lambda}$). Il s'agit du cas où un électron passe du dernier niveau peuplé E_{p+1} au niveau E_{p+2} immédiatement au-dessus. L'énergie du photon absorbé est : $E = E_{p+2} - E_{p+1} = (p+2)^2 \frac{h^2}{8m_e L^2} - (p+1)^2 \frac{h^2}{8m_e L^2} = (2p+3) \frac{h^2}{8m_e L^2} =$

$$\frac{2p+3}{(2p+2)^2} \frac{h^2}{8m_e a^2}, \text{ et sa longueur d'onde : } \lambda = \frac{hc}{E} = \frac{(2p+2)^2}{2p+3} \frac{8m_e a^2 c}{h}.$$

5. En appliquant la formule théorique on trouve pour $p = 2, 3, 4$, et 5 respectivement : $\lambda = 327 \text{ nm}, 452 \text{ nm}, 578 \text{ nm}$ et 705 nm . Ces résultats sont assez différents des valeurs expérimentales ce qui n'est pas étonnant étant donnée la grande simplicité du modèle. Cependant on trouve la bonne évolution.

6. La molécule avec $p = 2$ absorbe dans le violet. Sa couleur est donc la couleur complémentaire : jaune.

CHAPITRE 6 – INTRODUCTION AU MONDE QUANTIQUE

Circuits électriques dans l'ARQS



L'électrocinétique concerne l'étude du mouvement de particules chargées dans la matière. Dans ce chapitre, on définit les notions fondamentales comme le courant et la tension et on précise les lois générales de l'électricité, dans le cadre de l'approximation des régimes quasi-stationnaires.

1 Intensité du courant électrique

1.1 Charge électrique

Certains corps sont susceptibles d'accepter ou de perdre des particules chargées : on dit qu'ils s'électrisent. On peut citer :

- le verre qui, frotté avec de la soie, perd des électrons,
- l'ébonite ou l'ambre qui, frottés avec une fourrure, acquièrent des électrons.

L'électrisation obéit à plusieurs lois qualitatives :

- les corps électrisés exercent des actions mécaniques : ils attirent par exemple des objets légers comme des petits bouts de papier (c'est ce type d'expériences, mettant en évidence l'électrisation des corps, qui a permis la découverte de l'électricité dès l'Antiquité),
- l'électrisation peut se transférer d'un corps à un autre,
- il existe deux types d'électrisation qui seront qualifiés conventionnellement de positive et de négative,
- deux corps de même type d'électrisation se repoussent, tandis que deux corps de type différent d'électrisation s'attirent,
- tout corps électrisé peut attirer un corps non électrisé.

Ces résultats expérimentaux s'interprètent par la notion de charge électrique obtenue grâce aux travaux de Coulomb¹ de 1785 et à la découverte de l'électron en 1881 par Thomson².

1. Charles-Augustin Coulomb, 1736 – 1806, ingénieur français, célèbre pour ses lois du frottement en mécanique et ses expériences sur la force exercée entre charges électriques.

2. Sir Joseph John Thomson, 1856 – 1940, physicien anglais qui découvrit l'électron et inventa la spectrométrie de masse. Il reçut le prix Nobel de Physique en 1906.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

La **charge électrique** est une grandeur scalaire positive ou négative vérifiant les propriétés suivantes.

- Elle peut exister sous deux formes qu'on qualifie de positive et de négative. Le choix de l'électron comme porteur d'une charge négative est purement conventionnel, mais admis de tous. Une charge sera donc positive si elle est attirée par un électron et négative si elle est repoussée par ce dernier. Ceci permet de satisfaire les phénomènes d'attraction et de répulsion observés expérimentalement.
- La charge électrique est une grandeur **additive** dans la même acceptation que celle qui sera adoptée en thermodynamique : la charge ne dépend que de l'état du système et elle est égale à la somme algébrique des charges élémentaires qui la constituent. On peut formuler cette définition en considérant un système formé de l'association de deux sous-systèmes, l'un de charge électrique q_1 , l'autre de charge électrique q_2 . La charge totale q du système est $q = q_1 + q_2$.
- La charge électrique est une grandeur **conservative** au sens où la charge électrique totale d'un système isolé est constante au cours du temps. Les variations de la charge d'un système ne sont donc dues qu'aux échanges avec l'extérieur. Pour un système isolé, il n'y a ni création ni disparition de charges, sa charge électrique ne varie pas.
- La charge électrique se mesure en coulomb, unité de symbole C.
- Les lois de l'électrolyse découvertes par Michael Faraday (1791-1862) s'interprètent par l'existence d'une charge électrique élémentaire notée $e = 1,6 \cdot 10^{-19}$ C. L'ensemble des expériences qui ont été réalisées à ce jour indique que toute charge électrique rencontrée dans la nature est un multiple entier de cette charge, ce qui justifie le fait de parler de charge élémentaire : $q = \pm Ze$ avec $Z \in \mathbb{Z}$. On introduit ainsi la notion de quantification de la charge électrique : celle-ci ne peut prendre qu'un certain nombre de valeurs qui sont les multiples de la charge élémentaire. Cette idée de la quantification de la charge électrique est apparue lors de la découverte de la structure de l'atome. La charge élémentaire joue donc le rôle de quantum pour les charges électriques.
- Le fait d'attribuer un signe + ou un signe - à une charge ou à un porteur de charge est une pure convention mais il est important de s'y conformer afin que tout le monde parle le même langage.

1.2 Porteurs de charges

a) Rappel sur les atomes

La matière est constituée d'atomes qui sont formés :

- d'un noyau comportant :
 - des neutrons, particules non chargées, de masse $m_n = 1,675 \cdot 10^{-27}$ kg,
 - des protons, particules chargées positivement par convention, leur charge étant la valeur de la charge élémentaire $q_p = e = 1,6 \cdot 10^{-19}$ C et de masse $m_p = 1,673 \cdot 10^{-27}$ kg,
- d'un cortège électronique constitué d'électrons qui sont des particules chargées négativement par convention, dont la charge est en valeur absolue la charge élémentaire $q_e = -e = -1,6 \cdot 10^{-19}$ C et de masse $m_e = 9,109 \cdot 10^{-31}$ kg.

Les atomes sont électriquement neutres : leur charge totale est nulle et il y a donc autant de protons que d'électrons dans les atomes.

b) Conduction métallique

Les métaux sont des empilements réguliers d'atomes. La conduction métallique est liée à l'existence d'**électrons libres** ou **électrons de conduction**. En moyenne, un atome du métal conducteur comme le cuivre libère un électron de conduction. Ce dernier peut se déplacer au sein du métal entre les différents atomes : il n'est pas lié à un atome d'où le nom de libre qui lui est donné.

c) Solutions électrolytiques

On s'intéresse en particulier à la conduction dans des solutions, par exemple aqueuses. Les porteurs de charge sont dans ce cas les ions. Les atomes décrits précédemment peuvent céder ou gagner des électrons : on parle d'ionisation de l'atome. Les ions ainsi obtenus ont donc par définition une charge non nulle contrairement aux atomes. On distingue deux types d'ions en fonction du signe de leur charge totale :

- les cations qui ont une charge positive, ce qui correspond à une perte d'électrons, par exemple H_3O^+ ,
- les anions qui ont une charge négative, ce qui correspond à un gain d'électrons, par exemple $\text{S}_4\text{O}_6^{2-}$.

La charge des ions est un multiple de la charge élémentaire suivant le nombre d'électrons qui ont été gagnés ou cédés. Il convient de noter que la matière est globalement neutre conformément au principe de conservation de la charge. Le phénomène d'ionisation correspond à une modification de la répartition des charges mais il n'y a ni apparition ni disparition de charges.

Dans tous les cas :

La charge électrique d'un électron, d'un proton, d'un atome ou d'une molécule est **quantifiée**. Elle est un multiple entier relatif de la charge élémentaire $e = 1,6 \cdot 10^{-19} \text{ C}$.

1.3 Le courant électrique

a) Définition

Le **courant électrique** est un déplacement de charges.

On prend l'exemple d'une solution de sel de cuisine NaCl , dissocié en Na^+ et Cl^- en solution aqueuse. Si on la soumet à un dispositif qui met en mouvement les charges, par exemple un champ électrique entre deux électrodes, comme le montre la figure 7.1, on observe que les charges positives et les charges négatives se déplacent en sens opposé. L'ensemble du mouvement des charges de signes opposés forment le courant électrique. On convient du sens conventionnel suivant :

Le sens du courant est conventionnellement le sens de déplacement des charges positives.

Dans le cas de la conduction métallique (voir figure 7.2), assurée par des électrons, ceux-ci bougent dans le sens opposé à celui du courant.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

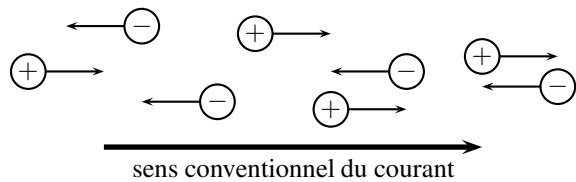


Figure 7.1 – Déplacement de quelques charges et sens conventionnel du courant.

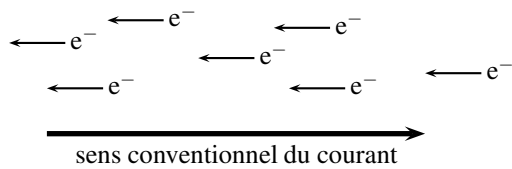


Figure 7.2 – Déplacement des électrons et sens conventionnel du courant.

b) Le courant, grandeur algébrique

Dans un fil électrique, les charges peuvent se déplacer dans les deux sens. On choisit donc, arbitrairement, un sens positif pour le courant. Si celui-ci se déplace dans le sens positif, alors le courant est positif, s'il se déplace dans l'autre sens, alors il est négatif. Ce sens positif arbitraire est signalé par une flèche sur le fil électrique :



Figure 7.3 – Sens arbitraire positif du courant.

1.4 Intensité du courant

On peut développer une analogie hydraulique de l'intensité du courant. Soit un fleuve, de débit D , exprimé en $\text{m}^3 \cdot \text{s}^{-1}$, qui représente le volume qui traverse une section du fleuve par unité de temps. Ce débit est le rapport entre le volume ΔV d'eau qui traverse une section S de référence pendant la durée Δt , et cette durée Δt , comme le montre la figure 7.4 :

$$D = \frac{\Delta V}{\Delta t}.$$

On considère de même un fil conducteur et une section S de référence. La charge qui traverse cette section S pendant la durée Δt est notée Δq . On définit alors l'**intensité** I d'un courant permanent, c'est-à-dire invariable dans le temps, par :

$$I = \frac{\Delta q}{\Delta t}.$$

C'est le débit de charge dans le fil. La formule ci-dessus donne le débit moyen sur la durée Δt . Dans le cas où le courant varie dans le temps, pour avoir l'intensité instantanée i on considère une durée dt infiniment petite et la charge écoulee dq est également infiniment

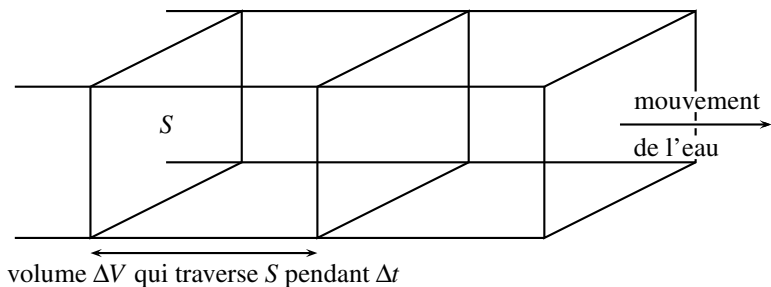


Figure 7.4 – Analogie hydraulique de l'intensité du courant.

petite. L'intensité instantanée se définit alors par :

$$i = \frac{dq}{dt}.$$

1.5 Mesure de l'intensité d'un courant

L'intensité d'un courant se mesure avec un **ampèremètre**. Il ne s'agit pas ici de développer la technologie des ampèremètres, mais plus simplement d'indiquer leur branchement.

L'ampèremètre doit compter la charge qui traverse une section d'un fil en cours du temps. Pour cela, on le branche en série sur le fil dont on souhaite mesurer l'intensité du courant. L'ampèremètre est symbolisé par un $\textcircled{A}$. De la même manière que le débit hydraulique à travers un tuyau, sans fuite, est une constante, l'intensité d'un courant dans un fil unique reste constante, quelque soit la position dans le fil.

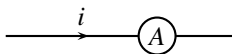


Figure 7.5 – Branchement d'un ampèremètre.

L'intensité du courant se mesure en ampère³, unité de symbole A : 1 A = 1 C.m⁻¹.

Quel est l'ordre de grandeur de l'intensité i du courant ? Elle est très variable. On définit les domaines de :

- l'électronique signal, où les intensités des courants sont de l'ordre de grandeur du mA. Cette électronique est celle des ordinateurs ou des téléphones portables.
- l'électrotechnique, où les courants peuvent atteindre 10³ A. Citons dans ce cas les moteurs électriques des TGV (500 à 10³ A) ou la consommation électrique d'usine.
- les phénomènes naturels, par exemples les éclairs d'orages où l'intensité du courant peut atteindre 50.10³ A, pendant une durée très brève.

3. en l'honneur d'André-Marie Ampère, 1775 – 1836, chimiste, mathématicien et physicien français. Il introduisit l'usage des mathématiques en physique et contribua au développement de l'électromagnétisme. Il fait partie des 72 scientifiques français des XVIII^e et XIX^e siècles dont le nom est inscrit en lettres d'or au premier étage de la tour Eiffel, à Paris.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

1.6 Loi des nœuds

On considère un embranchement électrique, par exemple celui réalisé par une prise multiple, où un fil se sépare en deux. Cet embranchement est nommé **nœud** en électricité.

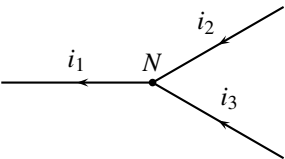


Figure 7.6 – Nœud en électricité.

On oriente les courants des trois fils dans les sens positifs de la figure 1.6. Pendant la durée dt , une charge dq_2 traverse une section du fil 2 et une charge dq_3 une section du fil 3. Alors, les fils électriques ne stockant aucune charge, une charge $dq_2 + dq_3$ doit passer pendant la même durée dans le fil 1, c'est-à-dire :

$$dq_2 + dq_3 = dq_1.$$

En divisant par la durée dt :

$$i_2 + i_3 = i_1.$$

Ce résultat constitue la loi des nœuds. Elle est analogue à la conservation du débit lorsque deux cours d'eau se rejoignent : le débit en aval du confluent est la somme des deux débits en amont.

Un nœud est un embranchement électrique entre plusieurs fils. La **loi des nœuds** stipule qu'en un nœud électrique, la somme des intensités des courants entrant est égale à la somme des intensités des courants sortant du nœud.

Dans un schéma électrique, un nœud est représenté par un point épais.

1.7 Intensité constante, intensité variable

L'intensité d'un courant peut être constante ou dépendre du temps, par exemple avec une loi harmonique $i(t) = I_0 \cos(\omega t)$. Il est d'usage de noter une intensité constante avec une lettre majuscule I , et une intensité variable en fonction du temps avec une lettre minuscule i .

L'usage consacre la lettre majuscule I pour des intensités constantes et la minuscule i pour des intensités variables dans le temps.

2 Tension électrique

2.1 Analogie hydraulique

Le courant électrique dans un fil est analogue à un courant d'eau dans un tuyau. Ce courant d'eau entre deux récipients n'existe spontanément que s'il existe une différence de niveau

entre les deux surfaces. Par exemple, dans la figure 7.7 (à gauche), le récipient A se vide spontanément dans le B ; on observe alors un débit d’eau D dans le tuyau qui les relie.

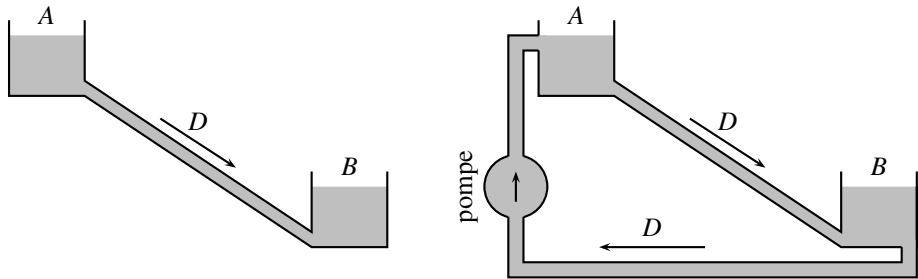


Figure 7.7 – Analogie hydraulique : courant d’eau spontané (à gauche), mouvement entretenu (à droite).

Toutefois, le récipient A finit par se vider complètement. Pour entretenir le mouvement d’eau, il faut fermer le circuit et ajouter une pompe qui assure le débit de fluide. Un générateur hydraulique impose le débit d’eau.

Un circuit électrique fonctionne de la même façon. Un générateur électrique assure le débit de charge, c’est-à-dire la circulation d’un courant électrique, à travers un récepteur électrique, par exemple une lampe, comprise entre les points A et B, comme le montre la figure 7.8.

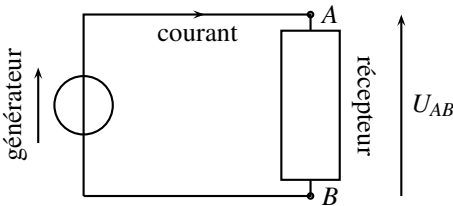


Figure 7.8 – Notion de circuit électrique.

La flèche à côté du générateur indique dans quel sens celui-ci fait circuler le courant électrique. C’est ce qui s’appelle, dans la suite, la convention générateur.

2.2 Notion de tension

Comme le montre la figure 7.8, un courant circule du point A vers le point B, à travers le récepteur. On dit alors qu’il existe une tension électrique entre le point A et le point B, notée U_{AB} . Cette tension est représentée sur le schéma par une flèche qui part de B et pointe vers A. On observe une tension entre deux points d’un circuit si un courant le traverse. Un générateur électrique permet de maintenir des tensions dans le circuit, afin d’imposer la circulation d’un courant électrique.

La **tension** U_{AB} est représentée par une flèche qui part de B et pointe vers A.



Attention à ne pas intervertir la direction de la flèche qui symbolise la tension U_{AB} .

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

2.3 Mesure d'une tension

La tension entre deux points se mesure avec un **voltmètre**⁴. Il ne s'agit pas ici de développer la technologie des voltmètres, mais plus simplement d'indiquer leur branchement.

Le voltmètre mesure la tension entre deux points, il doit donc être relié à ces deux points. De plus, il ne doit pas perturber le circuit, c'est-à-dire ne pas prélever de courant. Par construction, un voltmètre n'est traversé par aucun courant, d'où le 0 sur la figure 7.9 qui explique que le courant qui y passe est bien nul.

Le voltmètre est symbolisé par un $\bigcirc V$. Il mesure la tension U_{AB} sur la figure suivante.

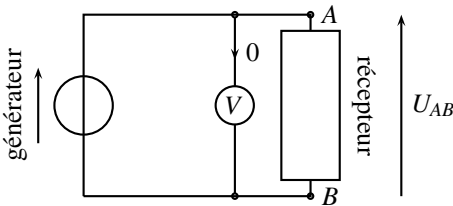


Figure 7.9 – Branchement d'un voltmètre.

Une tension s'exprime en volt, unité de symbole V.

Quel est l'ordre de grandeur d'une tension ? Elle est très variable. On peut citer :

- la tension entre les extrémités d'un éclair lors d'un orage, jusqu'à 500 MV,
- la tension domestique délivrée par EDF, 230 V,
- la tension dans les circuits d'un téléphone portable, quelques mV à quelques centaines de mV.

2.4 Tension constante, tension variable

Une tension peut être constante ou dépendre du temps, par exemple avec une loi harmonique $u(t) = u_0 \cos(\omega t)$. Il est d'usage de noter une tension constante avec une lettre majuscule U , et une tension variable en fonction du temps avec une lettre minuscule u .

L'usage consacre la lettre majuscule U pour des tensions constantes et la minuscule u pour des tensions variables dans le temps.

2.5 Différence de potentiel et tension

Sur la figure 7.10 suivante, la pompe permet de maintenir constantes les hauteurs d'eau z_A et z_B en régime permanent. La pompe impose la différence d'altitude $h_{AB} = z_A - z_B$. De même, le générateur impose la tension U_{AB} . Par analogie, on définit le **potentiel électrique** d'un point, noté V_A pour le point A, V_B pour le point B, par :

$$U_{AB} = V_A - V_B.$$

4. En l'honneur d'Alessandro Volta, 1745 – 1827, physicien italien dont les travaux portèrent sur l'électricité et qui inventa la première pile.

La tension électrique est alors identique à une différence de potentiel. Le potentiel a donc la même unité que la tension, le volt.

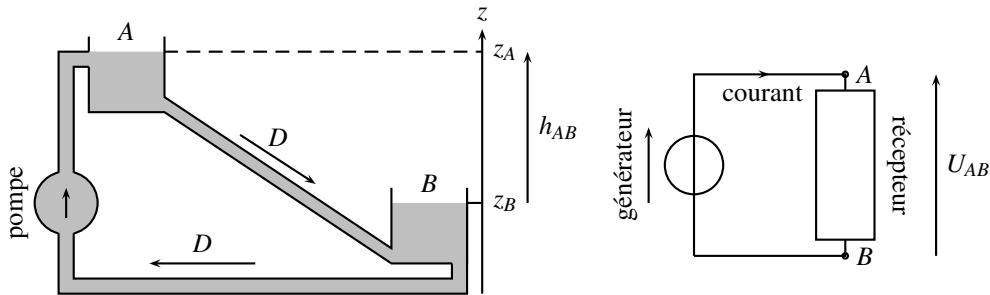


Figure 7.10 – Analogie hydraulique et différence de potentiel.

Une tension U_{AB} est égale à une différence de potentiel : $U_{AB} = V_A - V_B$. Le potentiel V en un point s'exprime en volts.



Attention à ne pas intervertir les indices dans la formule $U_{AB} = V_A - V_B$.

Comme pour l'intensité du courant ou la tension, un potentiel constant est noté avec la majuscule V , alors qu'un potentiel variable dans le temps est noté avec la minuscule v .

2.6 Référence de potentiel

Par définition, une tension est égale à une différence de potentiel. Comme on a :

$$U_{AB} = V_A - V_B = (V_A - C) - (V_B - C),$$

quelque soit la valeur de la constante C , le potentiel est défini à une constante additive près. Ceci permet de décider que le potentiel d'un nœud particulier est nul. Ce nœud est appelé **masse**.

La masse est le point de potentiel nul. Sa position est un choix arbitraire.

Ce choix arbitraire de la masse est noté sur un circuit par des hachures penchées :

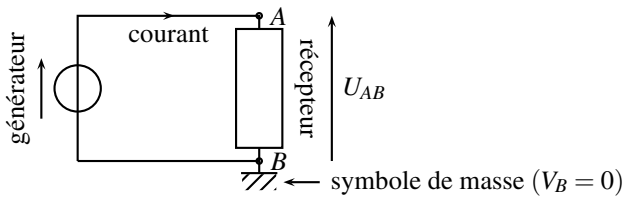


Figure 7.11 – Masse d'un circuit.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

2.7 Additivité des tensions, loi des mailles

Considérons le circuit électrique de la figure suivante. Il est constitué d'une **maille**, c'est-à-dire une boucle sans embranchement parcourue par un courant. Cette maille est constituée d'un générateur et de deux récepteurs. Chaque récepteur est soumis à une tension, ou différence de potentiel, U_{AB} ou U_{BC} .

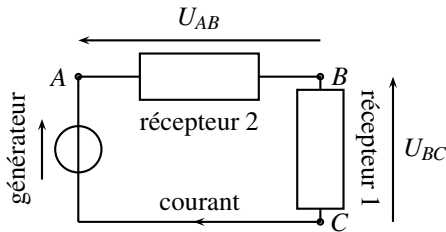


Figure 7.12 – Loi d'additivité des tensions.

Que vaut la tension U_{AC} aux bornes des deux récepteurs pris dans leur ensemble ? La réponse est évidente avec la définition de différence de potentiel :

$$\begin{cases} U_{AB} = V_A - V_B \\ U_{BC} = V_B - V_C \end{cases} \quad \text{donc} \quad U_{AB} + U_{BC} = V_A - V_B + V_B - V_C = V_A - V_C = U_{AC}.$$

C'est la loi d'additivité des tensions qui est analogue à la relation de Chasles des vecteurs.

La loi d'additivité des tensions stipule que : $U_{AC} = U_{AB} + U_{BC}$.

On en déduit immédiatement que $U_{AB} = -U_{BA}$ ainsi que :

$$U_{AD} = U_{AB} + U_{BC} + U_{CD} \quad \text{soit} \quad U_{DA} + U_{AB} + U_{BC} + U_{CD} = 0.$$

La loi des mailles stipule que : $U_{AB} + U_{BC} + U_{CD} + U_{DA} = 0$.

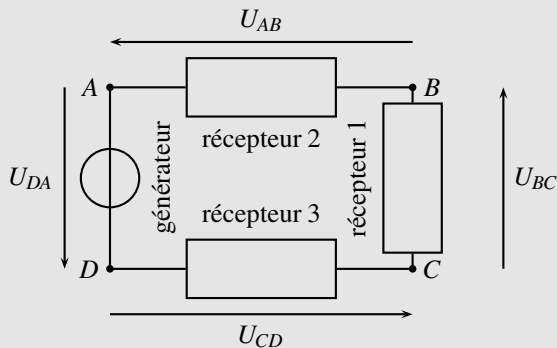


Figure 7.13 – Loi des mailles.

2.8 Notation simplifiée de la tension

Dans la notation U_{AB} , les indices expliquent clairement les deux points entre lesquels on considère la tension. Il souvent d’usage, lorsqu’un schéma électrique est dessiné, d’omettre ces indices. En effet, la représentation de la tension sous forme d’une flèche permet sans ambiguïté de visualiser les deux potentiels qu’on soustrait.

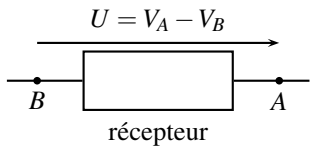


Figure 7.14 – Notation simplifiée, sans indice, d’une tension.



Attention à ne pas intervertir les potentiels. Une tension, symbolisée par une flèche, est le potentiel du bout de la flèche, moins celui de la base de la flèche.

3 Circuit électrique

3.1 Circuit fermé

Un circuit électrique associe toujours un élément actif, ou générateur, qui délivre un certain courant ou une certaine tension, et un élément passif, ou récepteur, qui les utilise. Par exemple pour une lampe, le générateur est l’alimentation EDF et le récepteur la lampe ; pour un téléphone portable, le générateur est l’accumulateur et le récepteur toute l’électronique du téléphone.

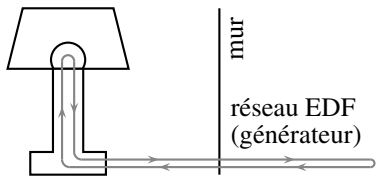


Figure 7.15 – Circuit électrique d’une lampe : en gris, le cheminement du courant électrique.

Un point important doit être souligné :

Un circuit électrique est un circuit **fermé**. Le générateur met en mouvement le courant, qui passe à travers le récepteur pour revenir dans le générateur.

C’est la raison pour laquelle les fils électriques possèdent deux fils côte-à-côte, et que les prises électriques présentent deux fiches. Une fiche sert à conduire le courant vers le récepteur, l’autre à le ramener vers le générateur. Expérimentalement, une fiche EDF, nommée **phase**, est au potentiel de 240 V, l’autre, nommée **neutre**, est au potentiel nul, c’est la masse.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

3.2 Schéma électrique simplifié

Un circuit électrique est électriquement fermé. Ce point est toujours vérifié, à tel point qu'il est possible de ne pas le dessiner sur un schéma électrique, comme le montre la figure suivante :

- la générateur n'est représenté que par la tension qu'il impose,
- les tensions sont mesurée par rapport à une masse arbitrairement choisie,
- le circuit se boucle par un fil non représenté.

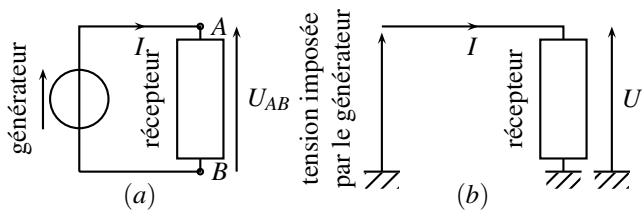


Figure 7.16 – Simplification d'un circuit électrique, (a) schéma réel, (b) schéma simplifié.

3.3 La terre

Certaines prises possèdent, outre la phase et le neutre, une troisième fiche métallique apparente, nommée la **terre**. Elle est directement reliée à un conducteur métallique planté dans le sol. Cette prise de terre est obligatoire dans le cas d'équipement comportant une carcasse métallique, comme une machine à laver par exemple.

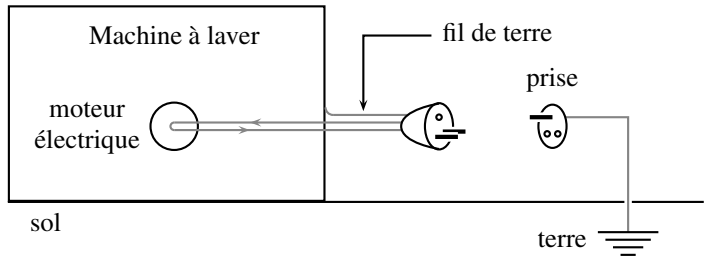


Figure 7.17 – Branchement d'une prise de terre.

La prise de terre est reliée à la carcasse métallique. Sans elle, en cas de faux contact, l'enveloppe de l'appareil peut être portée au potentiel de la phase, et électrocuter un utilisateur qui la toucherait. Avec la prise de terre, la carcasse reste au potentiel de la terre. Les pieds et la main de l'utilisateur qui touche l'appareil, sont au même potentiel ; donc il ne court aucun danger.

3.4 Convention générateur, convention récepteur

Le sens des tensions et des courants sont liés en ce qui concerne les générateurs, qui font circuler le courant, et les récepteurs (appelés parfois charges) qui l'utilisent pour fonctionner.

Dans le circuit suivant, le générateur impose une tension U_{AB} . On observe que cette tension et le courant sont dans la même direction pour le générateur, alors qu'ils sont de sens opposés dans le cas du récepteur. Cette observation correspond aux conventions générateur et récepteur.

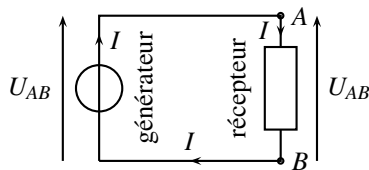
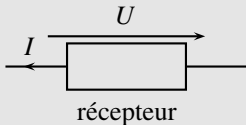
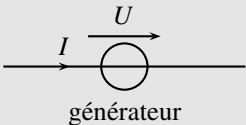


Figure 7.18 – Convention générateur ou récepteur.

La **convention générateur** stipule que pour un générateur, les orientations conventionnellement positives de la tension et du courant sont dans le même sens.



La **convention récepteur** stipule que pour un récepteur, les orientations conventionnellement positives de la tension et du courant sont dans des sens opposés.

3.5 L'approximation des régimes quasi-stationnaires

Lorsqu'on allume un interrupteur, le courant électrique ne traverse pas immédiatement la lampe. Il existe un délai de propagation qui est égal à environ $\frac{L}{c}$, où L est la longueur du fil entre l'interrupteur et la lampe et c la vitesse de la lumière dans le vide. Quel est l'ordre de grandeur ? Avec $L = 10 \text{ m}$, $c = 3.10^8 \text{ m.s}^{-1}$, on trouve $\frac{L}{c} = 3,3.10^{-7} \text{ s}$, ce qui est négligeable par rapport à l'échelle de temps de variation du courant 50 Hz : $T = \frac{1}{f} = 2.10^{-2} \text{ s}$.

Pour une antenne FM, de longueur $L = 1 \text{ m}$, alimentée par un courant de fréquence $f = 100 \text{ MHz}$, $\frac{L}{c} = 3,3.10^{-9} \text{ s}$, qui n'est pas négligeable devant $T = \frac{1}{f} = 10^{-8} \text{ s}$. Dans ce cas, quand l'intensité varie en un bout de l'antenne, elle n'a pas encore varié à l'autre bout. L'**approximation des régimes quasi-stationnaire**, ou **ARQS**, est vérifiée dès que tous les points d'un circuit sont « immédiatement » au courant des variations de la source, c'est-à-dire que la durée de propagation du signal est négligeable devant le temps caractéristique de variation de la source :

$$\Delta t \ll T, \quad \text{donc} \quad L \ll cT.$$

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

Exemple

L'ARQS est-elle vérifiée pour une ligne à haute tension EDF, d'une longueur $L = 300$ km, qui sépare une centrale électrique d'un particulier ? Le temps caractéristique de variation de la tension est l'inverse de la fréquence EDF $f = 50$ Hz :

$$L = 300 \text{ km} \ll cT = \frac{c}{f} = \frac{3 \cdot 10^8}{50} = 6 \cdot 10^3 \text{ km} : \text{ARQS vérifiée.}$$

De même, pour un circuit électronique, de taille caractéristique $L = 10$ cm, qui travaille à 10 MHz :

$$L = 10^{-1} \text{ m} \ll cT = \frac{c}{f} = \frac{3 \cdot 10^8}{10^7} = 30 \text{ m} : \text{ARQS vérifiée.}$$

Dans l'ARQS, toute variation de la source se répercute immédiatement à tous les points du circuit.

Si, par exemple, on débranche la source, le récepteur s'arrête immédiatement de fonctionner. De même, dès qu'on branche un générateur, tous les points d'un même fil sont immédiatement au même potentiel ; c'est la raison pour laquelle, sur la figure 7.18, la tension aux bornes du générateur est la même qu'aux bornes du récepteur.

4 Dipôles

4.1 Définition

En électricité, un dipôle est tout élément qui possède deux bornes. Le courant entre par une borne et sort par l'autre.

4.2 Dipôles passifs

Un **dipôle passif** est un dipôle qui ne peut pas créer lui-même du courant. S'il n'est soumis à aucune tension, aucun courant ne le traverse. Sa loi courant-tension, c'est-à-dire le graphe du courant qui le traverse en fonction de la tension à ses bornes, passe par l'origine ($I = 0, U = 0$).

Il existe de nombreux dipôles passifs. Les trois suivants sont les plus utilisés, car ils permettent de modéliser le comportement des autres dipôles.

a) La résistance

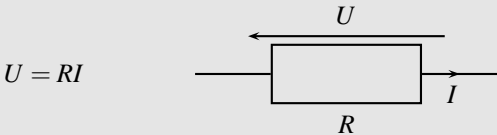
C'est le dipôle passif le plus simple, où la tension appliquée est proportionnelle à l'intensité du courant qui la traverse. La constante de proportionnalité est nommée **résistance** et notée R . La résistance s'exprime en **ohm**⁵, unité dont le symbole est la lettre grecque oméga majuscule Ω . Une résistance est symbolisée par un rectangle dans les schémas électriques.

5. En l'honneur de Georg Ohm, 1789 – 1854, physicien allemand dont les travaux portèrent sur l'électrocinétique et qui formula la loi d'Ohm.

Les valeurs numériques des résistances sont très variables, de quelques Ω à plus de $10^7 \Omega$. Une résistance « moyenne » en électronique est de l'ordre de $1000 \Omega = 1 \text{ k}\Omega$.

La loi de la résistance est $U = RI$, ou $u = Ri$ si le courant et la tension dépendent du temps. Cette loi est nommée **loi d'Ohm**.

La loi d'Ohm est la loi qui relie la tension aux bornes d'une résistance et l'intensité du courant qui la traverse. Avec R la valeur de la résistance du dipôle, qui s'exprime en Ω :



$U = RI$

Figure 7.19 – Résistance en convention récepteur.

b) Le condensateur

Les condensateurs sont des dipôles indispensables dans presque tous les circuits électriques en régime variable, c'est-à-dire pour toutes les applications qui dépendent du temps. Elles sont innombrables, comme dans un téléphone portable ou un ordinateur.

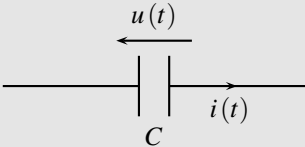
L'intensité du courant qui traverse un condensateur est proportionnelle à la variation de la tension à ses bornes. La constante de proportionnalité est nommée **capacité** et notée C . Ainsi :

$$i = C \frac{du}{dt}.$$

La capacité d'un condensateur s'exprime en **farad**, unité de symbole F. Les valeurs numériques des capacités sont très variables. Le farad est une unité énorme ; les valeurs les plus utilisées vont du pF (picofarad) au mF (millifarad).

Un condensateur est créé grâce à deux surfaces métalliques en regard, éventuellement enroulées. Le schéma électrique d'un condensateur rappelle cette technologie de base.

La tension aux bornes d'un condensateur et l'intensité du courant qui le traverse sont reliés par la loi :



$i = C \frac{du}{dt},$

où C est la capacité du consateur, exprimée en farad F.

Figure 7.20 – Condensateur en convention récepteur.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

c) La bobine

Les bobines sont des dipôles indispensables dans les applications de forte puissance (moteur de TGV par exemple). En traitement du signal, on préfère les simuler avec des circuits électroniques adéquats, beaucoup plus facilement miniaturisables.

La tension aux bornes d'une bobine est proportionnelle à la variation de l'intensité du courant qui la traverse. La constante de proportionnalité est nommée **inductance** et notée L . Ainsi :

$$u = L \frac{di}{dt}.$$

L'unité d'inductance est le **henry**, de symbole H. Les valeurs numériques des inductances varient du mH (millihenry) au henry.

Une bobine est créée grâce à un fil électrique enroulé. Le schéma électrique d'un condensateur rappelle cette technologie de base.

La tension aux bornes d'un condensateur et l'intensité du courant qui le traverse sont reliés par la loi :

$$u = L \frac{di}{dt},$$

où L est l'inductance de la bobine, exprimée en henry H.

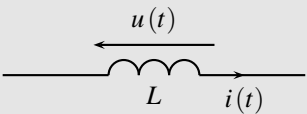


Figure 7.21 – Bobine en convention récepteur.

4.3 Dipôles actifs

Un **dipôle actif** est un dipôle qui permet le passage d'un courant dans le circuit. Il a une fonction génératrice, il génère le déplacement des électrons dans le circuit, on le nomme **générateur**.

Un **générateur idéal de tension** impose une tension U_0 , quel que soit le courant qui le traverse. La tension imposée est nommée tension du générateur ou force électromotrice. Un tel générateur est représenté dans un schéma électrique par un cercle traversé par le fil électrique.

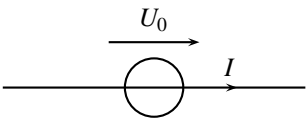


Figure 7.22 – Générateur idéal de tension.

Toutefois, quand on relève expérimentalement la loi courant-tension d'un générateur réel, on trouve la caractéristique de la page suivante. Plus l'intensité du courant débité par le géné-

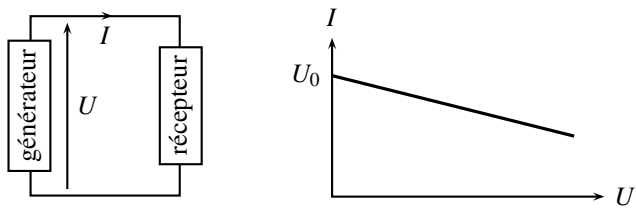


Figure 7.23 – Loi courant-tension d'un générateur réel.

rateur augmente, plus la tension du générateur chute. La loi courant-tension qui modélise ce comportement est alors :

$$U = U_0 - R_g I,$$

où U_0 est la tension à vide, c'est-à-dire quand aucun courant n'est débité, et R_g la résistance interne du générateur.

Le **générateur de tension réel** est représenté par un schéma électrique (nommé modèle de Thévenin) qui associe un générateur idéal de tension et une résistance R_g , appelée résistance interne du générateur :

Figure 7.24 – Schéma électrique de Thévenin d'un générateur réel.

5 Associations de dipôles

5.1 Association de résistances

Il existe deux types d'association de résistances : en série et en parallèle. Deux résistances en série sont parcourues par la même intensité ; deux résistances en parallèle sont soumises à la même tension.

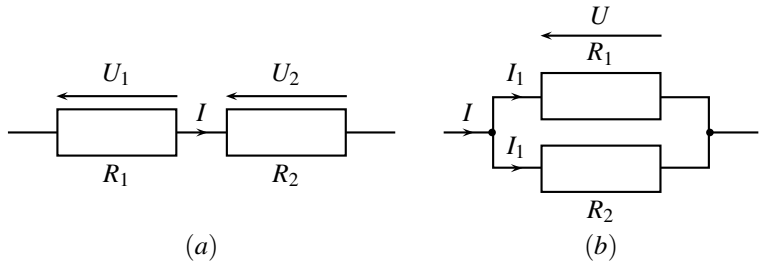


Figure 7.25 – Association de résistances (a) en série, (b) en parallèle.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

a) En série

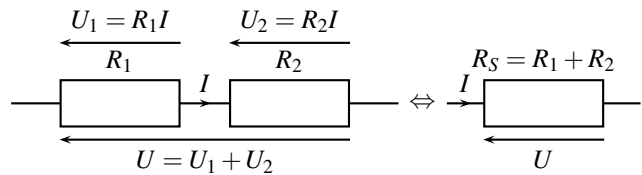


Figure 7.26 – Association de résistances en série.

La tension aux bornes des deux résistances est :

$$U = U_1 + U_2 = R_1 I + R_2 I = (R_1 + R_2) I = R_S I.$$

Les deux résistances sont équivalentes à une seule résistance, de valeur $R_S = R_1 + R_2$. Ce résultat est extrapolable à plus de deux résistances.

Plusieurs résistances R_i , montées **en série**, sont équivalentes à une unique résistance R_S , dont la valeur est la somme des résistances individuelles : $R_S = \sum_i R_i$.

b) En parallèle

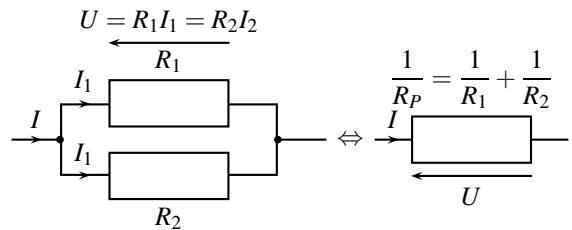


Figure 7.27 – Association de résistances en parallèle.

La loi des nœuds impose :

$$I = I_1 + I_2 = \frac{U}{R_1} + \frac{U}{R_2} = U \left(\frac{1}{R_1} + \frac{1}{R_2} \right) = \frac{U}{R_P}.$$

Les deux résistances sont équivalentes à une seule résistance R_P telle que : $\frac{1}{R_P} = \frac{1}{R_1} + \frac{1}{R_2}$. Ce résultat est extrapolable à plus de deux résistances.

Plusieurs résistances R_i montées **en parallèle**, sont équivalentes à une unique résistance R_P , telle que : $\frac{1}{R_P} = \sum_i \frac{1}{R_i}$.

5.2 Associations de générateurs

L'association de générateurs de tension est possible, sous certaines conditions, car il existe une configuration dangereuse et donc interdite.

L'association en série mène immédiatement à un générateur équivalent unique qui impose la tension $U_1 + U_2$ (voir figure 7.28). C'est ce que l'on fait lorsqu'on met des piles tête-bêche dans un appareil électrique.

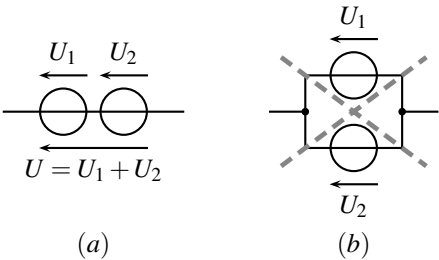


Figure 7.28 – Association de générateurs (a) permise en série, (b) interdite en parallèle.

Dans l'association en parallèle, que vaudrait la tension aux bornes de l'ensemble des deux générateurs ? U_1 ou U_2 ? Le résultat final ne serait ni l'un, ni l'autre, mais plus sûrement l'endommagement des deux générateurs. C'est pourquoi cette configuration est interdite.

5.3 Diviseur de tension

Un **diviseur de tension** est un circuit électrique qui permet, au prix d'une perte de puissance (voir paragraphe 8), de diminuer la valeur d'une tension.

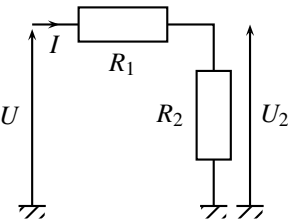


Figure 7.29 – Diviseur de tension.

La loi des mailles impose :

$$\begin{cases} U_2 = R_2 I \\ U = R_2 I + R_1 I \end{cases} \quad \text{donc}$$

$$U_2 = \frac{R_2}{R_1 + R_2} U.$$

La tension E est multipliée par un terme inférieur à un, d'où le nom de diviseur de tension.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS



Toute la démonstration repose sur le fait que le même courant, d'intensité I , circule dans les deux résistances, et que les résistances sont montées en série. Si tel n'est pas le cas, la formule du diviseur de tension est fausse.

La tension est récupérée aux bornes de la résistance R_2 . Ainsi, la puissance Joule $R_1 I^2$, vue ci-après, dissipée dans la résistance R_1 , représente des pertes. Ce montage est donc à proscrire dans toutes les installations où le rendement est primordial. On utilise alors des hacheurs, convertisseurs de haute technologie, inventés dans la seconde moitié du XX^e siècle.

5.4 Diviseur de courant

Un **diviseur de courant** permet de diviser un courant entre deux dipôles en parallèle (voir figure 7.30).

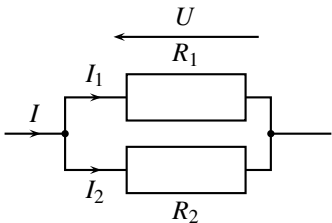


Figure 7.30 – Diviseur de courant.

La même tension U est aux bornes des deux résistances. La loi d'Ohm mène à :

$$U = R_1 I_1 = R_2 I_2.$$

La loi des nœuds impose :

$$I = I_1 + I_2 = I_1 + \frac{R_1}{R_2} I_1.$$

Finalement :

$$I_1 = \frac{R_2}{R_1 + R_2} I$$

et de même

$$I_2 = \frac{R_1}{R_1 + R_2} I.$$

6 Résistances de sortie et d'entrée

6.1 Résistance d'entrée

Un dispositif électrique, comme un ordinateur vu de l'alimentation, un moteur électrique, une lampe, peut être compliqué. Toutefois, dans tous les cas, un générateur extérieur impose une certaine tension et le dispositif électrique, qui constitue le récepteur, laisse circuler un certain courant à travers lui afin de fonctionner. On peut alors définir le rapport entre la tension imposée et l'intensité du courant, homogène à une résistance. Cette résistance est nommée **résistance d'entrée** du dispositif électrique.

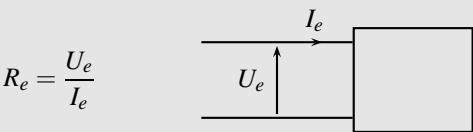


Figure 7.31 – Résistance d'entrée.

La **résistance d'entrée** d'un dipôle passif est la résistance équivalente vue de l'extérieur. Elle modélise électriquement le comportement du dipôle.

6.2 Résistance de sortie

On a vu au paragraphe 4.3 qu'un générateur de tension présente une certaine résistance R_g . Cette résistance interne au générateur est nommé **résistance de sortie**, car elle est placée en sortie du générateur.

La **résistance de sortie** d'un dipôle actif est la résistance du générateur de Thévenin équivalent vue de l'extérieur. Elle modélise électriquement la chute de tension du dipôle quand il débite un courant.

7 Point de fonctionnement d'un circuit

Soit un circuit constitué d'un dipôle actif, le générateur, et d'un dipôle passif, le récepteur. On cherche à trouver l'intensité du courant qui traverse le circuit ainsi que la tension aux bornes des dipôles. L'ensemble est modélisé par la figure 7.34 page suivante.

7.1 Caractéristique d'un dipôle

La caractéristique d'un dipôle est sa loi courant-tension. Elle peut prendre la forme d'une équation, $I(U)$, ou d'un graphe, où I est tracé en fonction de U . La représentation graphique est surtout utilisée dans le cas où le dipôle n'est pas linéaire, ou qu'elle ne se met pas sous la forme d'une équation.

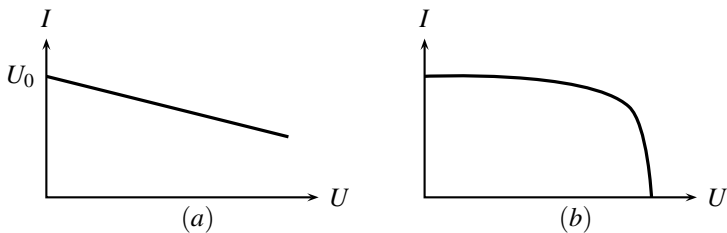


Figure 7.32 – Exemples de caractéristique de dipôle actif (a) linéaire, (b) non linéaire.

Le graphe de la caractéristique d'un dipôle actif linéaire correspond à l'équation

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

$U = U_0 - R_g I$. Un dipôle actif non linéaire peut avoir des formes diverses.
La caractéristique d'un dipôle passif passe par le point $(I = 0, U = 0)$ car l'intensité du courant qui le traverse est nulle quand il n'est pas alimenté. Un dipôle passif peut être tant linéaire, une résistance par exemple, que non linéaire.

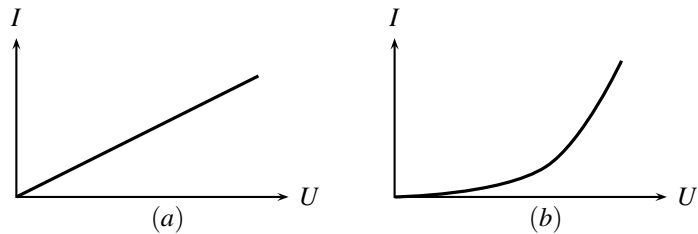


Figure 7.33 – Exemples de caractéristique de dipôle passif (a) linéaire, (b) non linéaire.

7.2 Résolution graphique

Le **point de fonctionnement** du circuit correspond à la valeur de l'intensité I du courant qui traverse les deux dipôles, et à celle de la tension à leurs bornes. On l'obtient graphiquement en superposant les deux caractéristiques. Par exemple :

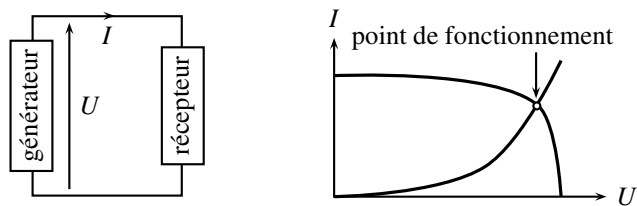


Figure 7.34 – Point de fonctionnement d'un circuit.

7.3 Résolution algébrique

La résolution algébrique n'est possible que si l'on est capable d'exprimer les fonctions $I(U)$ pour les dipôles. C'est possible dans les cas de dipôles linéaires simples, comme un générateur décrit par son modèle de Thévenin, ou une résistance, par exemple la résistance d'entrée R d'un récepteur :

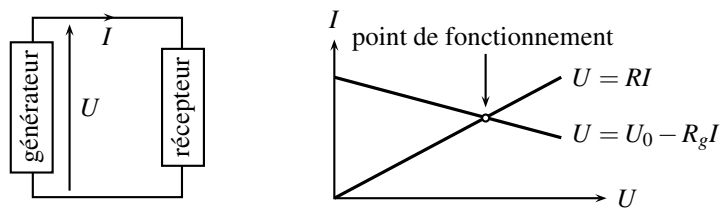


Figure 7.35 – Point de fonctionnement d'un circuit linéaire.

$$\left\{ \begin{array}{l} U = U_0 - R_g I \\ U = RI \end{array} \right. \text{ implique } \left\{ \begin{array}{l} U = \frac{R}{R + R_g} U_0 \\ I = \frac{U_0}{R + R_g} \end{array} \right.$$

Toutefois, de nombreux autres cas existent, où seule une résolution graphique est possible.

8 Puissance et énergie électriques

8.1 Les unités

L'**énergie** s'exprime en joule⁶, symbolisé par la lettre J majuscule. On montre en mécanique que le joule s'exprime, dans le système international, en $\text{kg} \cdot \text{m}^2 \cdot \text{s}^{-2}$.

Une **puissance** s'exprime en watt⁷, symbolisé par la lettre W majuscule. La puissance est, par définition, l'énergie développée par unité de temps. Par exemple, si un dispositif consomme une énergie de 20 J pendant 5 s, alors la puissance moyenne consommée vaut $\frac{20}{5} = 4 \text{ W}$.

Quand l'énergie disponible est une fonction $\mathcal{E}(t)$ du temps, on définit alors la puissance instantanée par :

$$\mathcal{P}(t) = \frac{d\mathcal{E}}{dt}.$$

Une énergie s'exprime en joule, unité de symbole J.

Une puissance est homogène à une énergie divisée par un temps ; elle s'exprime en watt, unité de symbole W.

8.2 Puissance électrique

Soit un circuit très simple, constitué d'un générateur et d'un dipôle passif. On mesure, au moyen d'un wattmètre, schématisé par un $\textcircled{W}$, la puissance délivrée par le générateur et consommée par le dipôle passif (voir figure 7.36). On remarque que le wattmètre ne modifie ni l'intensité du courant qui le traverse, ni la tension.

On observe expérimentalement que la puissance est proportionnelle tant à U qu'à I et s'exprime par :

$$\mathcal{P} = UI.$$

Si les grandeurs dépendent du temps, il est d'usage de les noter avec des lettres minuscules :

$$\mathcal{P}(t) = u(t)i(t).$$

6. En l'honneur de James Prescott Joule, 1818 – 1889, physicien britannique, auteur de travaux fondamentaux en thermodynamique, sur l'équivalence entre travail et transfert thermique. Il étudia de plus l'énergie dissipée dans une résistance, nommé depuis effet Joule.

7. En l'honneur de James Watt, 1736 – 1819, ingénieur écossais, qui développa les machines à vapeur, sources de la révolution industrielle.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

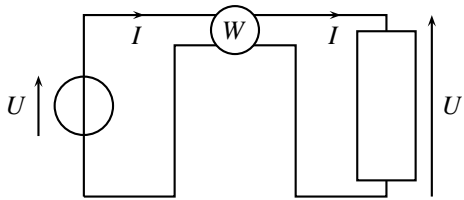


Figure 7.36 – Circuit de mesure de puissance.

8.3 Puissance Joule dans une résistance

Lorsqu'un électron parcourt un fil électrique, son mouvement n'est pas une ligne droite. Il est en effet en mouvement dans un milieu constitué d'atomes métalliques et de cations fixes. Il est alors soumis à de nombreuses interactions avec les ions du réseau ou avec les plans de rupture de l'empilement des atomes. Ces interactions, dénommées chocs, perturbent sa trajectoire. L'électron perd de l'énergie à chaque choc, au profit des atomes et cations du fil, qui s'échauffent alors. Cette perte d'énergie des électrons, et l'échauffement du fil qui en résulte, est nommée effet Joule.

Quelle est la puissance dissipée par effet Joule dans une résistance de valeur R ? Aux bornes de la résistance :

$$\begin{cases} \mathcal{P} = UI \\ U = RI \end{cases} \quad \text{implique} \quad \mathcal{P} = RI^2 = \frac{U^2}{R}.$$

L'**effet Joule** est la transformation d'énergie électrique en énergie thermique. La puissance Joule dissipée dans une résistance de valeur R , parcourue par un courant d'intensité I , soumise à une tension U , est :

$$\mathcal{P} = RI^2 = \frac{U^2}{R}.$$

8.4 Énergie stockée dans un condensateur ou une bobine

Les condensateurs et les bobines emmagasinent de l'énergie. Pour la calculer, on écrit, pour chaque composant, la puissance disponible, puis on en déduit l'énergie stockée.

a) Cas du condensateur

La puissance $\mathcal{P}$, disponible pour le condensateur est :

$$\begin{cases} \mathcal{P} = ui \\ i = C \frac{du}{dt} \end{cases} \quad \text{donc} \quad \mathcal{P} = Cu \frac{du}{dt} = \frac{d}{dt} \left(\frac{1}{2} Cu^2 \right).$$

Or la puissance instantanée $\mathcal{P}(t)$ représente l'énergie par unité de temps, c'est-à-dire $\mathcal{P}(t) = \frac{d\mathcal{E}}{dt}$. On en déduit ainsi que l'énergie du condensateur est :

$$\mathcal{E}_{\text{élec}} = \frac{1}{2}Cu^2.$$

Il convient d'insister sur le fait suivant :

Un condensateur ne consomme aucune énergie. Il stocke de l'énergie pour la délivrer ensuite.

b) Cas d'une bobine

La puissance $\mathcal{P}$, disponible pour la bobine est :

$$\begin{cases} \mathcal{P} = ui \\ u = L \frac{di}{dt} \end{cases} \quad \text{donc} \quad \mathcal{P} = Li \frac{di}{dt} = \frac{d}{dt} \left(\frac{1}{2} Li^2 \right).$$

On en déduit ainsi que l'énergie de la bobine :

$$\mathcal{E}_{\text{magné}} = \frac{1}{2} Li^2.$$

Il convient d'insister sur le fait suivant :

Une bobine ne consomme aucune énergie. Elle stocke de l'énergie pour la délivrer ensuite.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

SYNTHÈSE

SAVOIRS

- savoir que la charge électrique est quantifiée
- relier l'intensité du courant au débit de charge
- exprimer la condition de l'ARQS en fonction de la taille du circuit et de la fréquence
- ordres de grandeur d'intensités du courant
- ordres de grandeur de tensions
- relations entre l'intensité du courant et la tension pour une résistance, un condensateur, une bobine
- ordres de grandeur de résistances, capacités, inductances
- puissance dissipée par effet Joule dans une résistance
- énergies stockées dans un condensateur et une bobine
- modèle de Thévenin d'une source de tension non idéale
- établir la relation du diviseur de tension
- établir la relation du diviseur de courant
- connaître les notions de résistance d'entrée et de sortie

SAVOIR-FAIRE

- utiliser la loi des mailles
- algébriser les grandeurs électriques
- utiliser les conventions générateur et récepteur
- remplacer une association série ou parallèle de résistance par une résistance unique
- utiliser la relation du diviseur de tension
- utiliser la relation du diviseur de courant

MOTS-CLÉS

- | | | |
|-----------------------------|-----------------------|----------------------------|
| • charge électrique | • énergie | • résistance de sortie |
| • intensité du courant | • résistance | • caractéristique d'un di- |
| • tension | • condensateur | pôle |
| • potentiel | • capacité | • point de fonctionnement |
| • masse ou référence de po- | • bobine | d'un circuit |
| tentiel | • inductance | |
| • puissance | • résistance d'entrée | |

S'ENTRAÎNER

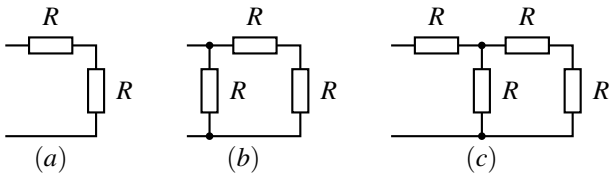
7.1 Électrolyse (★)

On électrolyse une solution de nitrate d'argent AgNO_3 : la cathode fixe des atomes d'argent, qui résultent de la capture d'un électron par des ions Ag^+ de la solution. Dans les conditions de l'expérience, on recueille 108 g d'argent sur la cathode en faisant passer une charge $Q = 96500 \text{ C}$ entre les électrodes.

Quelle durée τ permet de fixer 10 g d'argent à la cathode, sous un courant d'intensité $I = 5 \text{ A}$?

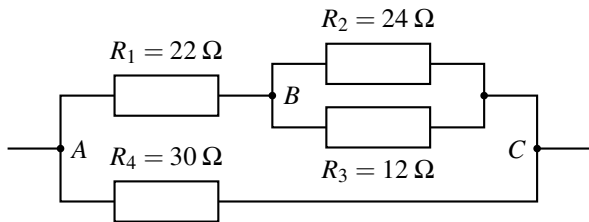
7.2 Associations de résistances (★)

Déterminer en fonction de R la résistance équivalente des dipôles suivants :



7.3 Association de résistances et intensités (★)

On donne $U_{AC} = 30 \text{ V}$. Déterminer :

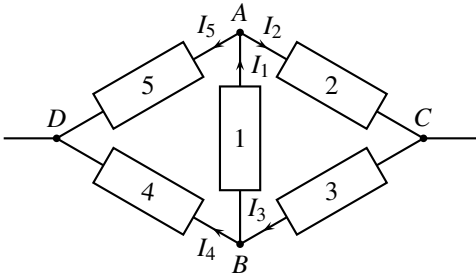


1. la résistance équivalente entre les nœuds A et C,
2. la valeur de la tension U_{BC} ,
3. les intensités des courants dans chaque résistance,
4. la puissance Joule dissipée dans R_4 .

7.4 Loi des nœuds et d'Ohm (★)

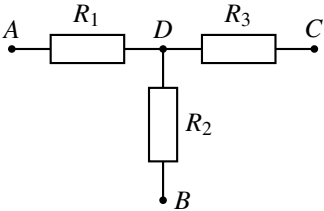
Toutes les résistances sont identiques et de valeur $R = 10^2 \Omega$. On donne $I_3 = 8 \text{ A}$, $I_4 = 10 \text{ A}$ et $I_5 = 5 \text{ A}$. Déterminer I_1 , I_2 , U_{AB} , U_{AC} et U_{BD} .

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS



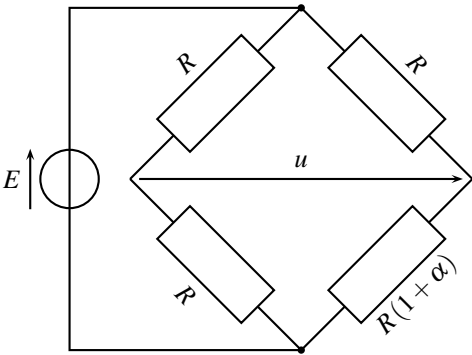
7.5 Loi des nœuds en terme de potentiels (★)

En partant de la loi des nœuds et de la loi d'Ohm, exprimer le potentiel au nœud D en fonction des potentiels des nœuds A , B , et C .



7.6 Capteur de déformation (★)

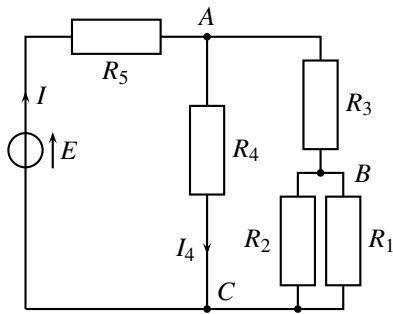
Les jauges de déformation sont des résistances variables. Elles permettent de réaliser des capteurs de force ou de pression et peuvent être utilisées localement afin de mesurer la déformation du corps sur lequel elles sont collées. Dans une telle jauge, la valeur de α change en fonction de la contrainte et on mesure la tension u .



1. Établir le lien entre u , E et α . On pourra nommer les potentiels aux quatre coins du losange et utilement préciser la masse dans le circuit.
2. Dans le cas où $\alpha = 0$, calculer la puissance consommée dans le circuit.

7.7 Circuit résistif (★)

Dans le circuit suivant, calculer U_{BC} , U_{CA} et I_4 . On précise $R_1 = 12\ \Omega$, $R_2 = 20\ \Omega$, $R_3 = 30\ \Omega$, $R_4 = 40\ \Omega$, $R_5 = 50\ \Omega$ et $E = 10\text{ V}$.



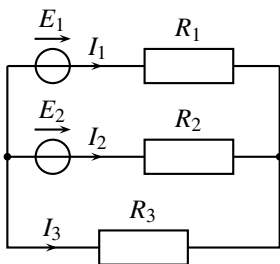
7.8 Modèle de pile (★)

Un générateur présente une différence de potentiel de 22 V quand il est traversé par une intensité du courant de 2 A. La différence de potentiel monte à 30 V lorsque l'intensité du courant descend à 1,2 A.

Préciser numériquement la résistance interne et la force électromotrice du modèle de Thévenin du générateur. Quelles sont les puissances, fournie par le générateur et perdue par effet Joule dans la seconde expérience ?

7.9 Deux générateurs réels (★)

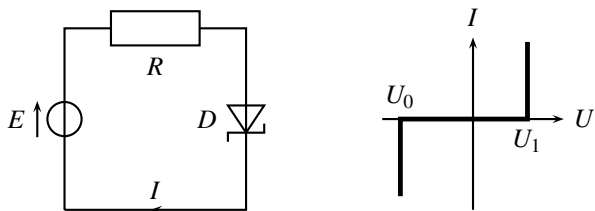
Dans le montage suivant, $E_1 = 24\text{ V}$, $E_2 = 32\text{ V}$, $R_1 = 2\ \Omega$, $R_2 = 4\ \Omega$ et $R_3 = 50\ \Omega$. Calculer les valeurs I_1 , I_2 et I_3 des trois intensités.



7.10 Point de fonctionnement d'un circuit à diode Zener (★)

Un générateur idéal de force électromotrice $E > 0$ est branché en série avec une résistance R et une diode Zener D dont la caractéristique courant-tension $I(U)$ est représentée ci-dessous.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS



- 1. Déterminer le point de fonctionnement du montage, c'est à dire l'intensité du courant et la tension aux bornes de la diode.
- 2. Même question lorsqu'on retourne la diode.

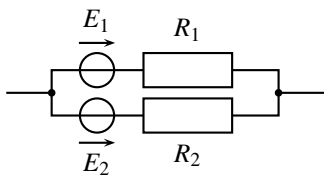
7.11 Batterie (★)

Une batterie est constituée de trois générateurs identiques en série, chacun modélisé par une tension à vide $E = 3 \text{ V}$ et de résistance interne $R = 0,1 \Omega$.

- 1. Quelle est la tension à vide E' de la batterie ? sa résistance interne R' ?
- 2. Quel est l'intensité du courant débité si la batterie est court-circuitée ? Conclusion ?

7.12 Générateurs en parallèle (★★)

Deux générateurs, modélisés par leurs schémas de Thévenin, de tensions à vide E_1 et E_2 , de résistances internes R_1 et R_2 , sont branchés en parallèle.

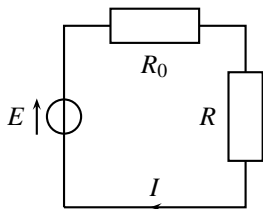


Établir que l'ensemble est équivalent à un unique générateur, de tension à vide $E = \frac{\alpha E_1 + \beta E R_2}{R_1 + R_2}$

et de résistance interne $R = \frac{\alpha \beta}{R_1 + R_2}$, où α et β dépendent de R_1 et R_2 .

7.13 Adaptation de puissance (★)

Un générateur présente une tension à vide E et une résistance interne R_0 . Il est branché en série avec une résistance de valeur R . Que doit valoir R afin que la puissance dissipée dans la résistance de valeur R soit maximale ? Que vaut alors la puissance Joule dans cette résistance ?



APPROFONDIR

7.14 Association de dipôles, utilisation de la caractéristique (★)

Soit un dipôle qu'on suppose caractérisé par une relation $u = f(i)$. La courbe représentative de cette relation est appelée caractéristique.

1. On considère que le dipôle est représenté en convention générateur. Indiquer en justifiant la réponse si le dipôle fonctionne en générateur ou en récepteur suivant le quadrant du plan donnant la tension u en fonction de l'intensité i dans lequel se trouve le point de la caractéristique.

2. Même question si le dipôle est en convention récepteur.

3. Soient deux dipôles D_1 et D_2 , le premier de caractéristique $u = E - Ri$ en convention générateur et le second de caractéristique $u = E' - R'i$ en convention générateur. On relie ces deux dipôles et on note u la tension à leurs bornes et i l'intensité les traversant. Peut-on utiliser pour les deux dipôles une unique convention ? Justifier la réponse.

On prendra dans la suite $E = 3,0 \text{ V}$, $E' = 6,0 \text{ V}$, $R = 1,0 \text{ k}\Omega$ et $R' = 6,0 \text{ k}\Omega$.

4. On suppose dans cette question que D_1 est en convention générateur et D_2 en convention récepteur.

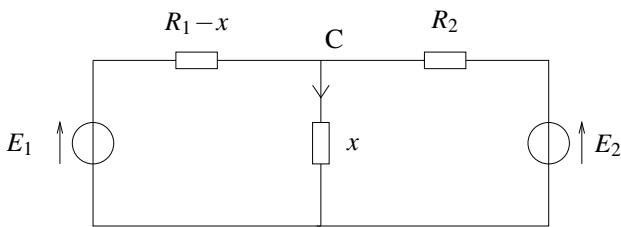
- a. Déterminer graphiquement puis par le calcul la tension u et l'intensité i .
- b. Préciser le caractère générateur ou récepteur de chacun des deux dipôles.

5. On inverse les bornes du dipôle D_2 .

a. Quelles en sont les conséquences pour l'équation de la caractéristique de ce dipôle ainsi que pour la convention utilisée pour ce dernier ?

- b. Déterminer graphiquement puis par le calcul la tension u et l'intensité i .
- c. Préciser le caractère générateur ou récepteur de chacun des deux dipôles.

7.15 Comparaison de deux tensions (★)



1. On considère le circuit représenté ci-dessus. On notera I_1 l'intensité du courant traversant $R_1 - x$ et I_2 celle du courant traversant R_2 . Peut-on appliquer la relation du pont diviseur de tension pour déterminer la tension aux bornes de x ? Justifier la réponse.

2. On cherche à calculer toutes les intensités du circuit. Montrer qu'il suffit de résoudre un système de deux équations à deux inconnues pour cela.

3. Écrire ce système en prenant I_1 et I_2 comme inconnues.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

4. Le résoudre.
5. En déduire l'intensité dans x .
6. On règle la valeur de x pour que I_2 soit nulle. Que vaut alors le rapport $\frac{E_2}{E_1}$?
7. Justifier qu'on puisse comparer deux tensions à l'aide de ce dispositif.
8. Que pensez-vous de son utilisation lorsque les tensions sont très différentes ?

CORRIGÉS

7.1 Électrolyse

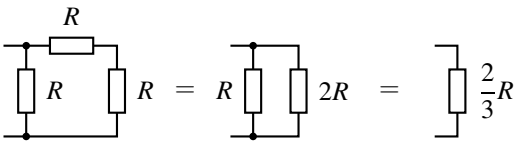
La charge qui traverse les électrodes pendant la durée τ est $q = I \times \tau$. Cette charge sert à fixer $m = 10 \text{ g}$ d'argent :

$$\begin{aligned} Q = 96500 \text{ C} &\longleftrightarrow M = 108 \text{ g} & \text{donc } q = \frac{m}{M} Q = I \tau. \\ q &\longleftrightarrow m = 10 \text{ g} \end{aligned}$$

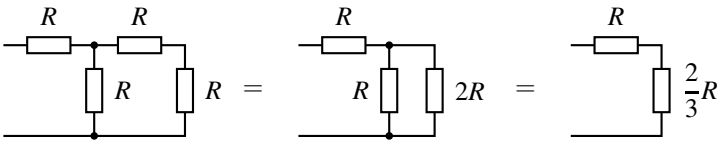
D'où $\tau = \frac{m}{M} \frac{Q}{I} = 1,78.10^3 \text{ s} \simeq 29 \text{ min}.$

7.2 Associations de résistances

Dans le circuit (a), les deux résistances sont en série, leur résistance équivalente est $R_a = 2R$. Dans le circuit (b), on utilise la résistance équivalente $2R$ des deux résistances en série pour trouver deux résistances en parallèle : $\frac{1}{R_b} = \frac{1}{R} + \frac{1}{2R} = \frac{3}{2R}$ donc $R_b = \frac{2}{3}R$.



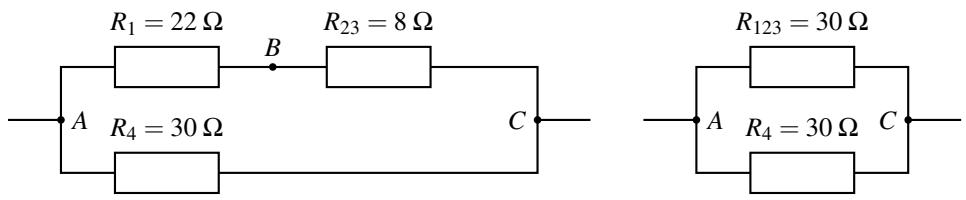
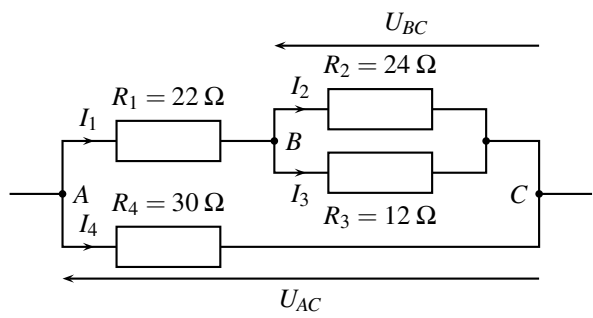
Dans le circuit (c), même méthode : $R_c = R + \frac{2}{3}R = \frac{5}{3}R$



7.3 Association de résistances et intensités

1. R_2 et R_3 sont en parallèle, équivalents à une unique résistance R_{23} telle que : $\frac{1}{R_{23}} = \frac{1}{R_2} + \frac{1}{R_3}$ donc $R_{23} = 8 \Omega$.
 R_1 et R_{23} sont alors en série, équivalents à une unique résistance $R_{123} = R_1 + R_{23} = 30 \Omega$.
Finalement R_4 et R_{123} sont en parallèle, équivalents à une unique résistance R_{1234} telle que : $\frac{1}{R_{1234}} = \frac{1}{R_4} + \frac{1}{R_{123}}$ donc $R_{1234} = 15 \Omega$.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS



2. Les résistances R_1 et R_{23} forment un diviseur de tension : $\frac{u_{BC}}{u_{AC}} = \frac{R_{23}}{R_{23} + R_1}$ donc $u_{BC} = \frac{R_{23}}{R_{23} + R_1} u_{AC} = 8 \text{ V}$.
3. $U_{AB} = U_{AC} - U_{BC} = 22 \text{ V}$ puis, avec la loi d'Ohm, avec le courant orienté en convention récepteur, $I_1 = \frac{U_{AB}}{R_1} = 1 \text{ A}$.
- De même $I_2 = \frac{U_{BC}}{R_2} = 0,33 \text{ A}$, $I_3 = \frac{U_{BC}}{R_3} = 0,67 \text{ A}$ et $I_4 = \frac{U_{AC}}{R_4} = 1 \text{ A}$.
4. $\mathcal{P}_J = R_4 I_4^2 = 30 \text{ W}$.

7.4 Loi des nœuds et d'Ohm

Loi des nœuds en B : $I_3 = I_1 + I_4$ donc $I_1 = -2 \text{ A}$.

Loi des nœuds en A : $I_1 = I_2 + I_5$ donc $I_2 = -7 \text{ A}$.

D'après la loi d'Ohm en convention récepteur : $U_{BA} = RI_1$ soit $U_{AB} = -U_{BA} = -RI_1 = 2.10^2 \text{ V}$.

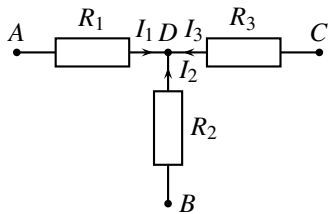
D'après la loi d'Ohm en convention récepteur : $U_{AC} = RI_2 = -7.10^2 \text{ V}$.

D'après la loi d'Ohm en convention récepteur : $U_{BD} = RI_4 = 10^3 \text{ V}$.

7.5 Loi des nœuds en terme de potentiels

La loi des nœuds indique que la somme des courants entrants dans un nœud est égale à celle sortant du nœud : $I_1 + I_2 + I_3 = 0$.

Exprimons ces intensités avec la loi d'Ohm, en faisant bien attention à se placer en convention récepteur : $\frac{u_{AD}}{R_1} + \frac{u_{BD}}{R_2} + \frac{u_{CD}}{R_3} = 0$.



Avec les potentiels : $\frac{V_A - V_D}{R_1} + \frac{V_B - V_D}{R_2} + \frac{V_C - V_D}{R_3} = 0$ d'où $V_D = \frac{\frac{V_A}{R_1} + \frac{V_B}{R_2} + \frac{V_C}{R_3}}{\frac{1}{R_1} + \frac{1}{R_2} + \frac{1}{R_3}}$.

7.6 Capteur de déformation

1. E impose une différence de potentiel. On choisit donc arbitrairement un potentiel nul au nœud sud du pont, donc un potentiel E au nœud nord. Soit V_E le potentiel du nœud est et V_O celui du nœud ouest : $u = V_E - V_O$.

Le courant d'intensité i qui circule dans la branche de gauche, vers le bas, s'exprime de deux façons différentes (attention à bien se placer en convention récepteur) : $i = \frac{E - V_O}{R} = \frac{V_O - 0}{R}$ d'où $V_O = \frac{E}{R}$.

Le courant d'intensité i' qui circule dans la branche de droite, vers le bas, s'exprime aussi de deux façons différentes (attention à bien se placer en convention récepteur) : $i = \frac{E - V_E}{R} = \frac{V_E - 0}{R(1 + \alpha)}$ d'où $V_E = \frac{(1 + \alpha)E}{2 + \alpha}$. Finalement : $u = \frac{(1 + \alpha)E}{2 + \alpha} - \frac{E}{2} = \frac{\alpha E}{2(2 + \alpha)}$.

2. E voit 2 résistances $2R$ en parallèle, c'est à dire une résistance R : $\frac{1}{R_{eq}} = \frac{1}{2R} + \frac{1}{2R} = \frac{1}{R}$.

D'où une puissance consommée : $\mathcal{P} = \frac{E^2}{R}$.

7.7 Circuit résistif

Il existe plusieurs méthodes pour résoudre cet exercice. On propose :

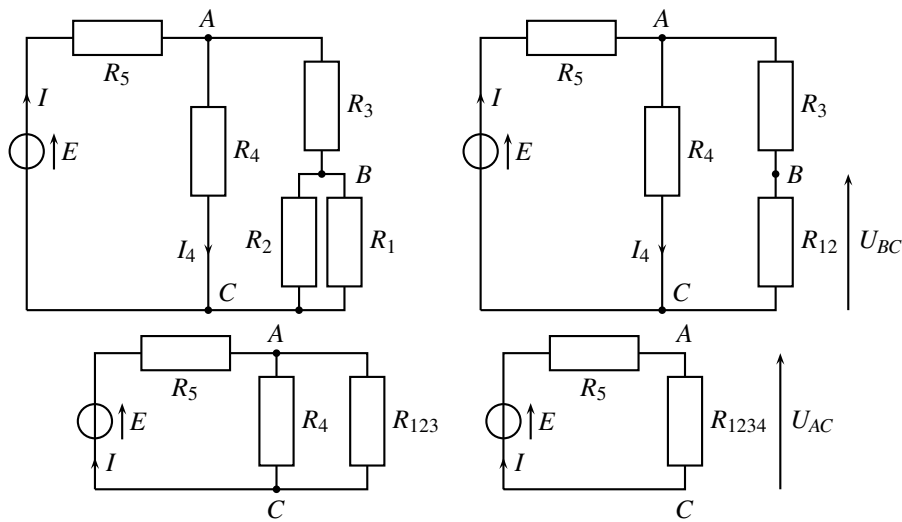
- calcul de la résistance équivalente à R_1, R_2, R_3 et R_4 entre A et C ,
- calcul de U_{CA} avec la formule du diviseur de tension entre E et U_{AC} (Attention ! Il est impossible d'appliquer directement un diviseur de tension entre R_4 et R_5 avec le circuit tel que dans l'énoncé. En effet, le courant doit être identique dans les deux résistances pour un diviseur de tension. Ce n'est pas le cas ici où $I_4 \neq I$),
- calcul de I_4 avec la loi d'Ohm aux bornes de R_4 ,
- calcul de U_{BC} avec la formule du diviseur de tension entre U_{AC} et U_{BC} ,

On a successivement : $\frac{1}{R_{12}} = \frac{1}{R_1} + \frac{1}{R_2}$ soit $R_{12} = \frac{R_1 R_2}{R_1 + R_2} = 7,5 \, \Omega$,

et : $R_{123} = R_3 + R_{12} = R_3 + \frac{R_1 R_2}{R_1 + R_2} = 37,5 \, \Omega$,

puis : $\frac{1}{R_{1234}} = \frac{1}{R_4} + \frac{1}{R_{123}}$ soit $R_{1234} = \frac{R_4 R_{123}}{R_4 + R_{123}} = 19,3 \, \Omega$.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS



Attendu que le même courant circule dans les résistances R_5 et R_{1234} , la formule du diviseur de tension est applicable. Il faut prêter attention aux sens des tensions, qui doivent être identiques : $\frac{U_{AC}}{E} = \frac{R_{1234}}{R_{1234} + R_5}$ d'où $U_{AC} = \frac{R_{1234}}{R_{1234} + R_5} E = 2,8 \text{ V}$.

En en déduit dès lors : $U_{CA} = -U_{AC} = -2,8 \text{ V}$.

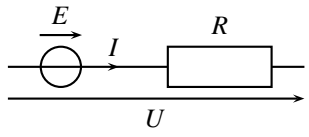
En convention récepteur, $U_{AC} = R_4 I_4$ donc $I_4 = \frac{U_{AC}}{R_4} = 7,0 \cdot 10^{-2} \text{ A}$.

Regardons le schéma en haut à droite, le même courant circule dans les résistances R_3 et R_{12} . Avec des tensions dans le même sens, la formule du diviseur de tension mène à :

$$\frac{U_{BC}}{U_{AC}} = \frac{R_{12}}{R_{12} + R_3} \quad \text{d'où} \quad U_{BC} = \frac{R_{12}}{R_{12} + R_3} U_{AC} = 5,6 \cdot 10^{-1} \text{ V}.$$

7.8 Modèle de pile

Le modèle de Thévenin du générateur est le suivant. On prend bien garde à orienter la tension à vide E et l'intensité I du courant en convention générateur.



L'énoncé précise les valeurs de U en fonction de I , liés par la relation $U = E - RI$. D'où le système :

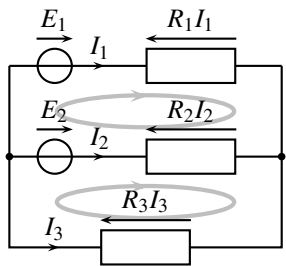
$$\begin{cases} E - R \times 2 = 22 \\ E - R \times 1,2 = 30 \end{cases} \quad \text{donc} \quad \begin{cases} E = 31,2 \text{ V} \\ R = 10 \Omega \end{cases}$$

Pour la seconde expérience, la puissance fournie est $\mathcal{P}_G = EI = 37,4 \text{ W}$ et celle dissipée par effet Joule est $\mathcal{P}_J = RI^2 = 14,4 \text{ W}$.

7.9 Deux générateurs réels

Trois inconnues sont à trouver : I_1 , I_2 et I_3 . Il faut donc trois équations, obtenues avec la loi des nœuds et la loi des mailles.

La loi des mailles (la somme des intensités des courants entrants dans le nœud est égale à la somme des intensités des courants sortants du nœud, qui est ici nulle) mène à : $I_1 + I_2 + I_3 = 0$. Le loi des mailles est écrite dans la maille du dessus, puis dans celle au dessous. On tourne dans le sens indiquées par les ellipses orientées grises : $E_1 - R_1 I_1 + R_2 I_2 - E_2 = 0$ et $E_2 - R_2 I_2 + R_3 I_3 = 0$.



On élimine I_3 avec la loi des nœuds ($I_3 = -I_1 - I_2$) :

$$\begin{cases} E_1 - R_1 I_1 + R_2 I_2 - E_2 = 0 \\ E_2 - R_2 I_2 + R_3 (-I_1 - I_2) = 0 = E_2 - (R_2 + R_3) I_2 - R_3 I_1 \end{cases}$$

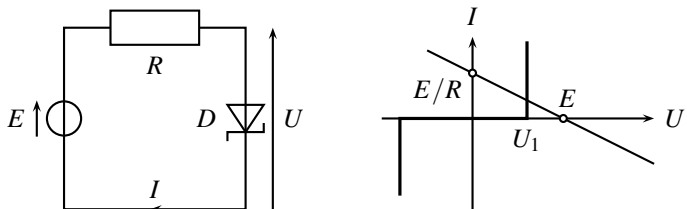
Et on résoud le système en I_1 et I_2 . On multiplie, par exemple, la première équation par R_3 , la seconde par $-R_1$, puis on les somme pour éliminer I_1 et trouver I_2 : $I_2 = \frac{E_2 R_1 - (E_1 - E_2) R_3}{R_2 R_3 + R_1 (R_2 + R_3)} = 1,5 \text{ A}$.

Et : $I_1 = \frac{E_2}{R_3} - \frac{R_2 + R_3}{R_3} I_2 = -1,0 \text{ A}$.

Puis : $I_3 = -I_1 - I_2 = -5,2 \cdot 10^{-1} \text{ A}$.

7.10 Point de fonctionnement d'un circuit à diode Zener

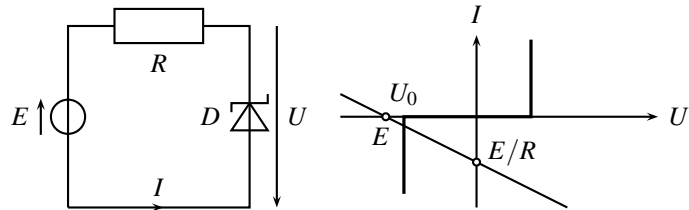
Notons U la tension aux bornes de la diode Zener, en convention récepteur. Le générateur et la résistance imposent $U = E - RI$, soit $I = \frac{E}{R} - \frac{U}{R}$.



1. Deux cas : si $E > U_1$ alors le point de fonctionnement est en $\left(U = U_1, I = \frac{E - U_1}{R} \right)$, si $E < U_1$, alors il est en $(U = E, I = 0)$.

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

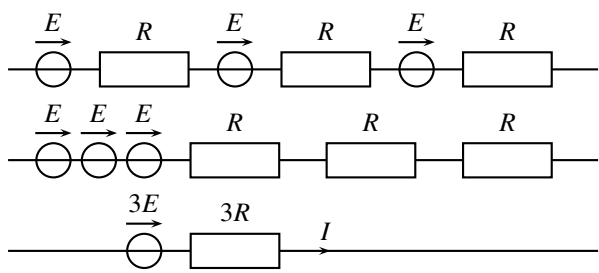
2. On change l'orientation de I et U , toujours en convention récepteur, afin de garder I dans le sens de la diode et conserver ainsi la caractéristique $I(U)$ inchangée. Aux bornes du générateur et de la résistance, $U = -E - RI$ et donc $I = -\frac{U+E}{R}$.



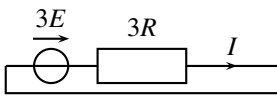
Deux cas : si $E < U_0$ alors le point de fonctionnement est en $\left(U = U_0, I = -\frac{E + U_0}{R} \right)$, si $E > U_0$, alors il est en $(U = E, I = 0)$.

7.11 Batterie

1. Dans un circuit série, la place des éléments n'a pas d'importance. On modifie alors le schéma pour arriver au modèle de Thévenin de la batterie. La tension à vide de la batterie, c'est à la tension lorsqu'aucun courant n'est débité, est $E' = 3E = 9 \text{ V}$; sa résistance interne vaut $R' = 3R = 0,3 \Omega$.



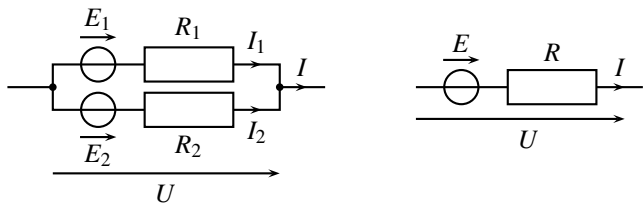
2. Si la batterie est court-circuitée, alors on impose une tension nulle à ses bornes. Dès lors : $3E - 3RI = 0$ et ainsi $I = E/R = 30 \text{ A}$. L'intensité du courant est très importante, la batterie se décharge très rapidement.



7.12 Générateurs en parallèle

Il faut mettre le dipôle sous la forme de droite. La méthode proposée consiste à établir les valeurs de I_1 ou de I_2 en fonction de I pour l'injecter dans la loi des mailles.

La loi des nœuds impose $I_1 + I_2 = I$. De plus, la tension U vaut : $U = E_1 - R_1 I_1 = U = E_2 - R_2 I_2 = E_2 - R_2 (I - I_1)$.



En égalant les deux expressions de U : $E_1 - R_1 I_1 = E_2 - R_2 I + R_2 I_1$ d'où $I_1 = \frac{E_1 - E_2 + R_2 I}{R_1 + R_2}$.

On remplace I_1 dans l'expression de U : $U = E_1 - R_1 I_1 = E_1 - R_1 \frac{E_1 - E_2 + R_2 I}{R_1 + R_2}$

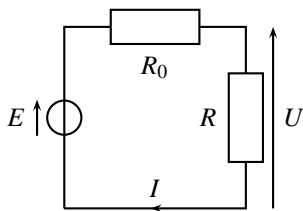
$$U = \frac{R_2 E_1 + R_1 E_2}{R_1 + R_2} - \frac{R_1 R_2}{R_1 + R_2} I.$$

Finalement : $E = \frac{R_2 E_1 + R_1 E_2}{R_1 + R_2}$ et $R = \frac{R_1 R_2}{R_1 + R_2}$.

7.13 Adaptation de puissance

La tension U aux bornes de R est obtenue avec un diviseur de tension : $\frac{U}{E} = \frac{R}{R + R_0}$. La

puissance aux bornes de R est : $\mathcal{P}_J = \frac{U^2}{R} = \frac{R}{(R + R_0)^2} E^2$.



Il faut maximiser la fonction $\mathcal{P}_J(R)$. Cette fonction est maximale quand sa dérivée est nulle :

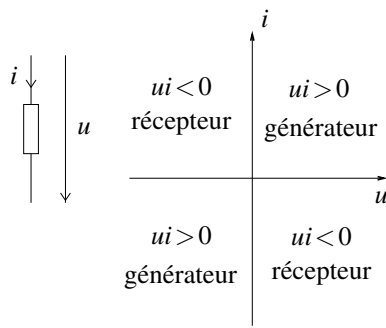
$$\frac{d\mathcal{P}_J}{dR} = \frac{d}{dR} \left(R(R + R_0)^{-2} \right) = (R + R_0)^{-2} + R(-2)(R + R_0)^{-3} = \frac{R_0 - R}{(R + R_0)^3}.$$

La dérivée s'annule quand $R = R_0$. Le graphe de la fonction $R \mapsto \frac{R}{(R + R_0)^2} E^2$ montre qu'elle n'admet qu'un maximum. La puissance dissipée dans R est donc maximale quand $R = R_0$, et alors $\mathcal{P}_J = \frac{E^2}{4R_0^2}$.

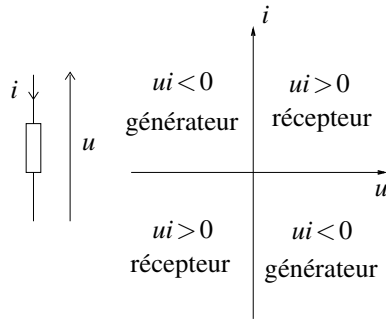
7.14 Association de dipôles, utilisation de la caractéristique

1. En convention générateur, la puissance fournie est donnée par $P_{\text{fournie}} = ui$. On a un générateur si la puissance fournie est effectivement fournie à savoir si $ui > 0$; sinon on a un récepteur. On en déduit donc :

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

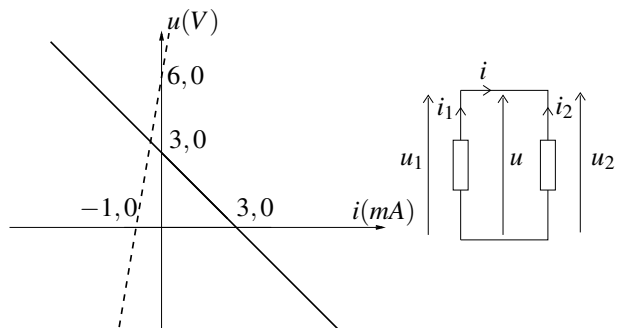


2. De même, en convention récepteur, la puissance reçue est donnée par $P_{\text{recue}} = ui$. On a un récepteur si la puissance reçue est effectivement reçue à savoir si $ui > 0$; sinon on a un générateur. On en déduit donc :



3. On effectue le choix d'une convention pour un dipôle. Dans la suite, on suppose par exemple qu'on prend la convention générateur pour le premier dipôle. Les sens de u et de i sont imposés alors pour le second dipôle se retrouvent dans l'exemple considéré en convention récepteur. Il n'est donc pas possible de prendre une unique convention pour tous les dipôles.

4. On trace la caractéristique du dipôle 1, ce qui ne pose pas de problème puisque sa caractéristique est fournie dans la bonne convention. En revanche, pour le dipôle 2, on a $i_2 = -i$ compte tenu de la convention prise. Il faut donc effectuer une symétrie de la caractéristique par rapport à l'axe des ordonnées :



a. Il convient de déterminer l'intersection des deux caractéristiques soit graphiquement c'est-à-dire en relevant les coordonnées du point d'intersection des deux droites soit en résolvant le système

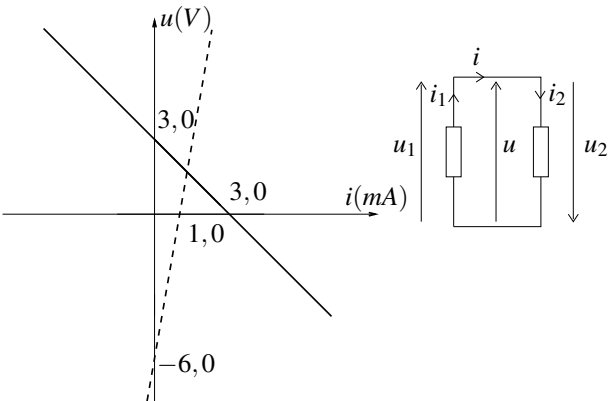
$$\begin{cases} u = 3,0 - 1,0 \cdot 10^3 i \\ u = 6,0 + 6,0 \cdot 10^3 i \end{cases}$$

On obtient dans les deux cas $i = -0,43 \text{ mA}$ et $u = 3,4 \text{ V}$.

b. On a la convention générateur pour le dipôle 1 et on obtient $ui < 0$ donc le dipôle 1 est récepteur.

On a la convention récepteur pour le dipôle 2 et on obtient $ui < 0$ donc le dipôle 2 est générateur.

5. a. Il n'y a toujours pas de soucis pour le dipôle 1 dont la caractéristique est fournie dans la « bonne » convention mais pour le dipôle 2, on a $u_2 = -u$ et $i_2 = i$: il faut donc procéder à une symétrie de la caractéristique par rapport à l'axe des abscisses :



b. On procède comme à la question précédente en résolvant $\begin{cases} u = 3,0 - 1,0 \cdot 10^3 i \\ u = -6,0 + 6,0 \cdot 10^3 i \end{cases}$ On obtient $i = 1,3 \text{ mA}$ et $u = 1,7 \text{ V}$.

c. On a la convention générateur pour le dipôle 1 et $ui > 0$ donc il est générateur. On a la convention récepteur pour le dipôle 2 et $ui > 0$ donc il est récepteur.

7.15 Comparaison de deux tensions

1. On ne peut pas appliquer la relation du pont diviseur de tension pour calculer la tension aux bornes de x car le courant circulant dans les résistances R_2 et $R_1 - x$ n'est pas nul.

2. Le circuit comprend trois branches donc il faut *a priori* déterminer trois intensités. Ces trois intensités ne sont pas indépendantes puisqu'elles sont reliées par une loi des nœuds. Finalement il reste deux inconnues soit un système 2×2 à résoudre.

3. L'écriture de deux lois des mailles donne $\begin{cases} R_1 I_1 + x I_2 = E_1 \\ x I_1 + (x + R_2) I_2 = E_2 \end{cases}$

CHAPITRE 7 – CIRCUITS ÉLECTRIQUES DANS L'ARQS

4. En utilisant la méthode de substitution ou en faisant des combinaisons linéaires de ces deux équations pour éliminer I_1 puis I_2 , on obtient $I_1 = \frac{x E_2 - (x + R_2) E_1}{x^2 - R_1 (x + R_2)}$ et $I_2 = \frac{x E_1 - R_1 E_2}{x^2 - R_1 (x + R_2)}$.
5. L'intensité dans x s'obtient alors par loi des nœuds soit $i = I_1 + I_2 = \frac{(x - R_1) E_2 - R_2 E_1}{x^2 - R_1 (x + R_2)}$.
6. Si la valeur de x est telle que $I_2 = 0$, on en déduit le rapport $\frac{E_2}{E_1} = \frac{x}{R_1}$.
7. La connaissance de x et R_1 permet donc de comparer les deux tensions E_1 et E_2 .
8. Il est possible d'effectuer cette comparaison si les tensions sont très différentes à condition de disposer de résistances de valeurs très différentes. Les gammes de valeurs de résistance seront donc un facteur limitant pour cette comparaison.

Circuit linéaire du premier ordre



Les intensités et tensions dans les circuits étudiés dans ce chapitre dépendent du temps. Pour les calculer, on est amené à résoudre une équation différentielle du premier ordre.

1 Exemple expérimental

1.1 Montage

On cherche à observer, puis prévoir par le calcul, le comportement du montage ci-dessous lorsqu'il est alimenté par une tension constante délivrée par un générateur basse fréquence, nommé dans la suite par son acronyme GBF. On observe simultanément les tensions délivrées par le GBF et aux bornes du condensateur sur les deux voies d'un oscilloscope.

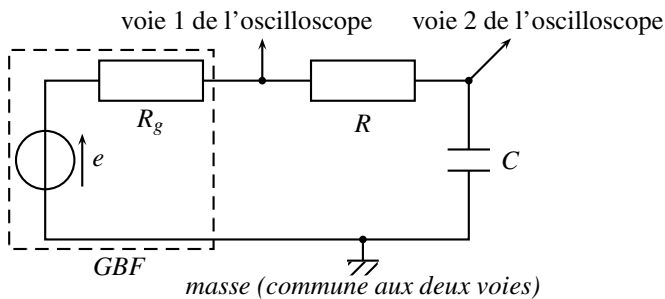


Figure 8.1 – Système étudié

1.2 Régimes transitoire puis permanent

On observe expérimentalement les formes d'onde de la figure 8.2, c'est-à-dire les graphes des tensions en fonction du temps, lorsque e passe de 0 à une tension positive (voie 1). La tension aux bornes du condensateur (voie 2) passe de 0 à 5 carreaux, c'est-à-dire de 0 à

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

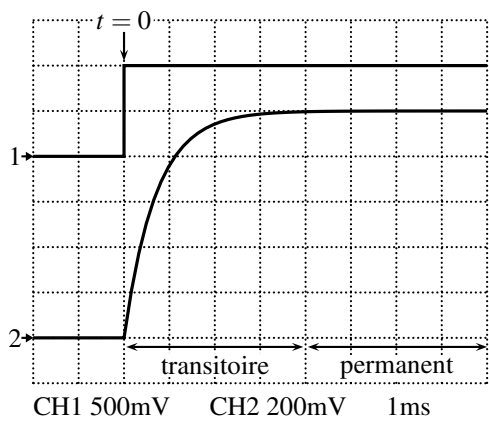
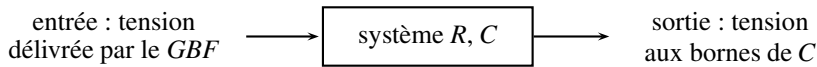


Figure 8.2 – Formes d’onde d’un essai indiciel.

$5 \times 200 \text{ mV par carreau} = 1 \text{ V}$, en approximativement $3,5 \text{ carreaux} \times 100 \text{ ms par carreau} = 350 \text{ ms}$. On définit alors deux régimes :

- le **régime établi**, ou **régime permanent**, pour lequel la sortie est de la même forme que l’entrée, ici constante,
- le **régime transitoire**, entre l’instant initial et le régime permanent.

Le montage est symboliquement représenté par :



Le signal délivré par le GBF est nommé **échelon de tension**. La réponse du système constitué de la résistance R et du condensateur de capacité C est nommé **réponse indicielle**.

2 Modélisation

2.1 Simplification du montage

Le GBF présente une résistance de sortie R_g , en série avec la résistance R . L’ensemble forme une unique résistance $R_{tot} = R_g + R$. Or R_g est très inférieure à R , donc $R_{tot} \simeq R$. Le montage devient, avec la représentation des différences de potentiel, ou tensions, entre les 4 coins :

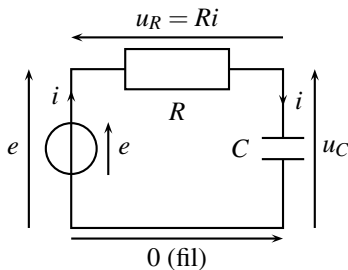


Figure 8.3 – Système simplifié.

2.2 Équation différentielle sur $u_C(t)$

La loi des mailles impose :

$$e = 0 + u_C + u_R.$$

Le GBF fait circuler un courant d'intensité i qui passe successivement à travers R et C . C'est le même courant à travers chacun des trois dipôles. Ainsi $u_R = Ri$ et $i = C \frac{du_C}{dt}$. La loi des mailles devient alors :

$$e = u_C + RC \frac{du_C}{dt}.$$

À chaque instant, la tension $u_C(t)$ est reliée à sa dérivée et à la tension e imposée par le GBF. Une telle relation est une **équation différentielle d'ordre un**, car elle contient u_C et la dérivée première de u_C .

2.3 Unités et homogénéité de l'équation différentielle

e et u_C sont des tensions en volts. L'unité de la grandeur X , dans le Système International, est notée $[X]$:

$$[e] = V \quad \text{et} \quad [u_C] = V.$$

La dérivée $\frac{du_C}{dt}$ est en volts divisé par des secondes car u_C est en volts et t en secondes :

$$\left[\frac{du_C}{dt} \right] = V.s^{-1}$$

On en déduit l'unité du produit RC en écrivant les unités de chaque terme dans l'équation différentielle :

$$e = u_C + RC \frac{du_C}{dt} \quad \text{donc} \quad V = V + [RC] \frac{V}{s}.$$

L'ensemble est homogène si $[RC] = s$.

Le produit RC est en secondes. Il représente la **constante de temps** du circuit.

Remarque

On retrouve ce résultat avec les unités des lois électrocinétiques aux bornes d'une résistance et d'un condensateur :

$$\begin{cases} u_R = Ri & \text{donc} & V = \Omega \times A \\ i = C \frac{du_C}{dt} & \text{donc} & A = F \times \frac{V}{s} \end{cases} \quad \text{implique} \quad \Omega \times F = s.$$

Le circuit RC présenté à la figure 8.3 est un système du premier ordre, de constante de temps $\tau = RC$.

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

2.4 Résolution de l'équation différentielle

a) Position du problème

Pour $t < 0$, la tension délivrée par le GBF est nulle. Sans entrée, la sortie est aussi nulle :

$$u_C(t < 0) = 0.$$

Cette solution vérifie bien l'équation différentielle : $e = 0$, $u_C = 0$ et $\frac{du_C}{dt} = 0$ (car u_C est une constante).

Pour $t > 0$, le GBF délivre un signal constant E_0 de 2 carreaux $\times$ 500 mV = 1 V dans l'exemple expérimental de la section précédente :

$$e(t > 0) = E_0 = 1 \text{ V}.$$

Il est courant d'alléger les notations en posant $RC = \tau$. Le système est alors décrit par l'équation différentielle pour $t > 0$:

$$\tau \frac{du_C}{dt} + u_C = E_0.$$

On a une équation différentielle du premier ordre. Il suffit alors d'une seule **condition initiale**, par exemple $u_C(0)$ pour déterminer complètement la solution. Elle est déterminée par la continuité de l'énergie dans le condensateur.

L'énergie $\frac{1}{2}Cu_C^2$ contenue dans le condensateur est continue, elle ne peut pas varier instantanément. Il en résulte que la tension u_C est, elle aussi, continue.

Ainsi $u_C(0^+) = u_C(0^-) = 0$. Finalement, on cherche la solution du couple {équation différentielle, condition initiale} :

$$\begin{cases} \tau \frac{du_C}{dt} + u_C = E_0 \\ u_C(0) = 0. \end{cases}$$

b) Établissement de la solution

L'équation à résoudre est une équation différentielle du premier ordre, avec un second membre non nul. D'après l'annexe sur les outils mathématiques, la solution est la somme :

- d'une **solution particulière**, notée $u_{C,p}$ (indice p pour « particulière »), cherchée sous la même forme que l'entrée E_0 ,
- et d'une **solution homogène**, notée $u_{C,h}$ (indice h pour « homogène »), solution de l'équation homogène :

$$\tau \frac{du_{C,h}}{dt} + u_{C,h} = 0.$$

La **solution particulière** est cherchée sous la même forme que l'entrée, ici constante : $u_{C,p} = \text{constante}$. Alors $\frac{du_{C,p}}{dt} = 0$ et l'équation différentielle devient $(\tau \times 0) + u_{C,p} = E_0$, donc $u_{C,p}(t) = E_0$.

La **solution homogène** est la solution de $\tau \frac{du_{C,h}}{dt} + u_{C,h} = 0$. On montre dans l'annexe mathématique qu'elle s'intègre en :

$$u_{C,h} = \mu \exp\left(-\frac{t}{\tau}\right),$$

où μ est une constante d'intégration inconnue à ce stade.

La solution de l'équation différentielle est donc, pour $t > 0$, la somme de la solution particulière et de la solution homogène :

$$u_C(t) = u_{C,p}(t) + u_{C,h}(t) = E_0 + \mu \exp\left(-\frac{t}{\tau}\right).$$

La constante d'intégration μ est calculée avec la **condition initiale** :

$$u_C(0) = 0 = E_0 + \mu \quad \text{donc} \quad \mu = -E_0.$$

Finalement, la solution du couple {équation différentielle, condition initiale} :

$$\begin{cases} \tau \frac{du_C}{dt} + u_C = E_0 \\ u_C(0) = 0 \end{cases}$$

est :

$$u_C(t) = E_0 \left(1 - \exp\left(-\frac{t}{\tau}\right)\right).$$

La solution obtenue est **homogène** quant aux unités. E_0 est en volt, t et τ en seconde donc leur rapport est sans unité. L'exponentielle opère sur un nombre sans dimension $\left(-\frac{t}{\tau}\right)$ et renvoie un nombre sans dimension, ainsi $1 - \exp\left(-\frac{t}{\tau}\right)$ est sans dimension. Finalement u_C a la même unité que E_0 , le volt.

2.5 Tracé et identification de la constante de temps

Le graphe obtenu sur la figure 8.4 est caractéristique d'un système du premier ordre alimenté par un échelon.

Que vaut la valeur finale ? Attendu que $\lim_{t \rightarrow \infty} \exp\left(-\frac{t}{\tau}\right) = 0$: $\lim_{t \rightarrow \infty} u_C(t) = E_0$.

On retrouve la solution particulière. Au bout d'un temps « suffisamment » long, à préciser dans la suite, le système a oublié la condition initiale ; la sortie ne dépend plus que de la valeur E_0 , imposée par l'entrée.

Quand le régime permanent, ou établi, débute-t-il ?

La valeur finale E_0 est atteinte à 1% près à la date θ telle que $u_C(\theta) = 0,99 \times E_0$:

$$u_C(\theta) = 0,99 \times E_0 = E_0 \left(1 - \exp\left(-\frac{\theta}{\tau}\right)\right) \quad \text{donc} \quad \theta = 4,6 \times \tau.$$

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

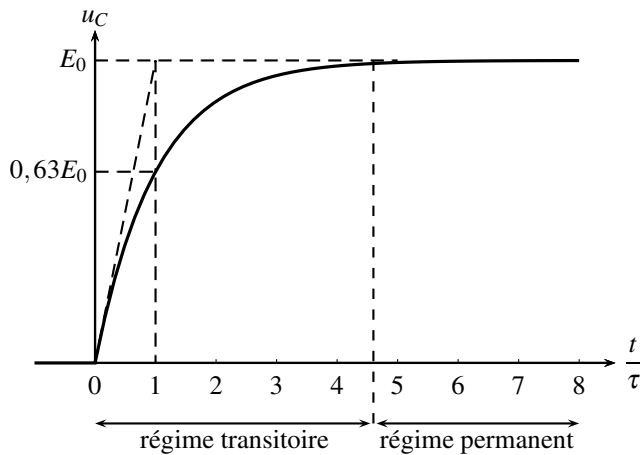


Figure 8.4 – Tension aux bornes du condensateur d'un système RC soumis à un échelon de tension.

Au bout d'une durée de $4,6\tau$, le système a atteint la valeur finale, à 1% près.

Le **temps de réponse** du système, c'est-à-dire la durée du régime transitoire, est conventionnellement de $4,6\tau$.

De plus, à la date $t = \tau$, $u_C(t)$ vaut $u_C(\tau) = E_0(1 - \exp(-1)) = E_0 \times 0,63$.

On en déduit une méthode d'**identification** de la valeur de τ pour un système du premier ordre :

- on repère la valeur finale de u_C , ici E_0 ,
- on note la date à laquelle u_C vaut $0,63 \times$ cette valeur finale : c'est τ .

Remarque

Il existe, en théorie, une autre méthode de détermination de τ , basée sur la valeur de la dérivée en $t = 0$:

$$\frac{du_C}{dt} = \frac{E_0}{\tau} \exp\left(-\frac{t}{\tau}\right) \quad \text{d'où} \quad \left(\frac{du_C}{dt}\right)_0 = \frac{E_0}{\tau}.$$

La tangente à l'origine, en pointillés sur le graphe, a donc pour équation $E_0 \frac{t}{\tau}$; elle prend la valeur E_0 en $t = \tau$. On en déduit qu'il suffit de tracer la tangente à l'origine et de déterminer quand elle vaut E_0 pour trouver la valeur τ . Toutefois, cette méthode est peu pratique expérimentalement. D'une part, la tracé de la tangente n'est pas possible sur un oscilloscope ; d'autre part, l'incertitude sur la position exacte de cette tangente est important lorsqu'on la trace avec des points expérimentaux.

2.6 Portrait de phase

Le **portrait de phase** du montage est la représentation graphique de la dérivée de u_C , notée $\frac{du_C}{dt}$ ou $\dot{u}_C$, en fonction de u_C .

Notons $y = \frac{du_C}{dt}$ et $x = u_C$. Alors l'équation différentielle $\tau \frac{du_C}{dt} + u_C = e$ devient $\tau y + x = e$, soit :

$$y = -\frac{x}{\tau} + \frac{e}{\tau} \quad \text{et pour } t > 0 : \quad y = -\frac{x}{\tau} + \frac{E_0}{\tau},$$

qui est l'équation d'une droite :

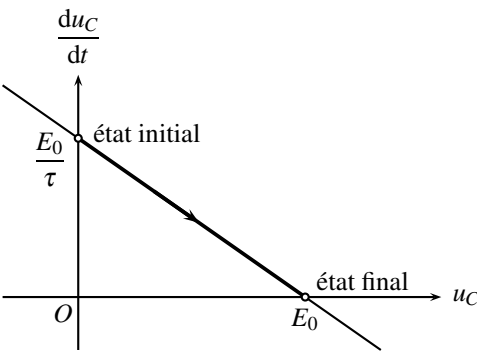


Figure 8.5 – Portrait de phase.

L'observation du portrait de phase permet de prévoir l'évolution du système avant toute résolution de l'équation différentielle. L'équation différentielle montre que le système se déplace sur une droite. La condition initiale, $u_C(0) = 0$, permet de placer le système à $t = 0$. Le système évolue ensuite jusqu'au point final où $\frac{du_C}{dt}$ est nul, c'est à dire avec l'intersection avec l'axe des abscisses ; il parcourt donc le segment de la gauche vers la droite.

Toutefois, le portrait de phase élimine le facteur temps. On ignore la durée prise par le système pour passer de l'état initial à l'état final.

2.7 Bilan de puissance

Multiplions la loi des mailles vue précédemment ($e = u_C + Ri$) par i :

$$e \times i = u_C \times i + Ri \times i.$$

Avec $i = C \frac{du_C}{dt}$, $u_C \times i = Cu_C \frac{du_C}{dt} = \frac{d}{dt} \left(\frac{1}{2} Cu_C^2 \right)$, ainsi :

$$ei = \frac{d}{dt} \left(\frac{1}{2} Cu_C^2 \right) + Ri^2.$$

La puissance délivrée par le GBF (ei) se répartit en puissance stockée par le condensateur

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

(l'énergie stockée par condensateur, sous forme électrique, est $\frac{1}{2}Cu_C^2$) et en puissance dissipée dans la résistance par effet Joule (Ri^2). Toutes ces puissances s'expriment en watt.

2.8 Bilan d'énergie

Attendu que la tension aux bornes du condensateur est $u_C(t) = E_0 \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$, l'intensité du courant qui la traverse vaut :

$$i(t) = C \frac{du_C}{dt} = \frac{CE_0}{\tau} \exp\left(-\frac{t}{\tau}\right) = \frac{E_0}{R} \exp\left(-\frac{t}{\tau}\right).$$

La puissance délivrée par le générateur est ei , donc l'énergie que délivre celui-ci entre l'instant initial, $t = 0$, et l'instant final $t \rightarrow \infty$, est :

$$\mathcal{E}_{\text{gén}} = \int_0^\infty E_0 i dt = \int_0^\infty \frac{E_0^2}{R} \exp\left(-\frac{t}{\tau}\right) dt = \frac{E_0^2}{R} \left[-\tau \exp\left(-\frac{t}{\tau}\right)\right]_0^\infty = \frac{\tau E_0^2}{R} = CE_0^2.$$

L'énergie électrique stockée par le condensateur est :

$$\mathcal{E}_{\text{élec}} = \int_0^\infty \frac{d}{dt} \left(\frac{1}{2} Cu_C^2 \right) dt = \frac{C}{2} [u_C^2]_0^\infty = \frac{CE_0^2}{2},$$

car $u_C(0) = 0$ et $u_C(\infty) = E_0$.

L'énergie dissipée par effet Joule dans la résistance est :

$$\mathcal{E}_{\text{Joule}} = \int_0^\infty Ri^2 dt = \frac{E_0^2}{R} \left[-\frac{\tau}{2} \exp\left(-2\frac{t}{\tau}\right)\right]_0^\infty = \frac{CE_0^2}{2}.$$

On retrouve le bilan :

$$\mathcal{E}_{\text{gén}} = \mathcal{E}_{\text{élec}} + \mathcal{E}_{\text{Joule}}.$$

Ainsi, quelles que soient les valeurs de R et de C , l'énergie délivrée par le générateur se répartit à parts égales entre l'énergie stockée par le condensateur et celle dissipée par effet Joule.

3 Réalisation concrète de l'échelon

3.1 Réponse à un signal créneau

Le GBF fournit une tension créneau, T -périodique, qui vaut E_0 sur une demi-période et 0 sur l'autre. On observe simultanément e et u_C sur la figure 8.6.

Pour $t \in \left[0, \frac{T}{2}\right]$, $e(t) = E_0$ et l'on observe la réponse indicielle. On laisse le temps au système de réagir afin de parvenir en régime permanent. C'est le cas dès que :

$$T \gg \tau$$

Puis, sur $t \in \left[\frac{T}{2}, T\right]$, e passe à zéro. On laisse le système évoluer sans entrée, ce qui correspond au **régime libre**.

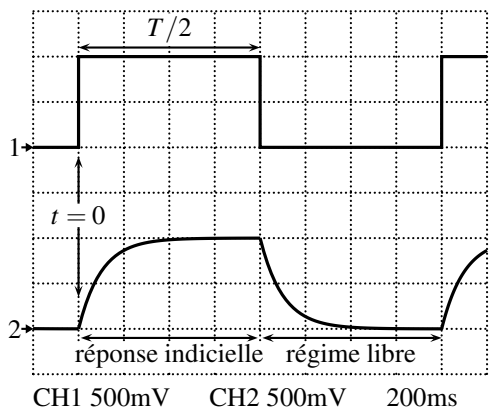


Figure 8.6 – Formes des signaux d'entrée et de sortie.

3.2 Modélisation du régime libre

L'équation différentielle qui régit l'évolution du système est toujours valable, avec $e = 0$:

$$\tau \frac{du_C}{dt} + u_C = 0.$$

Cette équation différentielle est homogène. Sa solution, déjà calculée précédemment, est :

$$u_C(t) = \lambda \exp\left(-\frac{t}{\tau}\right).$$

La constante d'intégration λ se calcule avec la condition initiale de ce régime, c'est à dire la valeur de u_C en $t = \frac{T}{2}$. Elle correspond à la fin de la réponse indicielle précédemment étudiée :

$$u_C\left(\frac{T}{2}\right) \simeq E_0 \quad (\text{ici } 1 \text{ V}).$$

Alors $u_C\left(\frac{T}{2}\right) = E_0 = \lambda \exp\left(-\frac{T}{2\tau}\right)$, donc $\lambda = E_0 \exp\left(\frac{T}{2\tau}\right)$. Finalement :

$$u_C(t) = E_0 \exp\left(\frac{T}{2\tau}\right) \exp\left(-\frac{t}{\tau}\right) = E_0 \exp\left(-\frac{t - \frac{T}{2}}{\tau}\right).$$

Le système réagit encore avec une constante de temps τ . Il atteint le régime permanent au bout de $4,6\tau$. En effet, en $t = \frac{T}{2} + 4,6\tau$ (instant origine + temps de réaction), le signal a perdu 99% de sa valeur initiale : $u_C\left(\frac{T}{2} + 4,6\tau\right) = E_0 \exp(-4,6) = E_0 \times 1,0.10^{-2}$.

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

Le **temps de réponse** du système, c'est à dire la durée du régime transitoire, est conventionnellement de $4,6\tau$, comme pour l'essai indiciel. C'est une caractéristique d'un système du premier ordre, quel que soit le signal d'entrée.

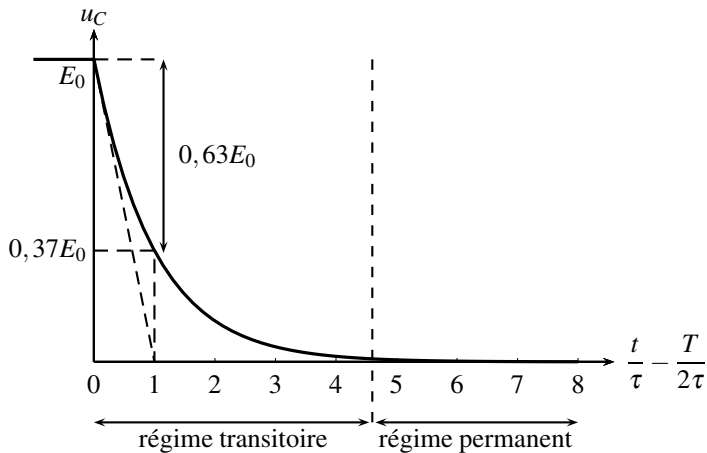


Figure 8.7 – Régime libre d'un circuit RC.

On retrouve la méthode d'**identification** de la valeur de τ :

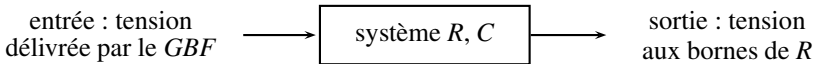
- on repère la valeur initiale de u_C , ici E_0 ,
- on note la date à laquelle u_C a perdu 0,63% de cette valeur : c'est τ .

La valeur de τ peut théoriquement être trouvée avec la tangente à l'origine à la courbe, en pointillés sur le graphe, de la même manière que plus haut.

4 Étude de la tension $u_R(t)$

4.1 Montage étudié et observations expérimentales

On change de point de vue pour étudier la tension aux bornes de la résistance, dans le montage du paragraphe 1.1 :



On observe simultanément la tension aux bornes de la résistance, u_R , sur la voie 2, et la tension e délivrée par le GBF, sur la voie 1. On peut encore définir :

- le **régime établi**, ou **régime permanent**, pour lequel la sortie est de la même forme que l'entrée, ici une constante, mais de valeur nulle,
- le **régime transitoire**, entre l'instant initial et le régime permanent.

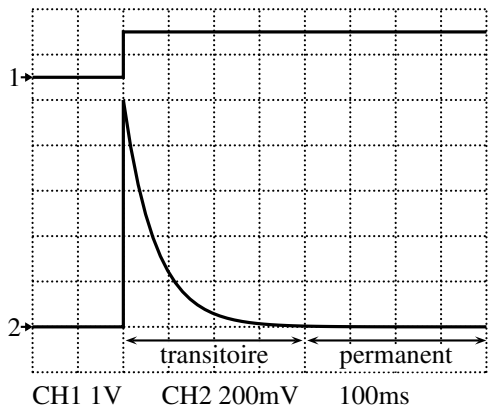


Figure 8.8 – Tensions d’entrée et aux bornes de la résistance.

4.2 Équation différentielle sur $u_R(t)$

L’entrée du montage est e , sa sortie est u_R . Pour $t < 0$, l’entrée est nulle et u_R aussi. On cherche, pour $t > 0$, un couple {équation différentielle sur u_R , condition initiale}.

Pour établir l’équation différentielle, on applique la loi des mailles :

$$e = 0 + u_R + u_C.$$

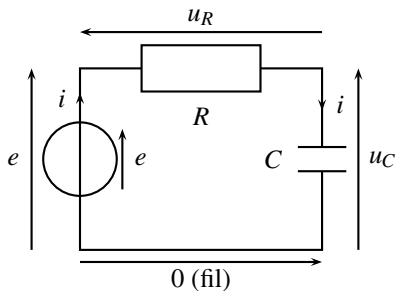


Figure 8.9 – Montage étudié.

Le même courant d’intensité i traverse le GBF, la résistance et le condensateur :

$$i = C \frac{du_C}{dt} \quad \text{et} \quad u_R = Ri.$$

On dérive la loi des mailles par rapport au temps :

$$\frac{de}{dt} = \frac{du_R}{dt} + \frac{du_C}{dt} = \frac{du_R}{dt} + \frac{i}{C},$$

Ainsi :

$$\frac{de}{dt} = \frac{du_R}{dt} + \frac{u_R}{RC}.$$

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

Comme $e(t)$ est constante pour $t > 0$, sa dérivée est nulle. Ainsi, avec $RC = \tau$, le système obéit à l'équation différentielle homogène, c'est-à-dire sans second membre :

$$\tau \frac{du_R}{dt} + u_R = 0.$$

4.3 Portrait de phase

On trace $\frac{du_R}{dt}$ en fonction de u_R , afin de prévoir l'évolution du montage avant toute résolution de l'équation différentielle.

Avec $y = du_R/dt$ et $x = u_R$, l'équation différentielle devient $\tau y + x = 0$. Le portrait de phase est donc une droite d'équation :

$$y = -\frac{x}{\tau}.$$

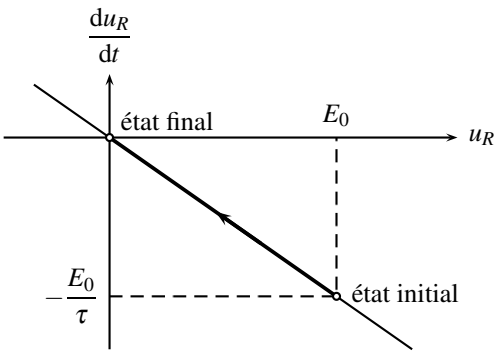


Figure 8.10 – Portrait de phase.

Où est l'état initial ? La tension u_R vaut, à l'origine, E_0 . On peut donc placer la point origine. Dans quel sens est parcourue la courbe ? Tant que le système évolue, sa dérivée est différente de zéro ; lorsque qu'elle atteint 0, le système ne varie plus. Ainsi le segment est parcouru de la droite vers la gauche :

- u_R décroît de E_0 à 0,
- $\frac{du_R}{dt}$ décroît de $-\frac{E_0}{\tau}$ à 0, qui marque l'état final.



Les systèmes pour lesquels le portrait de phase est calculable, décrit par une courbe dont l'équation $y(x)$ est analytiquement connue, demeurent des exceptions. Dans l'immense majorité des cas, on se contente de résolutions informatiques, comme dans le chapitre suivant.

4.4 Solution de l'équation différentielle

La solution générale de $\tau \frac{du_R}{dt} + u_R = 0$ est :

$$u_R(t) = \mu \exp\left(-\frac{t}{\tau}\right).$$

où μ est une constante d'intégration que l'on calcule en utilisant la condition initiale. Celle-ci est imposée par la continuité de la tension aux bornes du condensateur en $t = 0$:

$$u_C(0^-) = u_C(0^+).$$

D'après la loi des mailles, $u_C = e - u_R$, donc :

$$\underbrace{e(0^-)}_0 - \underbrace{u_R(0^-)}_0 = \underbrace{e(0^+)}_{E_0} - u_R(0^+) \quad \text{soit} \quad u_R(0^+) = E_0.$$

On observe bien que si le potentiel d'une borne du condensateur varie brutalement de E_0 , alors le potentiel de l'autre borne varie de la même valeur. Dès lors : $u_R(0^+) = E_0 = \mu$. Finalement :

$$u_R(t) = E_0 \exp\left(-\frac{t}{\tau}\right).$$

4.5 Régime libre

On observe successivement l'essai indiciel puis le régime libre en alimentant le circuit par un créneau, valant alternativement 0 et E_0 , de période T telle que :

$$T \gg \tau,$$

afin de laisser le temps au système d'atteindre, à la fin de chaque demi-période, le régime permanent.

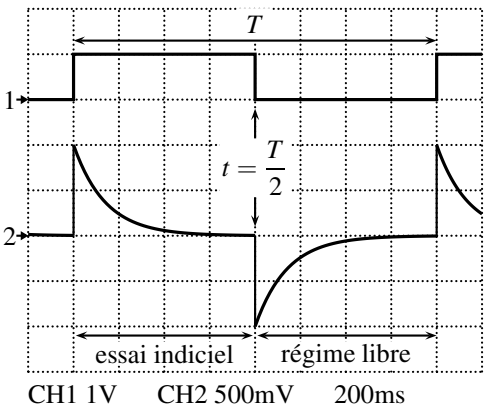


Figure 8.11 – Essai indiciel et régime libre aux bornes de la résistance.

L'équation différentielle qui régit l'évolution du système est, indépendamment de $e(t)$:

$$\tau \frac{du_C}{dt} + u_C = 0,$$

pour laquelle la solution est :

$$u_C(t) = \lambda \exp\left(-\frac{t}{\tau}\right).$$

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

La valeur de λ est trouvée avec la **condition initiale** en $t = \frac{T}{2}$.

La tensions aux bornes du condensateur est continue :

$$u_C\left(\frac{T^-}{2}\right) = u_C\left(\frac{T^+}{2}\right).$$

D’après la loi des mailles, $u_C = e - u_R$:

$$\underbrace{e\left(\frac{T^-}{2}\right)}_{E_0} - \underbrace{u_R\left(\frac{T^-}{2}\right)}_0 = \underbrace{e\left(\frac{T^+}{2}\right)}_0 - \underbrace{u_R\left(\frac{T^+}{2}\right)}_0.$$

Ainsi $u_R\left(\frac{T^+}{2}\right) = -E_0$ et :

$$u_R\left(\frac{T^+}{2}\right) = -E_0 = \lambda.$$

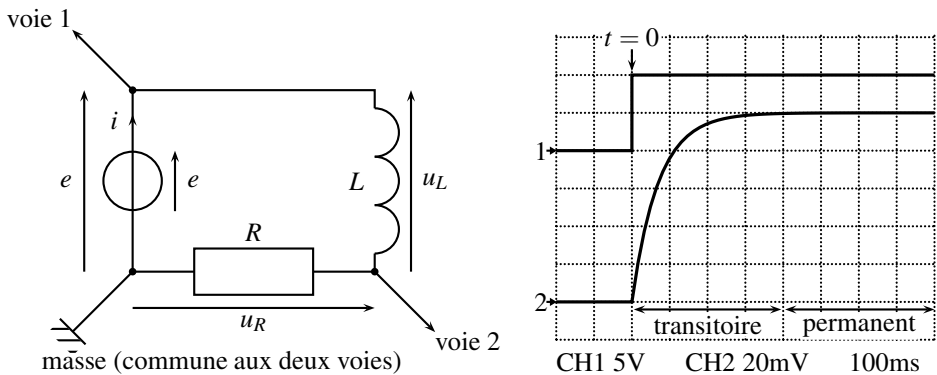
Finalement :

$$u_R(t) = -E_0 \exp\left(-\frac{t - \frac{T}{2}}{\tau}\right).$$

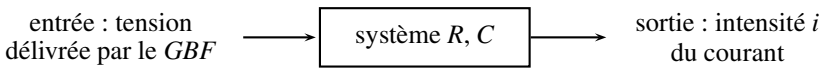
5 Exemple de circuit inductif

5.1 Schéma du montage

Le montage de la figure 8.12 suivante permet de visualiser en même temps à l’oscilloscope la tension e délivrée par le GBF et l’intensité i du courant, *via* la tension $u_R = Ri$ aux bornes de la résistance.



On s’intéresse au courant $i(t)$ qui circule dans le montage alimenté par e . L’entrée du système est e , sa sortie est i :



On observe expérimentalement la réponse indicielle sur la figure 8.12.

5.2 Équation différentielle sur $i(t)$

L'entrée du montage est e , sa sortie est u_L . On cherche un couple {équation différentielle sur u_L , condition initiale} pour $t > 0$. Pour $t < 0$, l'entrée est nulle et u_L aussi.

Pour établir l'équation différentielle, appliquons la loi des mailles :

$$e = 0 + u_L + u_R.$$

Le même courant d'intensité i traverse le GBF, la résistance et le condensateur :

$$u_L = L \frac{di}{dt} \quad \text{et} \quad u_R = Ri.$$

Ainsi :

$$e = Ri + L \frac{di}{dt}.$$

Soit :

$$\frac{L}{R} \frac{di}{dt} + i = \frac{e}{R}.$$

C'est une équation différentielle du premier ordre. Ce montage est donc un système du premier ordre, pour lequel on peut définir une constante de temps τ :

$$\tau = \frac{L}{R},$$

afin d'écrire l'équation différentielle :

$$\tau \frac{di}{dt} + i = \frac{e}{R}.$$

5.3 Homogénéité de l'équation différentielle

Les trois termes s'expriment en ampères A :

$$[i] = \text{A} \quad \text{et} \quad \left[\frac{e}{R} \right] = \frac{\text{V}}{\Omega} = \text{A}.$$

Or la dérivée di/dt est en ampères divisé par des secondes A.m^{-1} . L'ensemble est homogène si τ est en secondes :

$$[\tau] = \text{s}.$$

Le quotient $\frac{L}{R}$ est en seconde. Il représente la constante de temps du circuit.

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

Remarque

On retrouve ce résultat avec les unités des lois électrocinétiques aux bornes d'une résistance et d'une bobine :

$$\begin{cases} u_R = Ri & \text{donc } V = \Omega \times A \\ u_L = L \frac{di}{dt} & \text{donc } V = H \times \frac{A}{s} \end{cases} \text{ implique } \frac{H}{\Omega} = s.$$

5.4 Solution de l'équation différentielle

Le montage n'est pas alimenté pour $t < 0$, $e = 0$ et $i = 0$ de même.

Pour $t > 0$, $e = E_0$. Le montage est décrit par l'équation différentielle :

$$\tau \frac{di}{dt} + i = \frac{E_0}{R}.$$

Cette équation différentielle se résoud comme précédemment. Elle est la somme :

- d'une **solution particulière** i_p , cherchée sous la même forme que l'entrée, ici une constante. Avec $i_p = \text{constante}$, $\frac{di_p}{dt} = 0$ et l'équation devient $(\tau \times 0) + i_p = \frac{E_0}{R}$.
- d'une **solution homogène** $i_h(t)$, solution de l'équation homogène $\tau \frac{di_h}{dt} + i_h = 0$:

$$i_h(t) = \mu \exp\left(-\frac{t}{\tau}\right).$$

La solution générale de l'équation différentielle du circuit est donc :

$$i(t) = i_p(t) + i_h(t) = \frac{E_0}{R} + \mu \exp\left(-\frac{t}{\tau}\right).$$

La constante μ est déterminée par la **condition initiale**. Elle est imposée par la continuité du courant qui traverse la bobine d'inductance L .

L'énergie $Li^2/2$ contenue dans la bobine est continue, elle ne peut pas varier instantanément. L'intensité i dans une bobine est donc, elle aussi, continue

Ainsi $\underbrace{i(0^-)}_0 = i(0^+)$. D'où : $i(0) = \frac{E_0}{R} + \mu = 0$ donc $\mu = -\frac{E_0}{R}$.

Finalement :

$$i(t) = \frac{E_0}{R} \left(1 - \exp\left(-\frac{t}{\tau}\right)\right).$$

Cette solution est mathématiquement analogue à celle obtenue lors de l'étude de u_C pour le circuit RC .

5.5 Bilan de puissance

Multiplions la loi des mailles par i :

$$e \times i = u_L \times i + u_R \times i \quad \text{soir} \quad ei = Li \frac{di}{dt} + Ri^2$$

Or : $Li \frac{di}{dt} = \frac{d}{dt} \left(\frac{1}{2} Li^2 \right)$. Ainsi :

$$ei = \frac{d}{dt} \left(\frac{1}{2} Li^2 \right) + Ri^2.$$

On reconnaît :

- la puissance ei délivrée par le GBF,
- la dérivée de l'énergie $\mathcal{E}_{\text{magné}} = \frac{1}{2} Li^2$ emmagasinée dans la bobine d'inductance L ,
- la puissance Ri^2 dissipée par effet Joule dans la résistance.

La puissance délivrée par le générateur se répartit entre une puissance magnétique stockée par la bobine et une puissance perdue par effet Joule.

6 Généralisation : systèmes du premier ordre

Tous les systèmes du premier ordre sont décrits par une équation différentielle canonique :

$$\tau \frac{ds}{dt} + s = s_{\infty},$$

où s_{∞} est imposé par l'entrée, τ est la **constante de temps** du système et s est la **sortie** du système.

Le **temps de réponse** d'un système du premier ordre, c'est à dire la durée du régime transitoire, ou temps de réponse à 99%, est de $4,6 \tau$.

7 Régime permanent (PTSI)

Chaque circuit étudié dans ce chapitre a été modélisé par une équation différentielle, dont la solution est la somme :

- d'une **solution particulière**, cherchée sous la même forme que l'entrée,
- et d'une **solution homogène**, solution de l'équation homogène.

Dans le cas d'un système du premier ordre, la solution homogène est systématiquement sous la forme d'une exponentielle décroissante du temps, $\exp\left(-\frac{t}{\tau}\right)$, où τ est la constante de temps du système. La solution homogène est négligeable devant la solution particulière au bout d'une durée de $4,6 \tau$. Ainsi :

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

Le régime permanent est entièrement décrit par la solution particulière.

Dans le cas de la réponse à un échelon, cette solution particulière est cherchée sous la même forme que l'entrée, c'est-à-dire constante. Pour des signaux constants, les liens entre l'intensité du courant qui traverse un condensateur ou une bobine, et la tension à leurs bornes se simplifient. En effet, $u_L = L \frac{di_L}{dt} = 0$ dès que i_L est constant ; $i_C = C \frac{du_C}{dt} = 0$ dès que u_C est constant. Ainsi, en régime permanent, la tension aux bornes d'une bobine est nulle, alors que le courant qui traverse un condensateur est nul.

Pour un système soumis à un échelon, on en déduit les équivalents en régime permanent, qui assurent $u_L = 0$ et $i_C = 0$:

Lors de la réponse à un échelon, une bobine est équivalente, lors du régime permanent, à un fil.

Lors de la réponse à un échelon, un condensateur est équivalent, lors du régime permanent, à un interrupteur ouvert.

Par exemple, pour le circuit de la figure 8.3, remplacer le condensateur par un interrupteur ouvert impose $i = 0$. Alors $u_R = Ri = 0$ et $u_C = e = E_0$, qui représente bien la solution en régime permanent.

Dans le circuit de la figure 8.12, remplacer la bobine par un fil impose $u_L = 0$. Ainsi la loi des mailles mène à $e = u_R = Ri$ et donc $i = \frac{e}{R} = \frac{E_0}{R}$, qui est la solution en régime permanent.

SYNTHÈSE

SAVOIRS

- distinguer le régime transitoire du régime permanent ou établi
- continuité de u aux bornes de C
- continuité de i à travers L
- définition d'un échelon et de la réponse indicielle
- définition d'une réponse libre
- définition d'un portrait de phase
- équation différentielle canonique et constante de temps
- temps de réponse d'un système du premier ordre

SAVOIR-FAIRE

- établir l'équation différentielle du circuit
- résoudre une équation différentielle du premier ordre
- prévoir l'évolution d'un système avec son portrait de phase
- réaliser un bilan de puissance
- déterminer les grandeurs électriques en régime permanent en remplaçant les bobines et les condensateurs par des interrupteurs fermés ou ouverts (PTSI)

MOTS-CLÉS

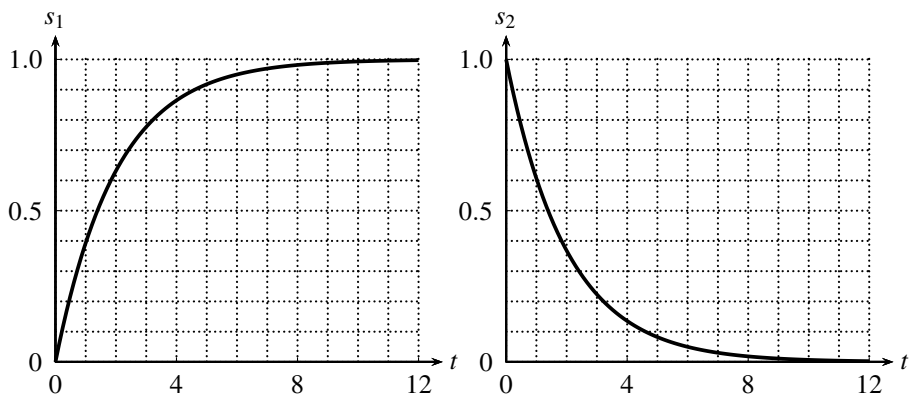
- | | | |
|------------------------------|-------------------------|---------------------|
| • régime transitoire | cielle | système |
| • régime permanent ou établi | • réponse libre | • portrait de phase |
| | • constante de temps | |
| • échelon et réponse indi- | • temps de réponse d'un | |

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

S'ENTRAÎNER

8.1 Constantes de temps (★)

On soumet deux systèmes, 1 et 2, à une entrée en échelon. Les sorties sont ci-après. Quelles sont les constantes de temps des deux systèmes ?



8.2 Réponse d'un système (★)

Un système linéaire obéit à l'équation différentielle : $2 \frac{df}{dt} + 8f(t) = 2e(t)$.

1. Quelle est la constante de temps du système ?
2. À l'état initial, $f(0) = 0$ et $e(t)$ passe alors de 0 à 5 V puis reste constant. Que vaut la sortie, vers quelle valeur converge-t-elle ?

8.3 Résistance de fuite d'un condensateur (★)

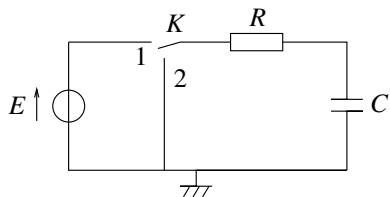
Soit un condensateur de capacité C et présentant une résistance de fuite R_f . On peut modéliser le condensateur par l'association en parallèle de R_f et C . On mesure la tension aux bornes du condensateur à l'aide d'un voltmètre électronique parfait (de résistance interne infinie). Ce condensateur ayant été chargé sous une tension E à l'aide d'une source idéale de tension, on ouvre le circuit. Au bout d'un temps T , on constate que la tension indiquée par le voltmètre n'est plus que de $E' < E$.

1. Comment peut-on expliquer ces observations ?
2. Donner l'expression de R_f en fonction de C , E , E' et T .

APPROFONDIR

8.4 Étude expérimentale d'un circuit RC (★)

On considère le circuit suivant comprenant une résistance de valeur R , un condensateur de capacité C et une alimentation stabilisée de tension à vide E :



- 1. À l'instant $t' = 0$, on place l'interrupteur K en position 1, le condensateur est déchargé. Décrire ce qui se passe.
- 2. On suppose que le régime permanent a été atteint. On place l'interrupteur K en position 2 à $t = 0$. Quelle est la nature de la courbe représentant $\ln\left(\frac{E}{u_C}\right)$ en fonction du temps ?
- 3. On réalise l'expérience avec les valeurs $R = 10\text{ M}\Omega$, $C = 10\text{ }\mu\text{F}$ et $E = 10\text{ V}$. On branche un voltmètre numérique (calibre 20 V) aux bornes du condensateur et on étudie la décharge du condensateur à partir de l'instant de date $t = 0$ où on place l'interrupteur en position 2. On relève les valeurs suivantes :

$t\text{ (s)}$	5	10	15	20	30	45	60	90	120	150
$u_C\text{ (V)}$	9,05	8,19	7,41	6,70	5,49	4,07	3,01	1,65	0,91	0,50

L'expérience est-elle en accord avec la théorie ?

- 4. Si oui, justifier. Si non, pour quel composant a-t-on utilisé un modèle inadapté aux caractéristiques du circuit ? On indiquera alors quel serait le modèle adapté et les valeurs des caractéristiques de ce dernier compte tenu des données expérimentales.
- 5. Quelle aurait été la condition à vérifier pour avoir accord parfait entre théorie et expérience ?
- 6. On souhaite visualiser sur un oscilloscope en mode DC les courbes théoriques de charge et de décharge du condensateur. On choisit comme valeurs $R = 10\text{ k}\Omega$ et $C = 10\text{ nF}$. L'alimentation stabilisée est remplacée par un générateur basse fréquence (GBF) et on place l'oscilloscope aux bornes du condensateur. Rappeler le modèle électrique du GBF et donner les valeurs des éléments du modèle.
- 7. Même question pour l'oscilloscope.
- 8. Dessiner la schéma en tenant compte des paramètres du GBF et de l'oscilloscope.
- 9. Déterminer les conditions qui permettent de n'avoir pas à tenir compte des paramètres du GBF et de l'oscilloscope. En déduire la justification du choix des valeurs employées pour R et C .
- 10. Quel signal du GBF doit-on choisir à la sortie du GBF pour observer sur l'oscilloscope la charge et la décharge du condensateur ?

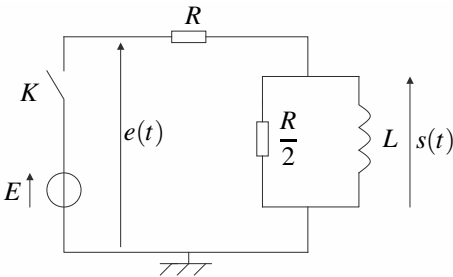
CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

11. Comment doit-on choisir la fréquence de ce signal pour observer pratiquement toute la charge puis de la décharge du condensateur ? Dans ce cas, représenter l'allure du signal obtenu.

8.5 Étude d'un circuit RL (★)

On considère dans le circuit représenté ci-dessous. A l'instant $t = 0$, on ferme l'interrupteur K qui était ouvert depuis très longtemps.

1. L'intensité du courant i traversant la résistance R est-elle continue en $t = 0$? Si non, donner ses valeurs en 0^- et en 0^+ .



2. Mêmes questions pour la tension s .
3. Que vaut $s(t)$ lorsque t tend vers l'infini ?
4. Établir l'équation différentielle vérifiée par $s(t)$.
5. En déduire l'expression de $s(t)$.
6. Tracer l'allure de $s(t)$.
7. Exprimer en fonction de L et de R le temps t_0 au bout duquel la tension s a été divisée par 10.
8. Proposer une méthode expérimentale pour déterminer t_0 à l'aide d'un oscilloscope. On précisera le montage à utiliser et la méthode de la mesure pratique.
9. On mesure $t_0 = 3,0 \mu s$ pour $R = 1000 \Omega$. En déduire la valeur de L . On donne $\frac{1}{\ln 10} = 0,43$.
10. On remplace le générateur idéal de tension par un générateur délivrant un créneau de période T . Quel est l'ordre de grandeur de la fréquence à utiliser pour pouvoir effectivement mesurer t_0 par la méthode proposée ci-dessus ?

8.6 Circuit à condensateur (★)

On considère un circuit composé d'une résistance R et d'un condensateur de capacité C en série aux bornes duquel on place un générateur de tension idéal de force électromotrice constante E et un interrupteur K . Initialement le circuit est ouvert et le condensateur est déchargé. On note u_C la tension aux bornes du condensateur. A l'instant $t = 0$, on ferme l'interrupteur K .

1. Déterminer la valeur de l'intensité du courant parcourant le circuit au bout d'un temps infini.
2. Même question pour la tension u_C .
3. En déduire le comportement du condensateur dans ces conditions.
4. Établir l'équation différentielle vérifiée par u_C .
5. On pose $\tau = RC$. Quel nom donne-t-on à cette constante ? Quelle est son unité dans le système international ?
6. Établir l'expression de $u_C(t)$ et retrouver la valeur limite pour les temps longs.
7. Tracer l'allure de $u_C(t)$ en précisant les éventuelles asymptotes.
8. Donner l'équation de la tangente à l'origine.
9. Calculer les coordonnées du point d'intersection de cette tangente avec l'asymptote.
10. Déterminer, en fonction de τ , l'expression du temps t_1 à partir duquel la charge du condensateur diffère de moins de 1 % de sa charge finale.
11. Donner l'expression de l'intensité parcourant le circuit en fonction du temps.
12. On mène maintenant l'étude énergétique de la charge du condensateur. Exprimer E_C l'énergie totale emmagasinée par le condensateur lors de sa charge en fonction de C et E .
13. Même question pour E_J l'énergie dissipée par effet Joule dans la résistance lors de cette charge.
14. Même question pour E_G l'énergie fournie par le générateur lors de cette charge.
15. Donner une relation liant E_C , E_J et E_G et proposer une interprétation physique de cette relation.
16. On définit le rendement énergétique de la charge du condensateur par $\rho = \frac{E_C}{E_G}$. Donner sa valeur ici.
17. Pour améliorer ce rendement, on effectue la charge en deux phases : la première en chargeant le condensateur avec un générateur de force électromotrice $\frac{E}{2}$ puis en basculant sur un générateur de force électromotrice E une fois que C a atteint son maximum de charge sous une tension $\frac{E}{2}$.
Déterminer pour la première phase les expressions de E_{C1} et E_{G1} des énergies respectivement emmagasinée par C et fournie par le générateur.
18. Établir l'expression de $u_C(t)$ lors de la seconde phase. On prendra pour nouvelle origine des temps l'instant de bascule entre les deux générateurs.
19. En déduire celle de $i(t)$.
20. Déterminer pour la seconde phase les expressions de E_{C2} et E_{G2} des énergies respectivement emmagasinée par C et fournie par le générateur.
21. En déduire le nouveau rendement énergétique.
22. Compte tenu des résultats obtenus, proposer une méthode pour faire tendre le rendement vers 1. On ne demande aucun calcul.

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

CORRIGÉS

8.1 Constantes de temps

La réponse indicielle vaut 63% de la valeur finale en $t = \tau$. On lit graphiquement pour le système 1 : $\tau = 2\text{s}$.

La réponse indicielle a perdu 63% de la valeur initiale en $t = \tau$. On lit graphiquement pour le système 2 : $\tau = 2\text{s}$.

8.2 Réponse d'un système

1. Il convient d'écrire l'équation différentielle sous sa forme canonique $\tau \frac{df}{dt} + 1 \times f(t) = \dots$ (le coefficient devant f doit être 1). Ici : $\frac{2}{8} \frac{df}{dt} + f(t) = \frac{2}{8} e(t)$ donc $\tau = \frac{2}{8} = 0,25\text{s}$.

2. La solution de l'équation différentielle est la somme de la solution homogène $f_h(t) = \mu \exp\left(-\frac{t}{\tau}\right)$ et de la solution particulière $f_p(t) = \frac{2}{8} \times 5 = 1,25\text{ V}$ (cherchée sous la même que l'entrée, constante). La constante d'intégration est trouvée avec la condition initiale en $t = 0$: $f(0) = \mu + 1,25 = 0$ donc $\mu = -1,25\text{ V}$. et $f(t) = 1,25 \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$ (V).

8.3 Résistance de fuite d'un condensateur

1. Lorsqu'on ouvre l'interrupteur, la capacité C se trouve uniquement relié à la résistance R_f dans laquelle elle se décharge. Ceci explique la diminution de la tension aux bornes du condensateur.

2. La décharge de C dans R_f vérifie l'équation différentielle suivante : $\frac{du}{dt} + \frac{u}{R_f C} = 0$ en notant u la tension aux bornes du condensateur. Compte tenu des conditions initiales : $u(0) = E$, la solution s'écrit : $u = E \exp\left(-\frac{t}{R_f C}\right)$. Or à $t = T$, $u = E'$ donc : $R_f = \frac{T}{C} \ln \frac{E}{E'}$.

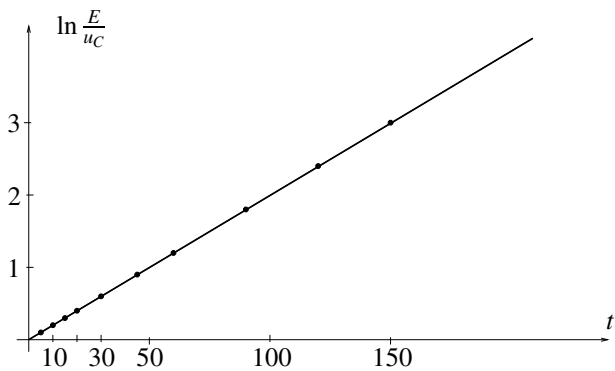
8.4 Étude expérimentale d'un circuit RC

1. On observe une charge du condensateur : une loi des mailles permet d'obtenir l'équation différentielle : $RC \frac{du_C}{dt} + u_C = E$ dont la solution est, compte tenu des conditions initiales, $u_C = E \left(1 - \exp\left(-\frac{t}{RC}\right)\right) = E \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$. La tension du condensateur tend vers E .

2. L'équation différentielle s'écrit maintenant : $RC \frac{du_C}{dt} + u_C = 0$. La solution est : $u_C = E \exp\left(-\frac{t}{\tau}\right)$. On en déduit : $\ln \frac{E}{u_C} = \frac{t}{\tau}$ On obtient donc une droite de pente $\frac{1}{\tau}$.

3. On obtient bien une droite comme l'annonce la théorie mais la pente vaut 0,02 au lieu de 0,01. L'écart est significatif.

t (s)	5	10	15	20	30	45	60	90	120	150
u_C (V)	9,05	8,19	7,41	6,70	5,49	4,07	3,01	1,65	0,91	0,50
$\ln \frac{E}{u_C}$	0,10	0,20	0,30	0,40	0,60	0,90	1,20	1,80	2,40	3,00



4. On n’a pas pris en compte les paramètres du voltmètre. Si on en tient compte, le nouveau montage est obtenu en plaçant en parallèle avec C la résistance R_V du voltmètre. On peut retrouver le montage initial en remplaçant R par $R_{eq} = \frac{RR_V}{R + R_V}$ et la pente p vaut alors $\frac{1}{R_{eq}C}$.

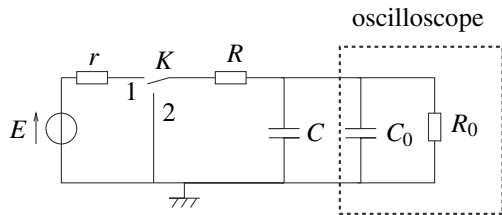
On peut en déduire la valeur de $R_V = \frac{R}{pRC - 1} = 10\text{ M}\Omega$. Ceci est tout à fait possible compte tenu des ordres de grandeur usuels des résistances des voltmètres.

5. L’effet de R_V aurait été négligeable si $R \ll R_V$ alors $R_{eq} \approx R$ et on retrouve les résultats théoriques.

6. Un GBF peut être modélisé par une source idéale de tension e en série avec une résistance interne r de $50\ \Omega$.

7. Pour l’oscilloscope, on distingue le mode DC où l’impédance d’entrée est constituée d’un condensateur de capacité C_0 de l’ordre de 50 pF en parallèle avec résistance R_0 de $1\text{ M}\Omega$ du mode AC où l’impédance d’entrée comporte en plus un condensateur de grande capacité en série avec l’impédance d’entrée du mode DC. Dans la suite de l’exercice, on utilise l’oscilloscope en mode DC.

8. On a donc le schéma suivant :

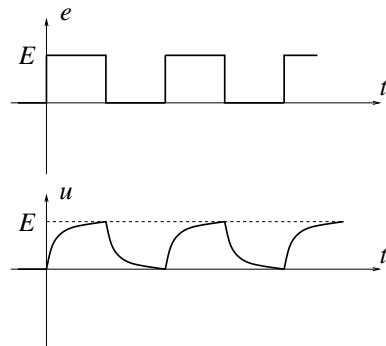


9. On vérifie avec les données numériques que $r = 50\ \Omega \ll R = 10\text{ k}\Omega \ll R_0 = 1\text{ M}\Omega$ et $C_0 = 50\text{ pF} \ll C = 10\text{ nF}$. On retrouve dans ce cas $R' = R$, $C' = C$ et $E' = E$.

10. On prend sur la sortie $50\ \Omega$ un signal créneau entre 0 V et $E\text{ V}$.

11. Pour observer complètement la décharge et la recharge, il faut que $T \gg \tau$ soit $f \ll 10\text{ kHz}$. On observe alors :

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE



8.5 Étude d'un circuit RL

1. L'écriture de la loi des mailles donne $e = Ri + s$ et celle de la loi des nœuds $i = i_R + i_L$. On la continuité de l'intensité traversant une inductance permet d'écrire la continuité de i_L soit $i_L(0^+) = 0$. Cette valeur implique dans la loi des nœuds que $i(0^+) = i_R(0^+)$. On obtient alors la tension s par la relation du pont diviseur de tension $s(0^+) = \frac{R}{2}i_R(0^+) = \frac{R}{2}i(0^+)$. Finalement en reportant dans la loi des mailles, on en déduit $E = Ri(0^+) + s(0^+)$ soit $i(0^+) = \frac{2E}{3R}$. Par conséquent, l'intensité i est discontinue puisque $i(0^-) = 0$.

La modélisation du système par sa fonction de transfert, alliée à l'utilisation du théorème de la valeur initiale, vu en cours de SI, permettra dans la suite de résoudre bien plus aisément ce genre de question.

2. D'après la question précédente, on a $s(0^+) = \frac{E}{3}$ qui est également discontinue puisque $s(0^-) = 0$.

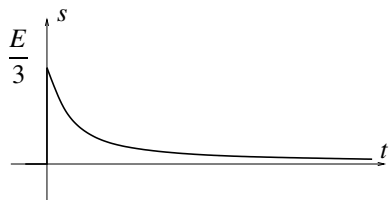
3. Lorsque le temps t tend vers l'infini, on obtient le régime permanent et les dérivées temporelles sont donc nulles. On obtient donc $s = L \frac{di}{dt} = 0$ pour un temps infini.

4. La loi des mailles $e = Ri + s$ avec $s = L \frac{di_L}{dt}$ et la loi des nœuds $i = i_L + i_R$ permet d'obtenir en dérivant $\frac{de}{dt} = R \frac{di_L}{dt} + R \frac{di_R}{dt} + \frac{ds}{dt} = R \frac{s}{L} + R \frac{d}{dt} \left(\frac{2s}{R} \right) + \frac{ds}{dt}$. L'équation différentielle demandée est donc $3 \frac{ds}{dt} + \frac{R}{L} s = \frac{de}{dt}$.

La modélisation du système par sa fonction de transfert permettra dans la suite de résoudre bien plus aisément ce genre de question.

5. On résout cette équation différentielle dont la solution est la somme de la solution de l'équation homogène associée $S \exp\left(-\frac{R}{3L}t\right)$ et d'une solution particulière $s = 0$. On détermine la constante d'intégration S en utilisant la condition initiale soit finalement $s = \frac{E}{3} \exp\left(-\frac{R}{3L}t\right)$.

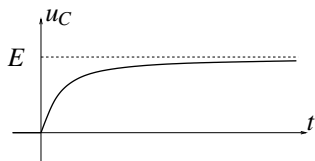
6. L'allure de $s(t)$ est donc la suivante :



7. Le temps t_0 cherché s'obtient en résolvant $s(t_0) = \frac{s(0^+)}{10}$ dont la solution est $t_0 = \frac{3L \ln 10}{R}$.
8. On branche l'oscilloscope entre la masse et le point reliant R, L et $\frac{R}{2}$. On utilise le décalibrage vertical de l'oscilloscope de manière à placer le maximum de s soit S sur la ligne 100 % et la valeur minimale de s soit 0,0 sur la ligne 0 % de l'écran de l'oscilloscope. On recherche ensuite la valeur de t_0 en déplaçant le curseur vertical jusqu'à ce qu'il soit sur l'intersection de la tension avec la ligne 10 %.
9. D'après ce qui précède, on déduit $L = \frac{Rt_0}{3 \ln 10} = 4,3 \cdot 10^{-4} \text{ H}$.
10. Pour mesurer t_0 , il faut laisser le temps au régime transitoire de se terminer, ce qui impose $\frac{T}{2} \gg t_0$ et en termes de fréquences $f \ll \frac{1}{2t_0} \simeq 10^5 \text{ Hz}$.

8.6 Circuit à condensateur

1. Pour un temps infini, on obtient le régime permanent qui est ici continu. Par conséquent, les dérivées temporelles sont nulles soit $i = 0$.
2. En écrivant la loi des mailles avec $i = 0$, on a $u_C = E$.
3. Dans ces conditions, la capacité C se comporte comme un interrupteur puisqu'il n'y a pas de courant.
4. En écrivant une loi des mailles ainsi que la relation $i = C \frac{du_C}{dt}$, on obtient $\frac{du_C}{dt} + \frac{u_C}{RC} = \frac{E}{RC}$.
5. La quantité τ est homogène à un temps dit caractéristique donc τ s'exprime en secondes.
6. En résolvant l'équation différentielle avec la condition initiale $u_C = 0$, on obtient $u_C = E \left(1 - \exp \left(-\frac{t}{RC} \right) \right)$ et $u_C = E$ pour un temps infini.
7. On a une asymptote $u = E$ et l'allure est la suivante :



8. L'équation de la tangente à l'origine s'écrit $u_C = \frac{du_C}{dt}(t=0)t + u_C(0) = \frac{E}{RC}t$.
9. L'intersection avec l'asymptote $u_C = E$ s'obtient pour $t = \tau$.
10. Il faut résoudre l'équation $u_C(t) = 0,99E$. On obtient $t_1 = RC \ln 100$.

CHAPITRE 8 – CIRCUIT LINÉAIRE DU PREMIER ORDRE

11. L'intensité s'obtient par la relation $i = C \frac{du_C}{dt} = \frac{E}{R} \exp\left(-\frac{t}{RC}\right)$.
12. L'énergie emmagasinée dans C s'écrit $E_C = \int_0^{+\infty} u_C(t)i(t)dt = \frac{CE^2}{2}$.
13. L'énergie dissipée par effet Joule est $E_J = \int_0^{+\infty} Ri^2(t)dt = \frac{CE^2}{2}$.
14. L'énergie fournie par le générateur vaut $E_G = \int_0^{+\infty} Ei(t)dt = CE^2$.
15. On constate que $E_G = E_C + E_J$, ce qui traduit la conservation de l'énergie.
16. On a $\rho = \frac{CE^2}{2CE^2} = 0,5$.
17. Par le même raisonnement que précédemment, on obtient $E_{G1} = \frac{CE^2}{4}$ et $E_{C1} = \frac{CE^2}{8}$.
18. Par la résolution de la même équation différentielle avec une nouvelle condition initiale $u_C = \frac{E}{2}$, on obtient $u_C = E \left(1 - \frac{\exp\left(-\frac{t}{RC}\right)}{2}\right)$.
19. On en déduit l'intensité $i = C \frac{du_C}{dt} = \frac{E}{2R} \exp\left(-\frac{t}{RC}\right)$.
20. Pour la seconde phase, on a $E_{G2} = \frac{CE^2}{2}$ et $E_{C2} = \frac{3CE^2}{8}$.
21. On en déduit $\rho' = 0,66$.
22. On pourrait faire tendre le rendement énergétique vers 1 en décomposant la charge en N charges de « pas » $\frac{E}{N}$.

Circuit linéaire du second ordre

9

Dans ce chapitre, on étudie la réponse temporelle de systèmes d'ordre deux, qu'ils soient électriques ou mécaniques. L'allure des courbes est ensuite analysée en fonction de la pulsation caractéristique et du coefficient d'amortissement du système.

1 Exemple expérimental

1.1 Montage

On cherche à observer, puis prévoir par le calcul, le comportement du montage ci-dessous lorsqu'il est alimenté par un échelon de tension délivré par un GBF. On observe simultanément les tensions délivrées par le GBF et aux bornes du condensateur sur les deux voies d'un oscilloscope. On néglige la résistance de sortie R_g du GBF devant la résistance variable R .

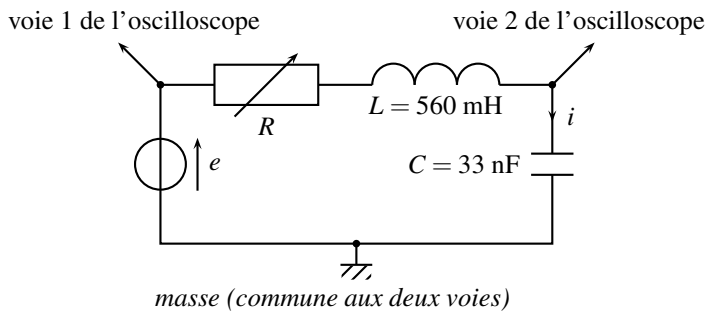


Figure 9.1 – Schéma du montage étudié.

1.2 Régimes transitoire puis permanent

On observe $e(t)$ sur la voie 1, qui passe de 0 à $E_0 = 600 \text{ mV}$, et la tension $u_C(t)$ aux bornes du condensateur sur la voie 2. Il existe plusieurs types de courbes lorsque R varie. Plus R augmente, plus le circuit est amorti, moins il oscille.

CHAPITRE 9 – CIRCUIT LINÉAIRE DU SECOND ORDRE

On observe dans chaque cas deux régimes :

- le **régime établi**, ou **régime permanent**, pour lequel la sortie est constante, comme l'entrée, égale à E_0 ,
- le **régime transitoire**, entre l'instant initial et le régime permanent. Il peut être oscillant ou non.

On observe que la durée du régime transitoire commence par diminuer avec la valeur de R , puis augmente après avoir atteint un minimum.

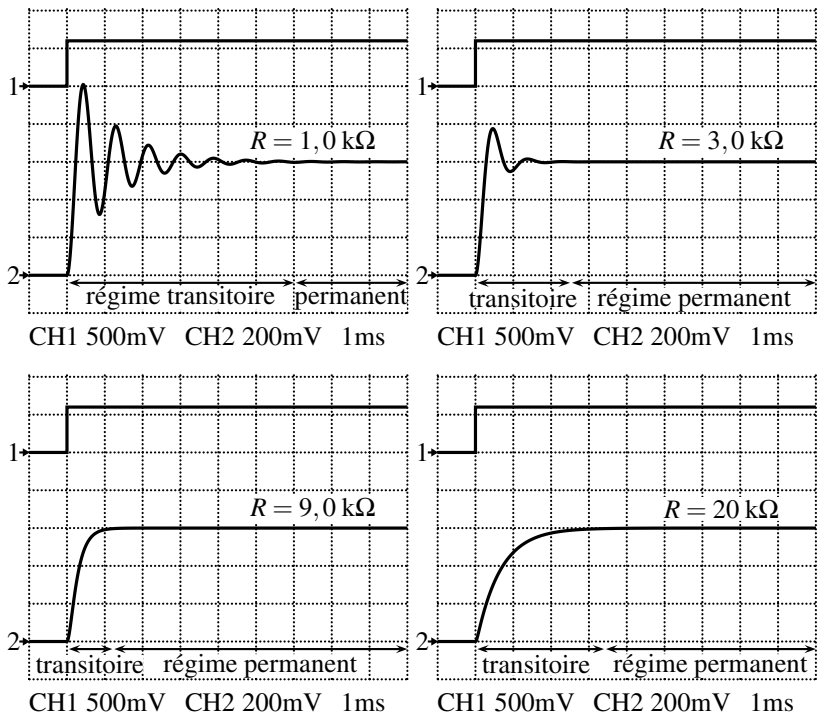


Figure 9.2 – Formes d'onde de l'essai indiciel en fonction de la valeur de R .

La courbe part avec une tangente nulle en l'origine, comme le montre le zoom pour $R = 3,0 \text{ k}\Omega$:

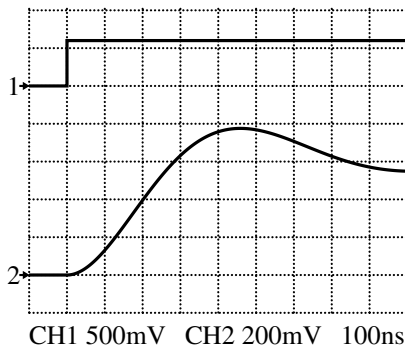


Figure 9.3 – Zoom d'une réponse indicielle en $t = 0$.

1.3 Portraits de phase

On peut étudier la dérivée temporelle de u_C (notée $\frac{du_C}{dt}$ ou $\dot{u}_C$), en fonction de u_C , pour chacun des quatre cas précédents. En notant E_0 la valeur de $e(t)$ pour $t > 0$, on trace alors $\frac{\dot{u}_C}{E_0}$ en fonction de $\frac{u_C}{E_0}$. L’allure des courbes est alors analogue quelque soit la valeur de E_0 .

L’état initial, connu avec les conditions initiales, est représenté par le point E_i de coordonnées $(u_C = 0, \dot{u}_C = 0)$ dans chaque cas. Dans l’état final, le condensateur est chargé et la tension $u_C(t)$ vaut E_0 et reste constante ; sa dérivée est alors nulle. Le portrait de phase converge vers le point E_f de coordonnées $(u_C = E_0, \dot{u}_C = 0)$.

Deux cas se présentent. Celui pour lequel la sortie est oscillante ; le portrait de phase converge alors vers le point E_f en tournant autour de lui (voir figure 9.4). Celui pour lequel la sortie n’oscille pas ; le portrait de phase converge directement vers E_f (voir figure 9.5).

Pour $R = 1,0\text{ k}\Omega$ et $R = 3,0\text{ k}\Omega$, le montage est oscillant. En effet, on observe que la dérivée est alternativement positive et négative. De plus, le portrait de phase tourne autour de la valeur finale $u_C = E_0$, $u_C(t)$ est donc alternativement supérieure et inférieure à E_0 , pour converger vers E_0 .

De plus, plus R augmente, plus le circuit est amorti, les valeurs extrémales de $\dot{u}_C$ diminuent, tout comme celles de u_C . Le système fait moins de tours autour de l’état final, le système oscille moins.

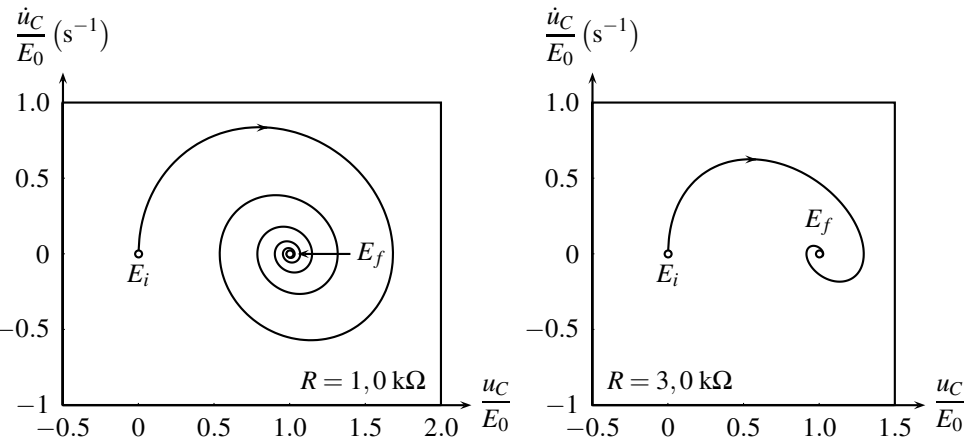


Figure 9.4 – Portraits de phase $\frac{\dot{u}_C}{E_0} = f\left(\frac{u_C}{E_0}\right)$ pour des systèmes pseudopériodiques.

Pour $R = 9,0\text{ k}\Omega$ et $R = 20\text{ k}\Omega$, le circuit n’est plus oscillant : le portrait de phase ne tourne pas autour de l’état final. En partant de l’état initial, on observe que $\dot{u}_C$ est positive, la tension u_C augmente. Puis $\dot{u}_C$ diminue pour arriver à 0, où le système ne bouge plus : il a atteint son état final.

De plus, la valeur maximale de $\dot{u}_C$ est de plus en plus faible quand R augmente. Le système est de plus en plus amorti.

CHAPITRE 9 – CIRCUIT LINÉAIRE DU SECOND ORDRE

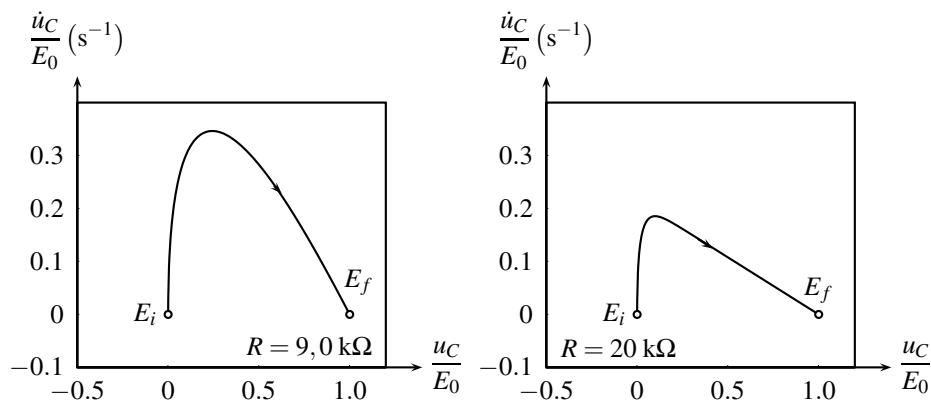


Figure 9.5 – Portraits de phase $\frac{\dot{u}_C}{E_0} = f\left(\frac{u_C}{E_0}\right)$ pour des systèmes apériodique ou critique.

Ces résultats expérimentaux sont modélisés dans le paragraphe suivant.

2 Équation différentielle sur la tension aux bornes du condensateur

2.1 Mise en équation

Le schéma électrique du circuit est le suivant :

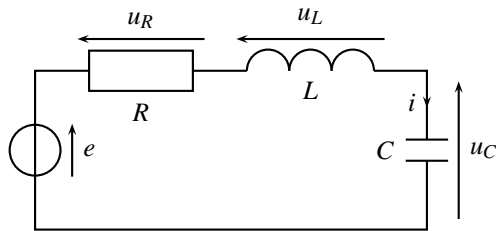


Figure 9.6 – Montage dont on étudie l'équation différentielle.

Fonctionnellement, le montage étudié est symboliquement représenté par :

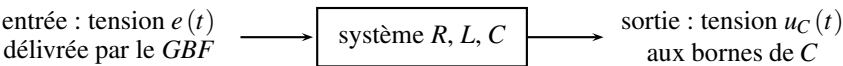


Figure 9.7 – Schéma-bloc du montage.

On cherche alors le lien entre $u_C(t)$ et $e(t)$. La loi des mailles impose :

$$e(t) = u_C(t) + u_L(t) + u_R(t).$$

De plus, R , L et C sont traversés par le même courant d'intensité $i(t)$:

$$u_R(t) = Ri(t), \quad u_L(t) = L \frac{di}{dt} \quad \text{et} \quad i(t) = C \frac{du_C}{dt}.$$

Ainsi :

$$u_R(t) = Ri(t) = RC \frac{du_C}{dt} \quad \text{et} \quad u_L(t) = L \frac{di}{dt} = LC \frac{d}{dt} \left(\frac{du_C}{dt} \right) = LC \frac{d^2 u_C}{dt^2}.$$

On remarque que c' est la dérivée seconde de $u_C(t)$ qui intervient dans $u_L(t)$.

La loi des mailles devient finalement :

$$e(t) = LC \frac{d^2 u_C}{dt^2} + RC \frac{du_C}{dt} + u_C(t).$$

Cette équation différentielle, du **deuxième ordre** (voir appendice mathématique), décrit l'évolution de la tension $u_C(t)$ dans le circuit.

2.2 Forme canonique de l'équation différentielle

Quelles sont les unités des différents termes de l'équation différentielle ?

$e(t)$ et $u_C(t)$ sont en volts :

$$[e(t)] = \text{V}, \quad [u_C(t)] = \text{V}.$$

Ainsi la dérivée $\frac{du_C}{dt}$ est en volts divisés par des secondes : $\left[\frac{du_C}{dt} \right] = \text{V.s}^{-1}$.

Et la dérivée seconde $\frac{d^2 u_C}{dt^2}$ est en volts divisés par des secondes au carré. C'est, en effet, une

dérivée temporelle (en s^{-1}) d'une dérivée première (en V.s^{-1}) : $\left[\frac{d^2 u_C}{dt^2} \right] = \text{V.s}^{-2}$.

Pour que l'équation différentielle soit homogène, c'est-à-dire que tous les termes sommés aient la même unité, le volt, il faut que :

$$[LC] = \text{s}^2 \quad \text{et} \quad [RC] = \text{s}.$$

Afin d'écrire toutes les équations différentielles du deuxième ordre de la même manière, on définit :

- la **pulsation caractéristique** ω_0 , qui s'exprime en rad.s^{-1} ,
- le **facteur d'amortissement** ξ , qui est sans dimension,
- la **facteur de qualité** $Q = \frac{1}{2\xi}$, qui est sans dimension.

Les équations différentielles du deuxième ordre sont alors écrites sous la forme canonique :

$$\frac{1}{\omega_0^2} \frac{d^2 u_C}{dt^2} + \frac{2\xi}{\omega_0} \frac{du_C}{dt} + u_C(t) = e(t).$$

CHAPITRE 9 – CIRCUIT LINÉAIRE DU SECOND ORDRE

ou bien, avec $Q = \frac{1}{2\xi}$:

$$\frac{1}{\omega_0^2} \frac{d^2 u_C}{dt^2} + \frac{1}{Q\omega_0} \frac{du_C}{dt} + u_C(t) = e(t).$$

Les valeurs de ω_0 et ξ du circuit étudié peuvent être identifiées en comparant l'équation différentielle et sa forme canonique :

$$\begin{cases} LC = \frac{1}{\omega_0^2} \\ RC = \frac{2\xi}{\omega_0} \end{cases} \quad \text{donc} \quad \begin{cases} \omega_0 = \frac{1}{\sqrt{LC}} \\ \xi = \frac{R}{2} \sqrt{\frac{C}{L}} \end{cases}$$



Le facteur de qualité vaut ici $Q = \frac{1}{2\xi} = \frac{1}{R} \sqrt{\frac{L}{C}}$.

2.3 Conditions initiales

On a vu dans le chapitre précédent que la résolution mathématique d'une équation différentielle du premier ordre $\tau \frac{ds}{dt} + s = s_\infty$ requiert une condition initiale sur le signal $s(t)$ cherché, en général la valeur $s(0)$ à l'instant initial. En effet, si l'équation différentielle pilote l'évolution du système, il convient d'expliquer de quel état il part.

De même, la résolution d'une équation différentielle du deuxième ordre requiert deux conditions initiales, comme dans le cas de l'oscillateur harmonique où l'on doit disposer de deux conditions initiales sur la position et la vitesse. Dans l'exemple du circuit RLC étudié, il s'agit des valeurs en $t = 0$ de $u_C(t)$ et de sa dérivée : $u_C(0)$ et $\left(\frac{du_C}{dt}\right)_{t=0}$.

Sur les oscillographes initiaux, il apparaît que $e(t)$ est nul pour $t < 0$, le circuit n'est pas alimenté, toutes les tensions sont, elles aussi, nulles, par exemple $u_C(t)$. Or la tension $u_C(t)$ aux bornes du condensateur de capacité C est continue ; elle était nulle pour $t < 0$, elle le reste en $t = 0$:

$$u_C(0) = 0.$$

La tension $u_R(t) = Ri(t)$, donc l'intensité $i(t)$ du courant, étaient nuls pour $t < 0$. Attendu que l'intensité du courant est continue dans une bobine d'inductance L , elle reste nulle en $t = 0$: $i(0) = 0$. Or $i(t) = C \frac{du_C}{dt}$ et ainsi : $\left(\frac{du_C}{dt}\right)_{t=0} = 0$.

Le couple {équation différentielle, conditions initiales} à résoudre est maintenant complet. Attendu que pour $t > 0$, $e(t) = E_0$:

$$\begin{cases} \frac{1}{\omega_0^2} \frac{d^2 u_C}{dt^2} + \frac{2\xi}{\omega_0} \frac{du_C}{dt} + u_C(t) = E_0 \\ u_C(0) = 0 \\ \left(\frac{du_C}{dt}\right)_{t=0} = 0. \end{cases}$$

3 Solution de l'équation différentielle

La solution générale de l'équation différentielle linéaire avec second membre non nul (voir annexe sur les outils mathématiques) est la somme d'une **solution particulière**, notée $u_{C,p}(t)$ (indice p pour « particulière ») et d'une **solution homogène**, notée $u_{C,h}(t)$ (indice h pour « homogène »).

3.1 Régime permanent ou établi

La solution particulière $u_{C,p}(t)$ est cherchée sous la forme d'une constante, comme l'entrée. Sa dérivée première est donc nulle, ainsi que sa dérivée seconde. L'équation différentielle devient donc $\frac{1}{\omega_0^2} \times 0 + \frac{2\xi}{\omega_0} \times 0 + u_{C,p}(t) = E_0$. Ainsi : $u_{C,p}(t) = E_0$.

En comparant aux graphes du paragraphe 1.2, où la sortie vaut E_0 en régime permanent, on voit que la solution particulière décrit le régime permanent ou établi. Ceci revêt un caractère général.

Le régime permanent ou établi est complètement décrit par la solution particulière.

3.2 Solutions homogènes

La solution homogène $u_{C,h}(t)$ est solution de l'équation homogène, c'est à dire avec un second membre nul :

$$\frac{1}{\omega_0^2} \frac{d^2 u_{C,h}}{dt^2} + \frac{2\xi}{\omega_0} \frac{du_{C,h}}{dt} + u_{C,h}(t) = 0.$$

Sa résolution passe l'équation caractéristique associée par (voir annexe sur les outils mathématiques page 1070) :

$$\frac{r^2}{\omega_0^2} + \frac{2\xi}{\omega_0} r + 1 = 0,$$

dont le discriminant Δ est $\Delta = \left(\frac{2\xi}{\omega_0}\right)^2 - 4 \times \frac{1}{\omega_0^2} \times 1 = \frac{4}{\omega_0^2} (\xi^2 - 1)$.

Le comportement physique du système est différent suivant le signe de Δ . On distingue trois cas :

a) Cas apériodique $\Delta > 0$

Ce cas correspond à $\xi^2 > 1$ ou, comme ξ est positif, $\xi > 1$ ou $Q < \frac{1}{2}$. Il correspond,

dans l'expérience du paragraphe 1.2, à $R > 2\sqrt{\frac{L}{C}} \simeq 8,2 \text{ k}\Omega$, c'est à dire aux deux dernières expériences. Le système est alors **fortement amorti**.

CHAPITRE 9 – CIRCUIT LINÉAIRE DU SECOND ORDRE

Les solutions de l'équation caractéristique associée sont alors :

$$\begin{aligned}r_1 &= \omega_0 \left(-\xi + \sqrt{\xi^2 - 1} \right), \\r_2 &= \omega_0 \left(-\xi - \sqrt{\xi^2 - 1} \right).\end{aligned}$$

Les solutions de l'équation différentielle sont dès lors, avec u_1 et u_2 deux constantes d'intégration, inconnues à ce stade de la résolution :

$$\begin{aligned}u_{C,h}(t) &= u_1 e^{r_1 t} + u_2 e^{r_2 t} \\&= u_1 \exp \left(\omega_0 t \left(-\xi + \sqrt{\xi^2 - 1} \right) \right) + u_2 \exp \left(\omega_0 t \left(-\xi - \sqrt{\xi^2 - 1} \right) \right) \\&= \exp(-\omega_0 \xi t) \left(u_1 \exp \left(\omega_0 \sqrt{\xi^2 - 1} t \right) + u_2 \exp \left(-\omega_0 \sqrt{\xi^2 - 1} t \right) \right).\end{aligned}$$

Ce cas est nommé **apériodique** car aucune des fonctions utilisées n'est périodique.

b) Cas critique $\Delta = 0$

Ce cas tire son nom de son caractère limite entre deux autres régimes. C'est le cas où $\xi = 1$

ou $Q = \frac{1}{2}$. Il correspond, dans l'expérience du paragraphe 1.2, à la valeur particulière

$R = 2\sqrt{\frac{L}{C}} \simeq 8,2 \text{ k}\Omega$. On parle d'**amortissement critique**.

Comme démontré dans l'annexe sur les outils mathématiques, la solution de l'équation homogène est alors :

$$u_{C,h}(t) = (u_1 + \mu t) \exp(-\omega_0 \xi t),$$

où u_1 et μ sont deux constantes d'intégration, inconnues à ce stade de la résolution.

c) Cas pseudopériodique $\Delta < 0$

C'est le cas où $\xi < 1$ ou $Q > \frac{1}{2}$, soit $R < 2\sqrt{\frac{L}{C}} \simeq 8,2 \text{ k}\Omega$. Le système est **faiblement amorti**, il correspond aux deux premières expériences du paragraphe 1.2.

Comme démontré dans l'annexe sur les outils mathématiques, la solution de l'équation homogène est alors :

$$u_{C,h}(t) = \exp(-\omega_0 \xi t) \left(u_1 \cos \left(\omega_0 \sqrt{1 - \xi^2} t \right) + u_2 \sin \left(\omega_0 \sqrt{1 - \xi^2} t \right) \right).$$

Des oscillations, dues aux fonctions trigonométriques de pulsation $\omega_0 \sqrt{1 - \xi^2}$, apparaissent dans ce cas, d'où le nom pseudopériodique. L'ensemble n'est toutefois pas périodique à cause de la présence du terme $\exp(-\omega_0 \xi t)$. On retiendra que la **pseudopulsation** est différente de ω_0 et vaut :

$$\omega = \omega_0 \sqrt{1 - \xi^2}.$$

3.3 Solution complète

La solution générale $u_C(t)$ est la somme des solutions particulière, $u_{C,p}(t)$, et homogène, $u_{C,h}(t)$:

$$u_C(t) = E_0 + u_{C,h}(t).$$



C'est sur cette solution complète qu'on identifie les constantes d'intégration par application des conditions initiales.

Il existe trois cas correspondants aux trois expressions différentes de $u_{C,h}(t)$ suivant le signe de Δ .

Exemple

À titre d'exemple, le cas où $\xi = 1$ est entièrement résolu :

$$u_C(t) = E_0 + (u_1 + \mu t) \exp(-\omega_0 \xi t).$$

Les conditions initiales imposent $u_C(0) = 0$ et $\left(\frac{du_C}{dt}\right)_{t=0} = 0$:

- $u_C(0) = 0 = E_0 + u_1$ donc : $u_1 = -E_0$;
- la dérivée de $u_C(t)$ par rapport à t est :

$$\frac{du_C}{dt}(t) = \mu \exp(-\omega_0 \xi t) + (u_1 + \mu t)(-\omega_0 \xi) \exp(-\omega_0 \xi t),$$

ainsi :

$$\left(\frac{du_C}{dt}\right)_{t=0} = 0 = \mu - \omega_0 \xi u_1,$$

d'où :

$$\mu = \omega_0 \xi E_0.$$

Finalement, la solution de l'ensemble {équation différentielle, conditions initiales}, dans le cas $\xi = 1$ est :

$$u_C(t) = E_0 + E_0(1 + \omega_0 \xi t) \exp(-\omega_0 \xi t).$$

Il est hors de question d'apprendre ce résultat par cœur. Seule la démarche doit être sue, ainsi que l'allure des graphes pour chaque cas, pour une tension d'entrée e en échelon.

On observe que les cas apériodique ou critique ont la même allure graphique. De plus :

Plus un système est amorti, c'est à dire plus son coefficient d'amortissement ξ est important, ou son facteur de qualité Q petit, moins il oscille.

Un oscillateur de grand facteur de qualité oscille beaucoup avant de s'arrêter.

CHAPITRE 9 – CIRCUIT LINÉAIRE DU SECOND ORDRE

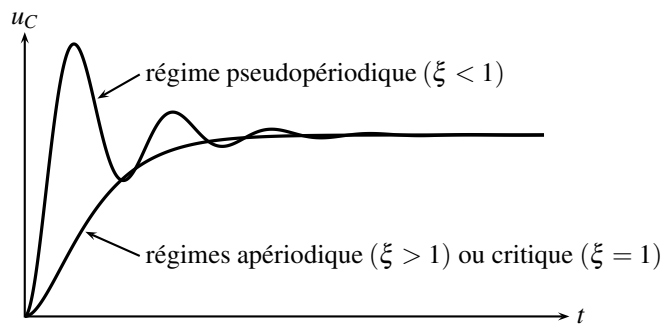


Figure 9.8 – Essais indiciels d'un système d'une deuxième ordre en fonction de ξ .

4 Durée du régime transitoire

4.1 Définition du temps de réponse T_R

On définit dans la littérature le **temps de réponse à 5%**. Le terme « littérature » renvoie ici à l'ensemble des publications scientifiques et techniques dans le domaine étudié. Ce temps de réponse à 5%, noté T_R , est la durée au bout de laquelle le système atteint sa valeur finale à moins de 5%, lors d'un essai indiciel (réponse à un échelon de hauteur E_0). Le signal reste alors compris entre 0,95 et 1,05 fois la valeur finale. Cette durée correspond à la durée du régime transitoire. Cette définition est illustrée par deux exemples ci-dessous.

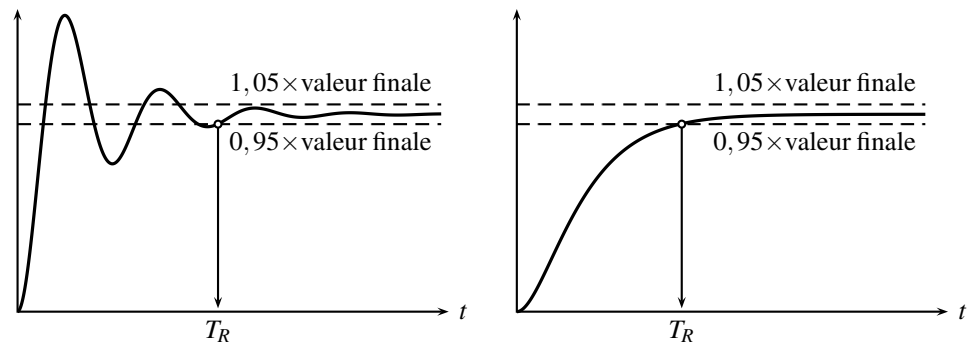
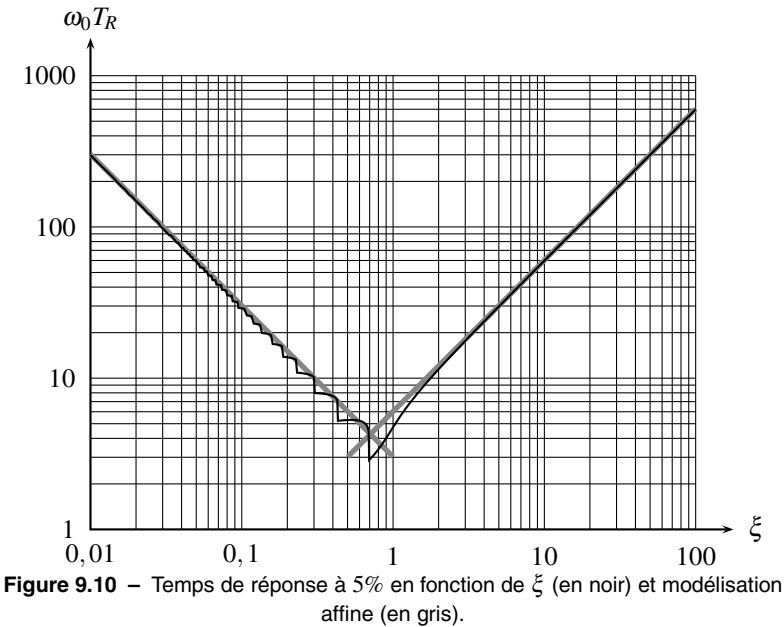


Figure 9.9 – Mise en évidence du temps de réponse à 5%.

Ce temps de réponse est calculable pour chaque valeur du facteur d'amortissement ξ . On trace ensuite le paramètre sans dimension $\omega_0 T_R$ en fonction du facteur d'amortissement ξ , en échelle logarithmique.

Qu'est-ce qu'un diagramme avec des échelles logarithmique ? On trace $\log(\omega_0 T_R)$ en fonction de $\log(\xi)$. Cependant, on continue à graduer les axes par les valeurs de $\omega_0 T_R$ et ξ pour plus de clarté. Les graduations ne sont alors plus équidistantes, mais se resserrent vers le haut et la droite.

Le graphe est tracé sur la figure 9.10. On observe que le minimum a pour coordonnées $(\xi = 0,69; \omega_0 T_R \approx 3)$.



Le temps de réponse minimum $T_{R,\min}$ est obtenu pour $\xi = 0,69$ et vaut $T_{R,\min} \approx \frac{3}{\omega_0}$.

Remarque

Cette valeur de $\xi = 0,69$ correspond à un régime pseudopériodique.

4.2 Complément : modélisation

La courbe en V du temps de réponse est approximée par deux droites de pentes opposées, qui se croisent en $\xi = 0,69$ (temps de réponse minimum). Ces droites sont des asymptotes. On observe que cette modélisation est d'autant plus vérifiée qu'on s'éloigne de $\xi = 0,69$.

Quelle est la relation entre $\omega_0 T_R$ et ξ pour chacune des deux droites ?

- Pour $\xi \in [10^{-2}, 1]$, on a tracé $\log(\omega_0 T_R) = 0,48 - \log(\xi)$:

$$\begin{aligned} \log(\omega_0 T_R) &= 0,48 - \log(\xi) \simeq \log(3,0) - \log(\xi) \\ \log(\omega_0 T_R) &\simeq \log\left(\frac{3,0}{\xi}\right) \quad \text{ou} \quad T_R \simeq \frac{3}{\xi \omega_0}. \end{aligned}$$

- Pour $\xi \in [5 \cdot 10^{-1}, 10^2]$, on a tracé $\log(\omega_0 T_R) = 0,78 + \log(\xi)$:

$$\begin{aligned} \log(\omega_0 T_R) &= 0,78 + \log(\xi) \simeq \log(6,0) + \log(\xi) \\ \log(\omega_0 T_R) &\simeq \log(6,0 \xi) \quad \text{ou} \quad T_R \simeq \frac{6 \xi}{\omega_0}. \end{aligned}$$

CHAPITRE 9 – CIRCUIT LINÉAIRE DU SECOND ORDRE

Ces deux expressions sont des approximations du temps de réponse T_R en fonction du coefficient d’amortissement ξ , valables approximativement pour $\xi < 0,2$ et $\xi > 2$, quand les droites grisées modélisent bien la courbe noire $\omega_0 T_R(\xi)$.

5 Réponse à un signal créniaux

5.1 Observations expérimentales

L’étude expérimentale de la réponse indicielle (réponse à un échelon de hauteur E_0) et du régime libre (entrée nulle mais conditions initiales différentes de zéro), s’effectue en alimentant le circuit par un signal créniaux de valeur moyenne différente de zéro et de période T . Le régime permanent doit avoir le temps de s’établir pour chaque demi-période, aussi doit-on imposer $T \gg T_R$, où T_R est le temps de réponse du système.

Pour le circuit série constitué de $C = 33 \text{ nF}$, $L = 560 \text{ mH}$ et $R = 3 \text{ k}\Omega$, déjà étudié en début de chapitre, on calcule :

$$\omega_0 = \frac{1}{\sqrt{LC}} \simeq 7,4.10^3 \text{ rad.s}^{-1} \quad \text{et} \quad \xi = \frac{R}{2} \sqrt{\frac{C}{L}} \simeq 0,37.$$

Avec le graphe de $\omega_0 T_R$ en fonction de ξ , on trouve $\omega_0 T_R \approx 8$, c’est à dire $T_R \approx 1,1 \text{ ms}$. Le résultat est le même avec la formule simplifiée $T_R \approx \frac{3}{\xi \omega_0}$. Le choix de $T = 8 \text{ ms}$, observé sur l’oscilloscope de la page suivante, est donc valable.

On observe la charge du condensateur lors de l’essai indiciel, puis sa décharge lors du régime libre.

Le portrait de phase est explicite. Le système part du point A de coordonnées $(\dot{u}_C = 0, u_C = 0)$ pour aboutir au point B de coordonnées $(\dot{u}_C = 0, u_C = E_0)$ lors de l’essai indiciel, en oscillant, car $\dot{u}_C$ est alternativement positif et négatif. Puis il revient vers A lors du régime libre, en oscillant de même.

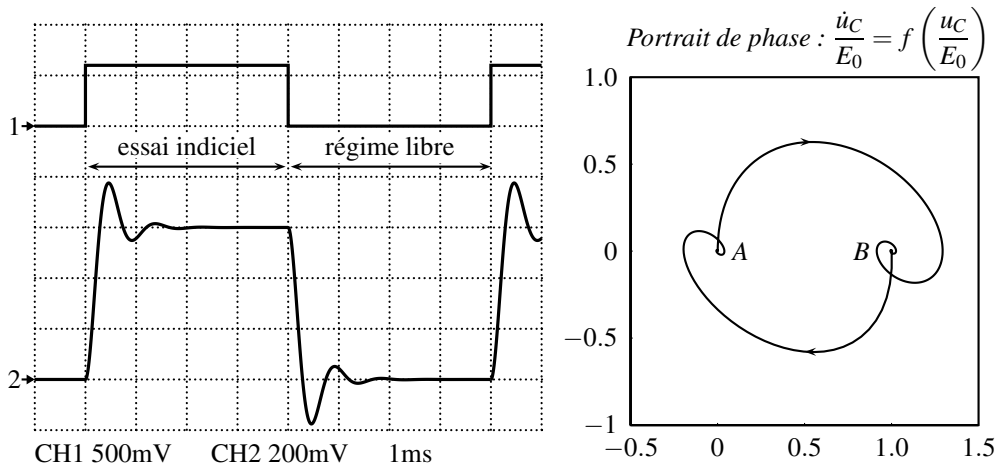


Figure 9.11 – Essai indiciel, régime libre et portrait de phase d’un système pseudopériodique.

5.2 Modélisation du régime libre

La mise en équation complète de la phase de régime libre requiert :

- l'équation différentielle qui modélise l'évolution du système,
- les conditions initiales qui précisent de quels points le système part.

De plus, on choisit de prendre pour nouvel instant initial $t = 0$, la date où ce régime commence. Ce choix de la date $t = 0$ est toujours arbitraire.

Dans un régime libre, le système n'est plus alimenté, $e(t) = 0$. L'équation différentielle sur $u_C(t)$ devient homogène (elle a un second membre nul) :

$$\frac{1}{\omega_0^2} \frac{d^2 u_C}{dt^2} + \frac{2\xi}{\omega_0} \frac{du_C}{dt} + u_C = 0 \quad \text{avec} \quad \begin{cases} \omega_0 \simeq 7,4.10^3 \text{ rad.s}^{-1} \\ \xi \simeq 0,37 < 1. \end{cases}$$

Le système est faiblement amorti, $\xi < 1$, il oscille donc.

Quelle est la condition initiale ? L'observation des formes d'onde montre qu'en fin d'essai indiciel, $u_C = E_0$ et $\frac{du_C}{dt} = 0$ pour $t < 0$. En effet, la tension aux bornes du condensateur reste constante, donc le courant qui le traverse est nul $i = C \frac{du_C}{dt} = 0$. À $t = 0^+$, au tout début du régime libre, $u_C(0^+) = E_0$, car la tension aux bornes du condensateur est continue ; de même $i(0^+) = C \left(\frac{du_C}{dt} \right)_{t=0^+} = 0$ car l'intensité du courant dans la bobine est continue. Ainsi, le problème devient :

$$\begin{cases} \frac{1}{\omega_0^2} \frac{d^2 u_C}{dt^2} + \frac{2\xi}{\omega_0} \frac{du_C}{dt} + u_C = 0 & (\omega_0 \simeq 7,4.10^3 \text{ rad.s}^{-1}) \\ u_C(0) = E_0 & \\ \left(\frac{du_C}{dt} \right)_{t=0} = 0 & (\xi \simeq 0,37 < 1) \end{cases}$$

La solution de l'équation différentielle, qui est homogène, est identique à celle calculée pour l'essai indiciel au paragraphe 3.3. Pour $\xi < 1$:

$$u_C(t) = \exp(-\omega_0 \xi t) \left(u_1 \cos(\omega_0 \sqrt{1 - \xi^2} t) + u_2 \sin(\omega_0 \sqrt{1 - \xi^2} t) \right).$$

La condition initiale sur $u_C(0)$ impose :

$$u_C(0) = E_0 = u_1.$$

La dérivée de $u_C(t)$ par rapport à t est :

$$\begin{aligned} \frac{du_C}{dt}(t) &= -\omega_0 \xi \exp(-\omega_0 \xi t) \left(u_1 \cos(\omega_0 \sqrt{1 - \xi^2} t) + u_2 \sin(\omega_0 \sqrt{1 - \xi^2} t) \right) \\ &\quad + \exp(-\omega_0 \xi t) \omega_0 \sqrt{1 - \xi^2} \left(-u_1 \sin(\omega_0 \sqrt{1 - \xi^2} t) + u_2 \cos(\omega_0 \sqrt{1 - \xi^2} t) \right). \end{aligned}$$

CHAPITRE 9 – CIRCUIT LINÉAIRE DU SECOND ORDRE

La condition initiale sur la dérivée de u_C impose :

$$\left(\frac{du_C}{dt}\right)_{t=0} = 0 = -\omega_0 \xi u_1 + \omega_0 \sqrt{1-\xi^2} u_2.$$

Ainsi :

$$u_2 = \frac{\xi}{\sqrt{1-\xi^2}} E_0 \simeq 0,40 \times E_0.$$

Finalement :

$$u_C(t) = E_0 \exp(-\omega_0 \xi t) \left(\cos(\omega_0 \sqrt{1-\xi^2} t) + \frac{\xi}{\sqrt{1-\xi^2}} \sin(\omega_0 \sqrt{1-\xi^2} t) \right).$$

Ce résultat n'est évidemment pas à connaître par cœur. Seule la méthode utilisée est à maîtriser.

La solution obtenue part bien de E_0 pour tendre vers 0, en oscillant avec une pseudopulsation $\omega_0 \sqrt{1-\xi^2} \simeq 6,9.10^3 \text{ rad.s}^{-1}$, ou une **pseudopériode** :

$$T = \frac{2\pi}{\omega_0 \sqrt{1-\xi^2}} \simeq 9,4.10^{-1} \text{ ms}.$$

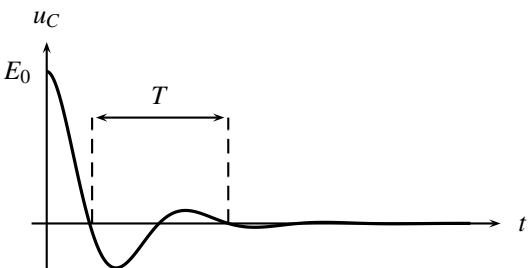


Figure 9.12 – Pseudopériode d'un système du deuxième ordre.

6 Bilan énergétique

Multiplions par $i(t)$ la loi des mailles qui permet d'établir l'équation différentielle modélisant le comportement du circuit au paragraphe 2.1 pour obtenir :

$$e \times i = u_C \times i + u_L \times i + u_R \times i,$$

soit :

$$ei = u_C \times C \frac{du_C}{dt} + L \frac{di}{dt} \times i + Ri \times i = \frac{d}{dt} \left(\frac{1}{2} C u_C^2 + \frac{1}{2} L i^2 \right) + R i^2.$$

Apparaissent dans cette équation l'énergie du condensateur, stockée sous forme électrique, $\mathcal{E}_{\text{élec}} = \frac{1}{2} C u_C^2$, ainsi que l'énergie de la bobine, stockée sous forme magnétique, $\mathcal{E}_{\text{magné}} = \frac{1}{2} L i^2$.

Dans le cas de l'**essai indiciel**, où le système est alimenté par une tension $e(t) = E_0$, le bilan de puissance s'écrit :

$$E_0 i = \frac{d}{dt} (\mathcal{E}_{\text{élec}} + \mathcal{E}_{\text{magné}}) + Ri^2.$$

La puissance délivrée par le générateur ($E_0 i$) alimente les pertes par effet Joule (Ri^2) et délivre de la puissance au condensateur et à la bobine. Le condensateur se charge, son énergie augmente au cours du temps.

Quant au **régime libre**, pour lequel le système évolue librement, lorsque le GBF ne délivre aucune tension ($e = 0$) :

$$\frac{d}{dt} (\mathcal{E}_{\text{élec}} + \mathcal{E}_{\text{magné}}) = -Ri^2.$$

Dans ce cas, toute l'énergie contenue dans le condensateur et la bobine est dissipée par effet Joule. L'état final est décrit par une énergie totale nulle, c'est à dire par $u_C(t) = 0$ et $i(t) = 0$, conformément aux observations expérimentales.

7 Analogies

Il existe de nombreux systèmes, tant en sciences physiques qu'en SII, décrits par des équations différentielles du deuxième ordre. Ces systèmes sont nommés **systèmes du deuxième ordre**.

Un exemple mécanique est celui d'une masse attachée à un ressort, qui oscille verticalement, freinée par des frottements fluides avec l'air. La position $x(t)$ de la masse est repérée par rapport à sa position d'équilibre. Ce système est étudié dans la partie mécanique. Son mouvement obéit à l'équation différentielle :

$$m \frac{d^2 x}{dt^2} = -kx - f \frac{dx}{dt}.$$

Cette équation différentielle peut être écrite sous forme canonique :

$$\frac{m}{k} \frac{d^2 x}{dt^2} + \frac{f}{k} \frac{dx}{dt} + x(t) = 0 \quad \text{soit} \quad \frac{1}{\omega_0^2} \frac{d^2 x}{dt^2} + \frac{2\xi}{\omega_0} \frac{dx}{dt} + x(t) = 0$$

sur laquelle on détermine la pulsation caractéristique ω_0 et le facteur d'amortissement ξ :

$$\left\{ \begin{array}{l} \frac{1}{\omega_0^2} = \frac{m}{k} \\ \frac{2\xi}{\omega_0} = \frac{f}{k} \end{array} \right. \quad \text{donc} \quad \left\{ \begin{array}{l} \omega_0 = \sqrt{\frac{k}{m}} \\ \xi = \frac{f}{2\sqrt{km}} \end{array} \right.$$

Le formalisme est alors totalement analogue entre électricité et mécanique.

CHAPITRE 9 – CIRCUIT LINÉAIRE DU SECOND ORDRE

SYNTHÈSE

SAVOIRS

- distinguer le régime transitoire du régime permanent ou établi
- équation différentielle canonique du deuxième ordre
- pulsation caractéristique d'un système du deuxième ordre
- coefficient d'amortissement d'un système du deuxième ordre
- facteur de qualité d'un système du deuxième ordre
- temps de réponse d'un système du deuxième ordre

SAVOIR-FAIRE

- tracer l'allure de la réponse indicielle ou du régime libre en fonction des valeurs de ξ et de Q
- établir l'équation différentielle du circuit
- écrire l'équation différentielle sous forme canonique
- résoudre une équation différentielle du deuxième ordre
- déterminer un ordre de grandeur de la durée du régime transitoire
- prévoir l'évolution d'un système avec son portrait de phase
- prévoir l'évolution d'un système à partir de considérations énergétiques
- réaliser un bilan de puissance

MOTS-CLÉS

- | | | |
|----------------------------|-----------------------------|----------------------|
| • régime transitoire | • réponse indicielle | ment |
| • régime permanent ou éta- | • réponse libre | • facteur de qualité |
| bli | • pulsation caractéristique | • portrait de phase |
| • échelon | • coefficient d'amortisse- | |

S'ENTRAÎNER

9.1 Paramètres caractéristiques d'un système (★)

Un système linéaire est décrit par son équation différentielle : $2\frac{d^2s}{dt^2} + \frac{ds}{dt} + 5s(t) = e(t)$.

- 1. Calculer ses paramètres caractéristiques (pulsation caractéristique, coefficient d'amortissement). En déduire son facteur de qualité.
- 2. Ce système est alimenté par un échelon. Quel est son temps de réaction ?

9.2 Influence d'un condensateur sur un circuit RL (★)

Soit un générateur de tension de force électromotrice E et de résistance interne r . Il est placé aux bornes d'une bobine d'inductance L et de résistance R . On suppose qu'à l'instant initial, l'ensemble fonctionne en régime permanent. On branche alors en parallèle avec la bobine un condensateur de capacité C . L'objet de ce problème est l'étude de l'intensité i traversant la bobine.

- 1. Déterminer la valeur de i à $t = 0^+$ juste après avoir branché le condensateur.
- 2. Même question pour $\frac{di}{dt}$.
- 3. Établir l'équation différentielle vérifiée par i .
- 4. On donne $L = 43 \text{ mH}$, $R = 9,1 \text{ }\Omega$, $r = 50 \text{ }\Omega$ et $E = 5,0 \text{ V}$. Quelle valeur doit-on prendre pour C pour observer un régime quasipériodique ?
- 5. Déterminer l'expression littérale de i si $C = 10 \text{ }\mu\text{F}$.

APPROFONDIR

9.3 Régime transitoire d'un circuit R, L, C parallèle (★)

Soit un circuit constitué de l'association en série d'une source idéale de tension E , d'une résistance R , d'un interrupteur K et de l'association en parallèle d'une résistance r , d'une inductance L et d'une capacité C . Initialement l'interrupteur est ouvert, la capacité est déchargée et tous les courants sont nuls. On ferme l'interrupteur K à l'instant $t = 0$. On notera i l'intensité du courant dans R , i_1 dans L , i_2 dans C et i_3 dans r ainsi que u la tension aux bornes de r (ou de C ou de L).

- 1. Déterminer, en les justifiant, les valeurs de u , i , i_1 , i_2 et i_3 juste après la fermeture de l'interrupteur K .
- 2. Même question lorsque le régime permanent est complètement établi.
- 3. Établir l'équation différentielle vérifiée par i_3 pour $t > 0$.
- 4. L'écrire sous la forme :

$$\frac{d^2i_3}{dt^2} + 2\lambda\omega_0\frac{di_3}{dt} + \omega_0^2i_3 = 0.$$

On donnera les expressions de ω_0 et λ en fonction de r , R , L et C .

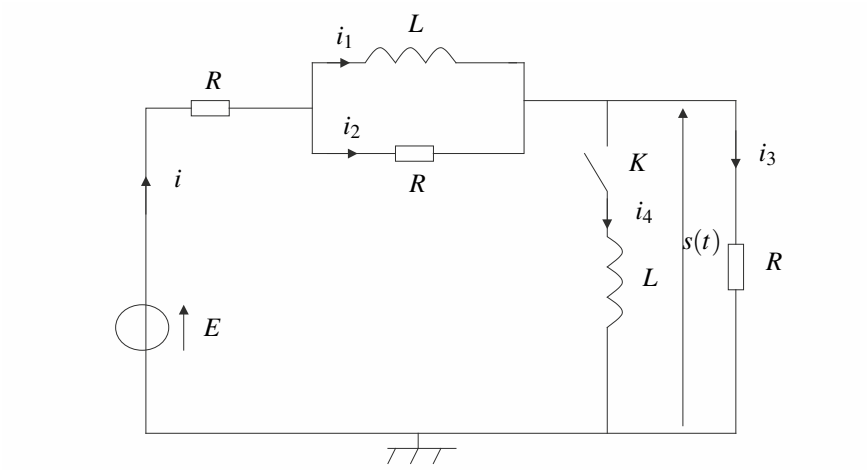
CHAPITRE 9 – CIRCUIT LINÉAIRE DU SECOND ORDRE

5. Établir la condition que doivent vérifier r , R , L et C pour observer un régime pseudo-périodique.
6. Que caractérise λ ?
7. Définir la pseudo-pulsation ω et la pseudo-période T et en donner les expressions en fonction de r , R , L et C .
8. Déterminer l'expression de i_3 en fonction du temps. On justifiera la détermination des constantes.
9. Déterminer le temps nécessaire pour que le régime permanent soit établi dans le circuit avec une précision d'un millième.
10. Déterminer, en les justifiant, les valeurs de u , i , i_1 , i_2 et i_3 juste après l'ouverture de l'interrupteur K .
11. Établir l'équation différentielle vérifiée par i_3 lorsqu'on ouvre l'interrupteur. On ne se préoccupera pas de ce qui se passe juste à l'instant de l'ouverture.
12. L'écrire sous une forme analogue à celle de la question 4.
13. Quelle est la nature du régime ? On suppose que la condition de la question 5. est vérifiée.
14. À quelle condition aura-t-on un signal non nul lors de l'ouverture de l'interrupteur ?

9.4 Étude d'un circuit RL (★)

On considère le circuit ci-après. L'interrupteur K est ouvert depuis très longtemps et on le ferme à l'instant $t = 0$.

1. Déterminer les valeurs des intensités i_4 , i_1 , i , i_2 et i_3 à l'instant $t = 0^+$.
2. Déterminer les valeurs de la tension s et des intensités i_2 , i , i_1 , i_3 et i_4 lorsque t tend vers l'infini.
3. Déterminer les relations entre s et i_3 puis entre s et i_4 .
4. Établir la relation entre R , L , s , $\frac{ds}{dt}$ et $\frac{di}{dt}$.
5. Même question pour R , L , i_2 , $\frac{di_2}{dt}$ et $\frac{di}{dt}$.
6. Déterminer la relation entre $\frac{dE}{dt}$, R , L , s , i_2 et $\frac{ds}{dt}$.
7. Établir la relation entre R , L , s , $\frac{ds}{dt}$, $\frac{d^2s}{dt^2}$, $\frac{dE}{dt}$ et $\frac{d^2E}{dt^2}$.
8. En déduire l'équation différentielle vérifiée par s sachant que le générateur de tension est idéal de tension à vide E .
9. En déduire l'expression de $s(t)$ en ne cherchant pas à déterminer les constantes d'intégration.
10. Donner une relation entre les deux constantes d'intégration.



CHAPITRE 9 – CIRCUIT LINÉAIRE DU SECOND ORDRE

CORRIGÉS

9.1 Paramètres caractéristiques d'un système

1. Il convient d'écrire l'équation différentielle sous la forme canonique (le terme facteur de $s(t)$ doit être 1) : $\frac{1}{\omega_0^2} \frac{d^2 s}{dt^2} + \frac{2\xi}{\omega_0} \frac{ds}{dt} + s(t) = \dots$. Dans notre cas : $\frac{2}{5} \frac{d^2 s}{dt^2} + \frac{1}{5} \frac{ds}{dt} + s(t) = \frac{e(t)}{5}$.

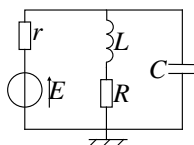
On identifie alors $\omega_0 : \frac{1}{\omega_0^2} = \frac{2}{5}$ donc $\omega_0 = 1,6 \text{ rad.s}^{-1}$. Puis $\xi : \frac{2\xi}{\omega_0} = \frac{1}{5}$ donc $\xi = 1,6 \cdot 10^{-1}$.

On en déduit : $Q = \frac{1}{2\xi} = 3,1$.

2. Connaissant ξ , on consulte la courbe en V du temps de réponse d'un système du deuxième ordre soumis à un échelon. Pour $\xi = 1,6 \cdot 10^{-1}$, on a $\omega_0 T_R \approx 20$, soit $T_R \approx 12,5 \text{ s}$.

9.2 Influence d'un condensateur sur un circuit RL

Le circuit étudié est le suivant :



1. L'intensité traversant l'inductance L pour $t < 0$ s'obtient en écrivant la loi des mailles sur la maille comprenant le générateur et la bobine soit $E - r i_r - R i + L \frac{di}{dt} = 0$. Or pour $t < 0$, le circuit est en régime permanent donc $\frac{di}{dt} = 0$ et il n'y a pas de courant dans la capacité donc $i_r = i$. On en déduit $i(t < 0) = \frac{E}{r+R}$. On utilise alors la continuité de l'intensité du courant traversant une inductance L et on obtient $i(0^+) = \frac{E}{r+R}$.

2. On utilise la continuité de la tension u aux bornes de la capacité C qui est aussi la tension aux bornes de la bobine ou de l'ensemble L et R . On écrit la loi des mailles à $t = 0^+$ soit $R i(0^+) + L \frac{di}{dt}(0^+) = 0$. On en déduit $\frac{di}{dt}(0^+) = -\frac{RE}{L(r+R)}$.

3. On écrit la loi des mailles sur la maille comprenant le générateur et la bobine soit $E - r i_r - R i + L \frac{di}{dt} = 0$ ainsi que la loi des nœuds, ce qui donne $i_r = i + i_C$. Par ailleurs, la relation entre intensité et tension pour une capacité fournit $i_C = C \frac{du}{dt}$ avec $u = R i + L \frac{di}{dt}$. On en déduit $\frac{d^2 i}{dt^2} + \left(\frac{1}{rC} + \frac{R}{L} \right) \frac{di}{dt} + \frac{R+r}{rLC} i = \frac{E}{rLC}$ en éliminant i_C et u .

4. Pour obtenir un régime quasipériodique ou pseudo-périodique, il faut que le déterminant de l'équation caractéristique associé à l'équation différentielle du second ordre à coefficients constants soit négatif. Ici on a $\Delta = \left(\frac{1}{rC} + \frac{R}{L} \right)^2 - \frac{4(r+R)}{rLC}$ qui doit être négatif. On peut

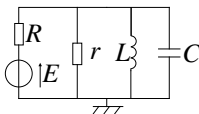
réécrire cette condition sous la forme $r^2R^2C^2 - 2rLC(2r + 3R) + L^2 < 0$. Il s'agit d'un trinôme du second degré en C qui doit être négatif, il faut donc que la valeur de C soit entre les racines de ce polynôme en C pour que l'inégalité soit vérifiée. Cette condition donne $\frac{rL(2r + 3R) \pm \sqrt{r^2L^2(2r + 3R)^2 - r^2R^2L^2}}{R^2r^2}$ soit numériquement $3,36 \mu\text{F} < C < 2,64 \text{ mF}$.

5. D'après la réponse à la question précédente, le régime est pseudopériodique et la résolution en tenant compte des conditions initiales établies au début du problème, on obtient

$$i = -\frac{2rRC \exp\left(-\left(\frac{1}{2rC} + \frac{R}{2L}\right)t\right)}{(r + R)\sqrt{4(r + R)rLC - (L + rRC)^2}} \sin \frac{\sqrt{4(r + R)rLC - (L + rRC)^2}}{2rLC} t.$$

9.3 Régime transitoire d'un circuit R, L, C parallèle

Le circuit étudié est le suivant :



1. On utilise la continuité de l'intensité du courant traversant une inductance et de la tension aux bornes d'une capacité pour déterminer les conditions initiales juste après la fermeture de l'interrupteur K . On obtient immédiatement $i_1(0^+) = u(0^+) = 0$. Par la loi d'Ohm, on a $i_3(0^+) = \frac{u(0^+)}{r} = 0$. En écrivant la loi des nœuds avec les résultats précédents, on a $i(0^+) =$

$i_2(0^+)$ et il suffit d'écrire une loi des mailles pour obtenir $i(0^+) = i_2(0^+) = \frac{E}{R}$.

2. En régime permanent, les grandeurs sont constantes et les dérivées temporelles nulles. Les relations intensité - tension pour les dipôles donnent $i_2(\infty) = u(\infty) = 0$. Alors en refaisant le même raisonnement qu'à la question précédente, on obtient $i_3(\infty) = 0$ et $i(\infty) = i_1(\infty) = \frac{E}{R}$.

3. On a $u = E - Ri = ri_3 = L \frac{di_1}{dt}$ et $u = \frac{du}{dt}$ soit en tenant compte du fait que E est constante, on obtient par dérivation $\frac{di}{dt} = -\frac{r}{R} \frac{di_3}{dt}$, $\frac{di_1}{dt} = \frac{r}{L} i_3$, $\frac{di_2}{dt} = rC \frac{d^2 i_3}{dt^2}$. On dérive la loi des nœuds $i = i_1 + i_2 + i_3$ et en utilisant les relations précédentes, on obtient $\frac{d^2 i_3}{dt^2} + \frac{r + R}{RrC} \frac{di_3}{dt} + \frac{1}{LC} i_3 = 0$.

4. Par identification, on en déduit $\omega_0 = \frac{1}{\sqrt{LC}}$ et $\lambda = \frac{R + r}{2Rr} \sqrt{\frac{L}{C}}$.

5. Le régime sera pseudo-périodique si le discriminant de l'équation caractéristique $r^2 + 2\lambda\omega_0 r + \omega_0^2 = 0$ est négatif soit $4\lambda^2\omega_0^2 - 4\omega_0^2 < 0$ qu'on peut écrire $\lambda < 1$ ou en explicitant la valeur de λ $\frac{R + r}{2rR} < \sqrt{\frac{C}{L}}$.

6. La grandeur λ caractérise l'amortissement des oscillations.

CHAPITRE 9 – CIRCUIT LINÉAIRE DU SECOND ORDRE

7. Dans le cas du régime pseudo-périodique considéré ici, la solution de l'équation caractéristique s'écrit $r_{\pm} = -\lambda \omega_0 \pm j\omega_0 \sqrt{1 - \lambda^2}$. On en déduit l'expression de la pseudo-pulsation

$$\omega = \omega_0 \sqrt{1 - \lambda^2} = \frac{\sqrt{4r^2 R^2 C - (R + r)^2 L}}{2rRC\sqrt{L}} \text{ et de la pseudo-période } T = \frac{2\pi}{\omega} \text{ soit}$$

$$T = 4\pi RrC\sqrt{L} / \sqrt{4R^2 r^2 C - (R + r)^2 L}.$$

8. Dans ces conditions, l'intensité i_3 s'écrit $i_3 = (A \cos \omega t + B \sin \omega t) \exp(-\lambda \omega_0 t)$. La détermination des constantes A et B s'obtient à partir des conditions initiales de la question 1. soit

$$i_3(0) = 0 \text{ et } \frac{di_3}{dt}(0) = \frac{E}{RrC}.$$

$$\text{Or } i_3(0) = A \text{ et } \frac{di_3}{dt} = ((-\omega A - \lambda \omega B) \sin \omega t + (\omega B - \lambda \omega A) \cos \omega t) \exp(-\lambda \omega t) \text{ donc } \frac{di_3}{dt}(0) = \omega B - \lambda \omega A. \text{ D'où } A = 0, B = \frac{E}{RrC\omega} \text{ et } i_3 = \frac{E}{RrC\omega} \sin \omega t \exp(-\lambda \omega_0 t).$$

9. Le régime est établi avec une précision de un millièmme si l'amplitude de i_3 est inférieure au millièmme de sa valeur maximale soit $\frac{E}{1000RrC\omega}$. On résout donc $i_3 = \frac{E}{RrC\omega} \exp(-\lambda \omega_0 t) < \frac{E}{1000RrC\omega}$. On obtient $t \geq \frac{\ln 1000}{\lambda \omega_0}$.

10. Pour déterminer les nouvelles conditions initiales juste après l'ouverture de l'interrupteur, on raisonne comme à la question 1. On obtient $i = u = i_3 = 0$ et $i_1 = \frac{E}{R} = -i_2$.

11. De même, avec un raisonnement analogue à celui de la question 3., on a $\frac{d^2 i_3}{dt^2} + \frac{1}{rC} \frac{di_3}{dt} + \frac{1}{LC} i_3 = 0$.

12. Par identification, on obtient $\omega_0 = \frac{1}{\sqrt{LC}}$ et $\lambda = \frac{1}{2r} \sqrt{\frac{L}{C}}$.

13. Comme $\frac{1}{2r} < \frac{r+R}{RrC} < \sqrt{\frac{C}{L}}$, on a un régime pseudo-périodique.

14. On aura un signal non nul uniquement si le courant dans la bobine est non nul.

9.4 Étude d'un circuit RL

1. On applique la continuité de l'intensité du courant traversant les inductances. Cela impose de façon immédiate $i_4(0^+) = 0$. Par ailleurs, à $t = 0^-$, le circuit est en régime permanent donc les dérivées temporelles sont nulles, ce qui donne pour la tension aux bornes de l'inductance parcourue par le courant i_1 $u(0^-) = L \frac{di_1}{dt}(0^-) = 0$. On en déduit par la loi d'Ohm que

$$i_2(0^-) = \frac{u(0^-)}{R} = 0 \text{ puis que } i(0^-) = i_1(0^-) = i_1(0^+) \text{ par continuité de l'intensité du courant traversant l'inductance. L'écriture de la loi des mailles donne } E = Ri(0^-) + Ri_3(0^-) = 2Ri(0^-) \text{ en utilisant la loi des nœuds } i(0^-) = i_3(0^-) + i_4(0^-) = i_3(0^-) \text{ soit } i(0^-) = \frac{E}{2R} = i_1(0^-). \text{ La continuité de l'intensité dans l'inductance permet alors d'écrire } i_1(0^+) = \frac{E}{2R}. \text{ La}$$

$$\text{loi des nœuds } i(0^+) = i_1(0^+) + i_2(0^+) = \frac{E}{2R} + i_2(0^+) = i_3(0^+) + i_4(0^+) = i_3(0^+) \text{ donne}$$

en reportant dans la loi des mailles $E = Ri(0^+) + Ri_2(0^+) + Ri_3(0^+)$ $i(0^+) = \frac{E}{2R}$ soit avec relations précédentes $i_2(0^+) = 0$ et $i_3(0^+) = \frac{E}{2R}$.

La modélisation du circuit par sa fonction de transfert, alliée à l'utilisation du théorème de la valeur initiale, vu en cours de SII, permettra dans la suite de répondre à ce genre de questions beaucoup plus aisément.

2. Pour un temps infini, le régime est permanent donc les dérivées temporelles sont nulles. On en déduit $s(\infty) = 0$ et $u(\infty) = Ri_2(\infty) = 0$. En reportant dans la loi des mailles, on en déduit $E = Ri(\infty) + Ri_2(\infty) + Ri_3(\infty)$ soit $i(\infty) = \frac{E}{R}$. Alors avec la loi des nœuds $i(\infty) = i_1(\infty) + i_2(\infty)$, on a par la loi des mailles $i_1(\infty) = \frac{E}{R}$ et $i_3(\infty) = \frac{s(\infty)}{R} = 0$. Enfin la loi des nœuds donne $i(\infty) = i_3(\infty) + i_4(\infty)$ et $i_4(\infty) = \frac{E}{R}$.

3. Par la loi d'Ohm, on a $s = Ri_3$ et la relation entre l'intensité et la tension pour une inductance donne $s = L \frac{di_4}{dt}$.

4. La loi des nœuds s'écrit $i = i_3 + i_4$ soit en dérivant et en utilisant les relations de la question précédente, on en déduit : $\frac{di}{dt} = \frac{1}{R} \frac{ds}{dt} + \frac{s}{L}$.

5. L'égalité des tensions pour deux dipôles en parallèle s'écrit $Ri_2 = L \frac{di_1}{dt}$ avec $i = i_1 + i_2$ par la loi des nœuds. On a donc : $\frac{di}{dt} = \frac{R}{L} i_2 + \frac{di_2}{dt}$.

6. Par une loi des mailles, on a $E = Ri + Ri_2 + Ri_3$, ce qui permet en dérivant et en utilisant les relations précédentes en dérivant et en utilisant les relations établies auparavant d'obtenir :

$$\frac{dE}{dt} = 2 \frac{ds}{dt} + \frac{R}{L} s + R \frac{di_2}{dt} = 3 \frac{ds}{dt} + 2 \frac{R}{L} s - \frac{R^2}{L} i_2.$$

7. En dérivant à nouveau cette relation et en utilisant encore les relations établies aux questions précédentes, on en déduit $\frac{d^2E}{dt^2} + \frac{R}{L} \frac{dE}{dt} = 3 \frac{d^2s}{dt^2} + 4 \frac{R}{L} \frac{ds}{dt} + \left(\frac{R}{L}\right)^2 s$.

8. Ici la tension E est constante donc sa dérivée est nulle et on a l'équation différentielle :

$$3 \frac{d^2s}{dt^2} + 4 \frac{R}{L} \frac{ds}{dt} + \left(\frac{R}{L}\right)^2 s = 0.$$

9. L'équation caractéristique de cette équation différentielle du second ordre est $3r^2 + 4 \frac{R}{L} r + \left(\frac{R}{L}\right)^2 = 0$. Son discriminant s'écrit $\Delta = \left(2 \frac{R}{L}\right)^2 > 0$. La solution est donc de la forme :

$$s(t) = A \exp\left(-\frac{R}{L}t\right) + B \exp\left(-\frac{R}{3L}t\right).$$

10. Par les conditions initiales, on a $s(0^+) = Ri_3(0^+) = \frac{E}{2}$. Or par la forme de la solution, on a $s(0^+) = A + B$. La relation demandée est donc $A + B = 0$.

CHAPITRE 9 – CIRCUIT LINÉAIRE DU SECOND ORDRE

Régime sinusoïdal

10

La réponse temporelle d'un système à une excitation constante a été étudiée précédemment. Toutefois, la majeure partie des signaux utiles ne sont pas constants mais variables dans le temps. Le but du présent chapitre est donc d'étudier comment un système réagit à un signal variable. Le choix d'une excitation sinusoïdale sera justifié dans le chapitre suivant.

1 Régimes transitoire et permanent sinusoïdal

On étudie un système du deuxième ordre, électrique ou mécanique, alimenté par une entrée sinusoïdale $e(t) = E_0 \cos(\omega t) = E_0 \cos(2\pi ft)$. On observe, avec un oscilloscope réglé en mode « monocoup », la sortie $s(t)$ sur la voie 2, après la fermeture du circuit, ainsi que l'entrée $e(t)$ sur la voie 1.

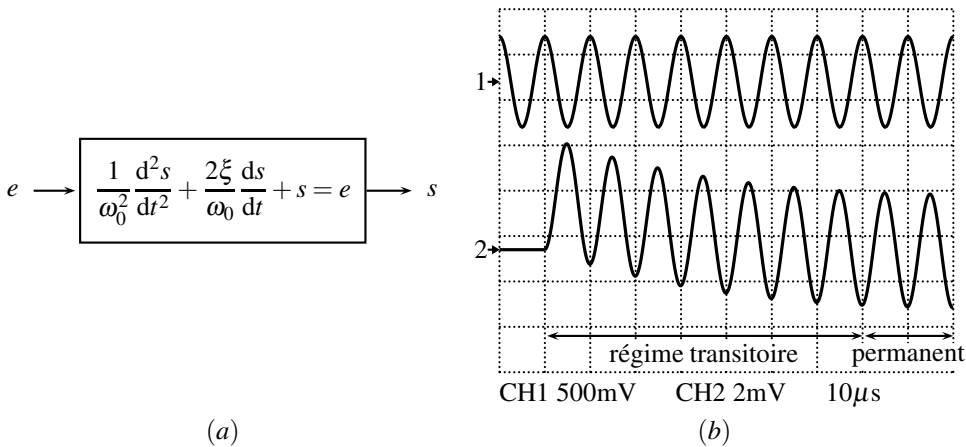


Figure 10.1 – Réponses d'un système du deuxième ordre : (a) Système étudié, (b) Système électrique $f_0 = 5,0 \text{ kHz}$, $\xi = 0,8$, $f = 100 \text{ kHz}$.

Les réponses sont très variées. On observe toutefois dans chaque cas deux régimes :

- le **régime établi**, ou **régime permanent**, pour lequel la sortie est de la même forme que

CHAPITRE 10 – RÉGIME SINUSOÏDAL

- l'entrée, ici sinusoïdale de valeur moyenne nulle. Ce régime est aussi nommé **régime sinusoïdal forcé**.
- Le **régime transitoire**, entre l'instant initial et le régime permanent. Ce régime a des formes extrêmement variées suivant les systèmes et les fréquences des signaux d'entrée.

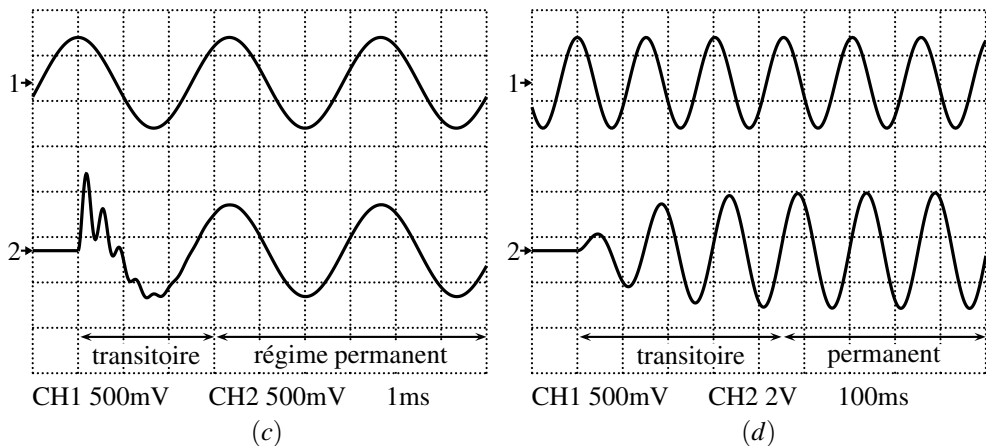


Figure 10.2 – Réponses d'un système du deuxième ordre : (c) Système électrique $f_0 = 2,7 \text{ kHz}$, $\xi = 0,1$, $f = 300 \text{ Hz}$, (d) Système mécanique $f_0 = 7 \text{ Hz}$, $\xi = 0,2$, $f = 6,6 \text{ Hz}$.

Le but de ce chapitre est de développer les outils pour analyser le régime permanent, ou sinusoïdal forcé, qui est le seul à être observé sur un oscilloscope en mode « normal » ou « auto ».

2 Régime permanent ou établi

2.1 Observation à l'oscilloscope

Lorsqu'on visualise un signal sinusoïdal sur un oscilloscope, l'origine des temps, c'est-à-dire la date $t = 0$, n'est pas précisée. On peut décaler cette origine et la placer arbitrairement où l'on veut. La **modélisation** du signal qui en résulte dépend alors de ce choix. Par exemple, sur l'oscillogramme de la figure 10.3, page suivante :

- avec $t = 0$ en 1, alors on modélise le signal $s(t)$ par : $s(t) = S_0 \sin(\omega t)$,
- avec $t = 0$ en 2, $s(t)$ devient : $s(t) = S_0 \cos(\omega t)$.

Les oscilloscopes numériques permettent tous de mesurer les caractéristiques des signaux : amplitude crête à crête (C–C), valeur moyenne, valeur efficace, période, fréquence, durée écoulée entre deux curseurs verticaux, tension entre deux curseurs horizontaux... Cette fonctionnalité est ici utilisée pour afficher directement la fréquence et l'amplitude crête à crête.

L'amplitude crête à crête est la différence entre les valeurs maximale et minimale du signal. Elle vaut deux fois l'amplitude S_0 .

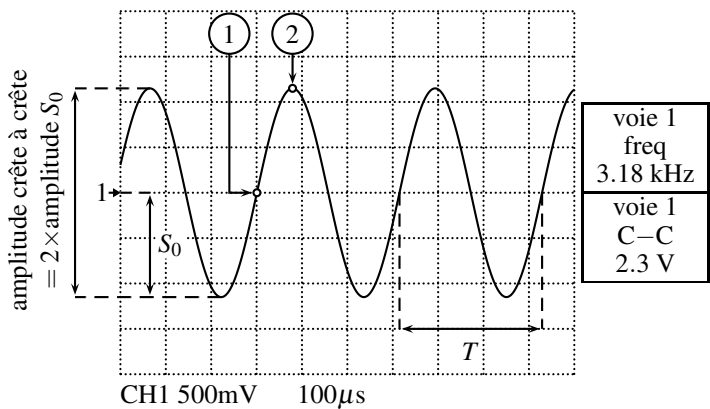


Figure 10.3 – Amplitude, période et modélisation d'un signal sinusoïdal.

2.2 Mesure d'un déphasage

On visualise sur un oscilloscope, pour le circuit RC suivant ($R = 1\text{ k}\Omega, C = 2,2\text{ }\mu\text{F}$), la tension $e(t)$, sur la voie 1 en gris, et la tension $u_C(t)$, sur la voie 2 en noir. Deux curseurs verticaux marquent les dates successives pour lesquelles les deux tensions passent par 0, avec une pente de même signe, ici positive :

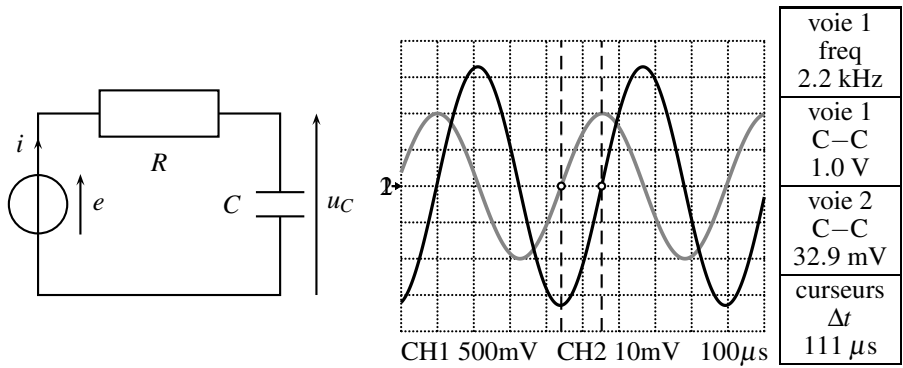


Figure 10.4 – Montage du premier ordre étudié et formes d'ondes des tensions. En gris, la voie 1, en noir, la voie 2.

On observe expérimentalement que :

- la fréquence de sortie est égale à celle d'entrée : c'est une caractéristique du régime permanent sinusoïdal d'un système linéaire,
- les maxima en sortie et en entrée ne se produisent pas pour les mêmes dates : les signaux sont **déphasés**,
- l'amplitude du signal de sortie est différente de celle de l'entrée,
- le déphasage et le rapport des amplitudes dépendent tous deux de la fréquence.

CHAPITRE 10 – RÉGIME SINUSOÏDAL

On observe sur l'oscilloscope que $e(t)$, en gris, passe par zéro avant $u_C(t)$, en noir. $e(t)$ est donc en avance sur $u_C(t)$, ou $u_C(t)$ est en retard sur $e(t)$ d'une durée Δt : $u_C(t)$ est **déphasé** par rapport à $e(t)$ d'un **déphasage** $\varphi < 0$.



Pour savoir quel signal est en avance sur quel autre, il faut regarder le passage des deux signaux par zéro, avec une pente de même signe.

Comment mesurer le déphasage φ entre les deux signaux ? Comme on l'a vu au chapitre 1, on le déduit du retard temporel Δt de $u_C(t)$ sur $e(t)$, mesuré par la distance entre les deux curseurs verticaux de l'oscilloscope, par la formule :

$$\varphi = -\omega \Delta t = -2\pi f \Delta t.$$

Exemple

Sur la figure 10.4, on mesure $\Delta t = 111 \mu s$ ce qui permet de calculer :

$$\varphi = -2\pi \times 2,2 \cdot 10^3 \times 111 \cdot 10^{-6} = -1,53 \text{ rad} = -88^\circ;$$

On modélisera donc les deux signaux par :

$$e(t) = E_0 \cos(\omega t) \quad \text{et} \quad u_C(t) = U_{C0} \cos(\omega t + \varphi).$$

3 Systèmes du premier ordre

Pour le circuit RC du paragraphe 2.2 précédent, comment prévoir *a priori* le déphasage entre les deux signaux et leur rapport d'amplitude en régime permanent sinusoïdal ? Il existe deux méthodes équivalentes dont la première est vectorielle, celle des **vecteurs de Fresnel**, introduite dans le chapitre 1, et la seconde algébrique.

3.1 Méthode des vecteurs de Fresnel (MPSI)

On a montré au chapitre 8 que, dans ce circuit, la tension $u_C(t)$ est liée à $e(t)$ par l'équation différentielle du premier ordre :

$$u_C(t) + RC \frac{du_C}{dt} = e(t).$$

Pour $e(t) = E_0 \cos(\omega t)$ on cherche ici une solution $u_C(t)$ sinusoïdale qui correspond au régime permanent ou établi.

Pour cela on associe à chaque terme de l'équation un vecteur de Fresnel. Le vecteur de Fresnel $\vec{E}$, associé à $e(t)$, de norme E_0 , est donc la somme du vecteur de Fresnel $\vec{U}_C$, associé à $u_C(t)$, de norme U_{C0} et de celui $RC \vec{DU}_C$, associé à $RC \frac{du_C}{dt}$, de norme $RC\omega U_{C0}$ et tourné de $\frac{\pi}{2}$:

$$\vec{U}_C + RC \times \vec{DU}_C = \vec{E}.$$

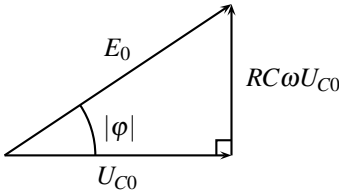


Figure 10.5 – Vecteurs de Fresnel du circuit RC.

On en déduit immédiatement les réponses aux questions posées en introduction :

- la valeur absolue du **déphasage** φ entre les signaux $e(t)$ et $u_C(t)$ se trouve par sa tangente :

$$\tan |\varphi| = \frac{RC\omega U_{C0}}{U_{C0}} = RC\omega.$$

Avec les données numériques du problème, le signal $u_C(t)$ est en retard sur $e(t)$, c'est-à-dire déphasé d'un angle : $\varphi = -\arctan(RC\omega) = -1,54 \text{ rad}$, ce qui est conforme aux observations expérimentales (le calcul et l'observation expérimentale diffèrent de $0,01 \text{ rad}$, ce qui n'est pas significatif attendu la précision de la mesure).

- le **rapport des amplitudes** des deux tensions résulte du théorème de Pythagore :

$$E_0^2 = U_{C0}^2 + (RC\omega U_{C0})^2 = U_{C0}^2 \left(1 + (RC\omega)^2 \right),$$

soit :

$$\frac{U_{C0}}{E_0} = \frac{1}{\sqrt{1 + (RC\omega)^2}}.$$

Avec les données numériques du problème : $U_{C0} = 32,9 \text{ mV}$, en accord avec les observations expérimentales.

3.2 Méthode complexe

a) Fondement de la méthode

Au signal harmonique $s(t) = S_0 \cos(\omega t + \varphi)$, on associe le **signal complexe** $\underline{s}(t)$:

$$\underline{s}(t) = S_0 \exp(j(\omega t + \varphi)).$$

Dans cette formule, le signal complexe est souligné pour affirmer son appartenance à $\mathbb{C}$ et j est le nombre complexe tel que $j^2 = -1$, noté i en mathématiques, mais qu'on évitera en physique afin de ne pas confondre avec l'intensité du courant. Le lien entre $s(t)$ et $\underline{s}(t)$ est la partie réelle :

$$s(t) = \text{Re}(\underline{s}(t)).$$

Le signal complexe se note aussi :

$$\underline{s}(t) = S_0 \exp(j\varphi) \exp(j\omega t) = \underline{S}_0 \exp(j\omega t),$$

où $\underline{S}_0 = S_0 \exp(j\varphi)$ est l'**amplitude complexe**.

CHAPITRE 10 – RÉGIME SINUSOÏDAL

b) Opérations mathématiques

Les combinaisons linéaires de signaux sinusoïdaux sont immédiates. Le signal complexe associé à :

$$s(t) = s_1(t) + s_2(t),$$

est :

$$\underline{s}(t) = \underline{s}_1(t) + \underline{s}_2(t),$$

Le signal complexe associé à la **dérivée** de $s(t) = S_0 \cos(\omega t + \varphi)$ est :

$$\frac{ds}{dt} = \frac{d}{dt} (\underline{S}_0 \exp(j\omega t)) = j\omega \times \underline{S}_0 \exp(j\omega t).$$

Dériver un signal complexe revient à le multiplier par $j\omega$.

Le signal complexe associé à une primitive du signal $s(t) = S_0 \cos(\omega t + \varphi)$ est :

$$\int \underline{s}(t) dt = \int \underline{S}_0 \exp(j\omega t) dt = \frac{\underline{S}_0}{j\omega} \exp(j\omega t).$$

Intégrer un signal complexe revient à le diviser par $j\omega$.

c) Application au circuit RC

L'équation différentielle qui régit l'évolution du circuit étudié est :

$$e(t) = RC \frac{du_C}{dt} + u_C(t).$$

Le signal complexe $\underline{e}$ associé à $e(t)$ est :

$$e(t) = E_0 \cos(\omega t) \quad \text{donc} \quad \underline{e} = E_0 \exp(j\omega t),$$

(on a arbitrairement choisit un déphasage nul pour $e(t)$, ce qui est toujours possible en décalant l'origine des temps) et $\underline{u_C}$ celui associé à $u_C(t)$:

$$u_C(t) = E_0 \cos(\omega t + \varphi) \quad \text{donc} \quad \underline{u_C} = E_0 \exp(j(\omega t + \varphi)).$$

L'équation différentielle devient :

$$\underline{e} = jRC\omega \underline{u_C} + \underline{u_C} = (1 + jRC\omega) \underline{u_C}.$$

On en déduit immédiatement les réponses aux questions posées en introduction.

Le **déphasage** φ entre $e(t)$ et $u_C(t)$ se définit par l'argument de $\underline{e} = (1 + jRC\omega) \underline{u_C}$, qui est :

$$\arg(\underline{e}) = \arg(1 + jRC\omega) + \arg(\underline{u_C}),$$

où :

$$\arg(\underline{e}) = \arg(E_0 \exp(j\omega t)) = \omega t,$$

et :

$$\arg(\underline{u_C}) = \arg(U_{C0} \exp(j(\omega t + \varphi))) = \omega t + \varphi.$$

Quant au complexe $\underline{z} = 1 + jRC\omega$, son argument passe par sa tangente :

$$\tan(\arg(1 + jRC\omega)) = \frac{RC\omega}{1} = RC\omega.$$

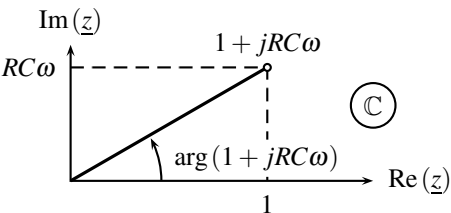


Figure 10.6 – Argumenet et module du nombre complexe $1 + jRC\omega$.

Finalement :

$$\omega t = \arctan(RC\omega) + \omega t + \varphi \quad \text{et} \quad \varphi = -\arctan(RC\omega).$$

L'angle φ est ici obtenu directement avec son signe, négatif, $u_C(t)$ est donc bien en retard par rapport à $e(t)$.

Le **rapport des amplitudes** des deux tensions passe par un calcul de module :

$$\left| \frac{\underline{u_C}}{\underline{e}} \right| = \left| \frac{U_{C0} \exp(j(\omega t + \varphi))}{E_0 \exp(j\omega t)} \right| = \frac{U_{C0}}{E_0},$$

car le module d'une exponentielle complexe vaut un. Ainsi :

$$\frac{U_{C0}}{E_0} = \left| \frac{\underline{u_C}}{\underline{e}} \right| = \left| \frac{1}{1 + jRC\omega} \right| = \frac{1}{|1 + jRC\omega|}.$$

Et le module du complexe $1 + jRC\omega$ est : $|1 + jRC\omega| = \sqrt{1^2 + (RC\omega)^2}$. Finalement :

$$\frac{U_{C0}}{E_0} = \frac{1}{\sqrt{1 + (RC\omega)^2}}.$$

4 Impédance

4.1 Impédance d'un dipôle

En régime permanent sinusoïdal, la relation entre la tension aux bornes d'un dipôle et l'intensité du courant qui le traverse s'exprime, en complexe, d'une manière aussi simple que $u_R = Ri$ pour une résistance.

En régime sinusoïdal permanent à la pulsation ω , imposée par la tension d'alimentation du montage, on définit les tensions et intensités complexes par :

$$\underline{u}(t) = \underline{U}_0 \exp(j\omega t) \quad \text{et} \quad \underline{i}(t) = \underline{I}_0 \exp(j\omega t),$$

où $\underline{U}_0$ et $\underline{I}_0$ sont les amplitudes complexes des tension et courant complexes, définies au paragraphe 3.2.

On définit alors l'**impédance** $\underline{Z}$ d'un dipôle par le rapport des signaux complexes $\frac{\underline{u}(t)}{\underline{i}(t)}$. L'unité de l'impédance est l'ohm (Ω).

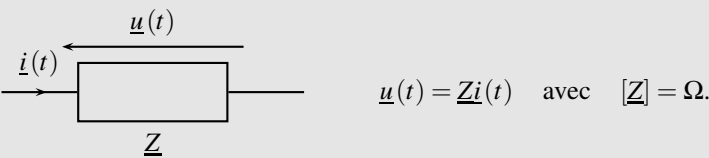


Figure 10.7 – Schéma, définition et unité d'une impédance en convention récepteur.

a) Cas d'une résistance

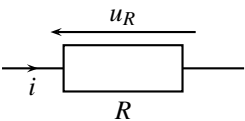


Figure 10.8 – Schéma d'une résistance en convention récepteur.

La loi d'Ohm s'écrit $\underline{u}_R(t) = R \underline{i}(t)$. Immédiatement :

L'impédance, en Ω , d'une résistance est $\underline{Z}_R = R$.

b) Cas d'une bobine

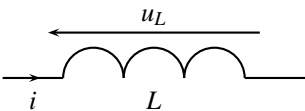


Figure 10.9 – Schéma d'une bobine d'inductance L en convention récepteur.

La relation entre $u_L(t)$ et $i(t)$ est $u_L = L \frac{di}{dt}$. Elle devient pour les signaux complexes :

$$\underline{u}_L = jL\omega \underline{i}.$$

L'impédance, en Ω , d'une bobine d'inductance L est $\underline{Z}_L = jL\omega$.

c) Cas d'un condensateur

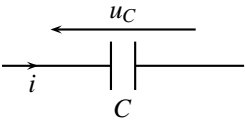


Figure 10.10 – Schéma d'un condensateur de capacité C en convention récepteur.

La relation entre $u_C(t)$ et $i(t)$ est $i = C \frac{du_C}{dt}$. Elle devient pour les signaux complexes :

$$\underline{u}_C = \frac{\underline{i}}{jC\omega}.$$

L'impédance, en Ω , d'un condensateur de capacité C est $\underline{Z}_C = \frac{1}{jC\omega}$.

d) Comportements limites d'un condensateur et d'une bobine

Condensateur On considère un condensateur, inséré dans un montage électrique en régime permanent constant. Toutes les tensions y sont constantes, en particulier u_C .

On a alors :

$$i = C \frac{du_C}{dt} = 0.$$

Ou bien, en régime sinusoïdal permanent, avec $\omega = 0$ (une tension de pulsation nulle est bien une tension constante car $\cos(0 \times t) = \text{constante}$) :

$$\underline{u}_C = \frac{\underline{i}}{jC\omega} \quad \text{donc} \quad \underline{i} = j\omega \underline{u}_C = 0.$$

Qu'en est-il en haute fréquence ? $\underline{Z}_C = \frac{1}{jC\omega}$ tend vers 0 quand ω tend vers l'infini. On retrouve l'impédance nulle d'un fil parfait, sans résistance.

En « basse fréquence », un condensateur est équivalent à un interrupteur ouvert.
En « haute fréquence », un condensateur est équivalent à un fil.

CHAPITRE 10 – RÉGIME SINUSOÏDAL

Bobine En régime sinusoïdal permanent, $\underline{Z}_L = jL\omega$. Cette impédance tend vers 0 quand ω tend vers 0, comme pour un fil parfait.

$\underline{Z}_L = jL\omega$ diverge quand ω tend vers 0, donc le courant $\underline{i} = \frac{\underline{u}_L}{\underline{Z}_L}$ qui la traverse est nul.

En « basse fréquence », une bobine est équivalente à un fil.
En « haute fréquence », une bobine est équivalente à un interrupteur ouvert.

 Une fréquence n'est jamais « haute » ou « basse » en elle-même ; elle est « haute » ou « basse » devant une autre fréquence de référence.

e) Association de deux impédances

La relation qui relie les tensions complexes et les intensités complexes aux bornes d'un dipôle, $\underline{u} = \underline{Z}\underline{i}$, est analogue à la loi d'Ohm pour une résistance $u_R = Ri$. Les lois d'association des impédances sont donc analogues à celle des résistances.

Association en série On envisage le cas d'une association de deux dipôles en série.

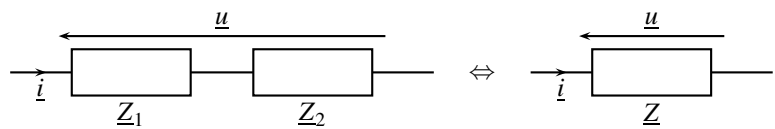


Figure 10.11 – Association d'impédances en série.

Les impédances $\underline{Z}_1$ et $\underline{Z}_2$ représentent des résistances, des bobines ou des condensateurs :

$$\underline{u} = \underline{u}_1 + \underline{u}_2 = \underline{Z}_1 \underline{i} + \underline{Z}_2 \underline{i} = (\underline{Z}_1 + \underline{Z}_2) \underline{i} = \underline{Z} \underline{i},$$

où $\underline{Z}$ est l'impédance équivalente du dipôle série.

Les impédances en série se somment : $\underline{Z} = \underline{Z}_1 + \underline{Z}_2$.

Association en parallèle (voir figure 10.12)

La tension $\underline{u}$ est la même aux bornes des deux impédances $\underline{Z}_1$ et $\underline{Z}_2$:

$$\underline{u} = \underline{Z}_1 \underline{i}_1 = \underline{Z}_2 \underline{i}_2.$$

Ainsi : $\underline{i}_1 = \frac{\underline{u}}{\underline{Z}_1}$ et $\underline{i}_2 = \frac{\underline{u}}{\underline{Z}_2}$. La loi des nœuds devient alors :

$$\underline{i} = \underline{i}_1 + \underline{i}_2 = \left(\frac{1}{\underline{Z}_1} + \frac{1}{\underline{Z}_2} \right) \underline{u} = \frac{1}{\underline{Z}} \underline{u},$$

où $\underline{Z}$ est l'impédance équivalente des deux dipôles en parallèle.

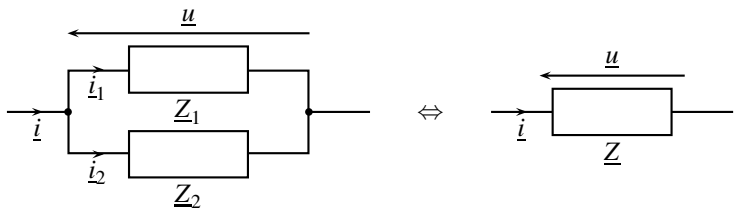


Figure 10.12 – Association d'impédances en parallèle.

L'inverse de l'impédance équivalente à deux impédances $\underline{Z}_1$ et $\underline{Z}_2$ en parallèle est la somme des inverses de $\underline{Z}_1$ et $\underline{Z}_2$: $\frac{1}{\underline{Z}} = \frac{1}{\underline{Z}_1} + \frac{1}{\underline{Z}_2}$.

4.2 Diviseur de tension

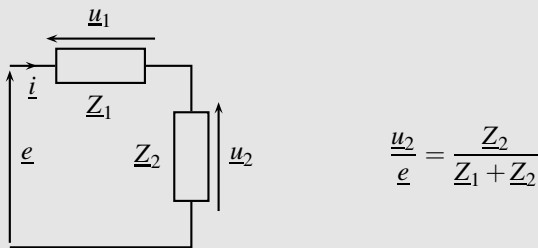


Figure 10.13 – Schéma et formule du diviseur de tension en régime permanent sinusoïdal.

La formule du diviseur de tension est encore valable avec les impédances, avec la même hypothèse d'un courant identique dans les deux dipôles.

Attendu que le courant d'intensité $\underline{i}$ est le même dans les deux impédances $\underline{Z}_1$ et $\underline{Z}_2$:

$$\underline{u}_1 = \underline{Z}_1 \underline{i} \quad \text{et} \quad \underline{u}_2 = \underline{Z}_2 \underline{i}.$$

La loi des mailles s'écrit :

$$\underline{e} = \underline{u}_1 + \underline{u}_2 = \underline{Z}_1 \underline{i} + \underline{Z}_2 \underline{i} = (\underline{Z}_1 + \underline{Z}_2) \underline{i}.$$

Par élimination de $\underline{i}$ dans les expressions de $\underline{e}$ et $\underline{u}_2$:

$$\underline{i} = \frac{\underline{u}_2}{\underline{Z}_2} = \frac{\underline{e}}{\underline{Z}_1 + \underline{Z}_2} \quad \text{d'où} \quad \frac{\underline{u}_2}{\underline{e}} = \frac{\underline{Z}_2}{\underline{Z}_1 + \underline{Z}_2}.$$

5 Résonance dans un système du deuxième ordre

Le circuit étudié dans ce paragraphe est un circuit *RLC* série, avec les valeurs des composants :

$R = 1,0\text{ k}\Omega$
 $L = 20\text{ mH}$
 $C = 51\text{ nF}$

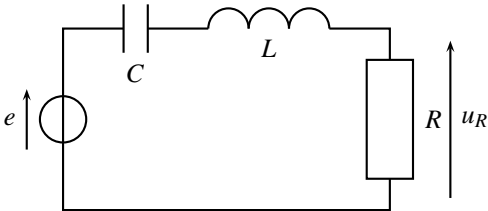


Figure 10.14 – Circuit *RLC* étudié.

5.1 Étude expérimentale de u_R

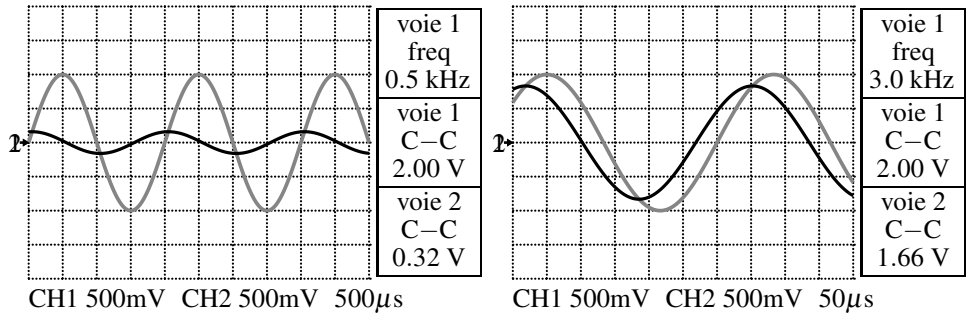


Figure 10.15 – Formes d’onde des tensions e , en gris sur la voie 1, et u_R , en noir sur la voie 2 ($f = 0,5\text{ kHz}$ et $f = 3,0\text{ kHz}$).

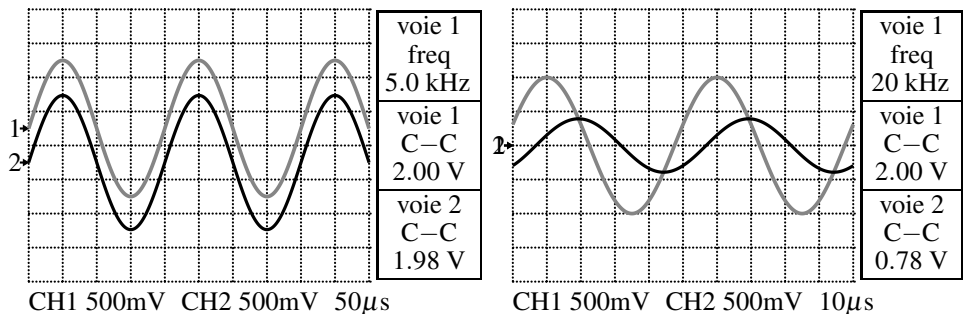


Figure 10.16 – Formes d’onde des tensions e , en gris sur la voie 1, et u_R , en noir sur la voie 2 ($f = 5,0\text{ kHz}$ et $f = 20\text{ kHz}$).

La forme d’onde de la tension u_R , aux bornes de la résistance, dépend de la fréquence f du signal d’entrée. On observe que pour $f < 5\text{ kHz}$, la tension u_R est en avance sur e alors que pour $f > 5\text{ kHz}$, u_R est en retard ; elles sont en phase en $f = 5\text{ kHz}$.

De plus, en mode XY de l’oscilloscope, on observe une ellipse pour une fréquence d’alimentation différente de 5 kHz, réduite à une droite lorsque $f = f_0 = 5 \text{ kHz}$ (voir page 38).

L’amplitude de u_R croît lorsqu’on se rapproche de $f = f_0 = 5 \text{ kHz}$, où elle vaut presque celle de l’entrée ; elle est d’autant plus faible que f est éloigné de f_0 . On visualise ce phénomène en traçant le rapport de l’amplitude U_{R0} de la tension u_R et de son amplitude maximale U_{R0max} en $f = f_0$, en fonction du rapport de la pulsation $\omega = 2\pi f$ à la pulsation $\omega_0 = 2\pi f_0$ (voir figure 10.17). On observe que l’allure de la courbe dépend fortement de la valeur de R , c’est-à-dire

de la valeur du **facteur de qualité** $Q = \frac{1}{R} \sqrt{\frac{L}{C}}$, introduit au chapitre précédent.

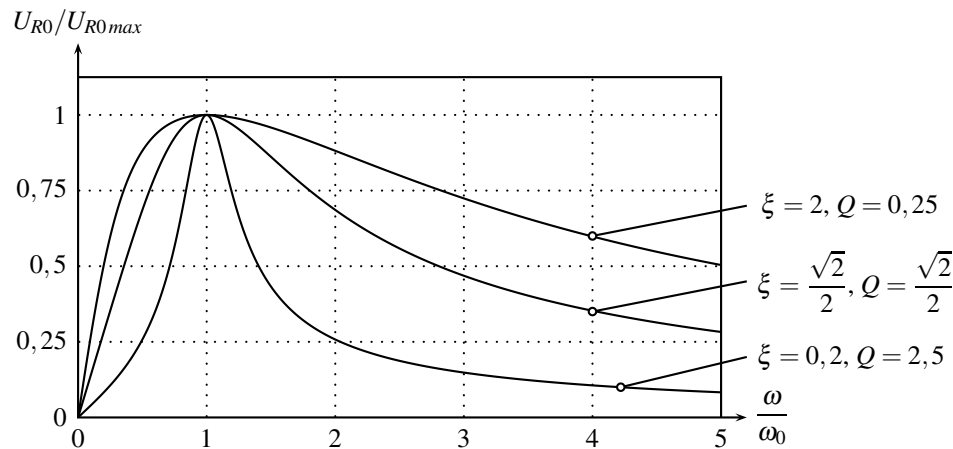


Figure 10.17 – Rapports de l’amplitude de u_R à sa valeur maximale (R variable).

Le passage de l’amplitude de u_R par un maximum est appelé **résonance**. Sur le graphe de la figure 10.17, plus le pic est mince autour de ω_0 , plus la résonance est dite aigue. L’acuité de la résonance dépend donc de Q ou de ξ .

D’après la loi d’Ohm, l’intensité i du courant qui traverse la résistance est reliée à u_R par $u_R = Ri$. Ainsi, lorsque u_R est maximum, i l’est aussi. La résonance est donc aussi nommée résonance en courant.

Il y a **résonance** quand l’amplitude du signal passe par un maximum pour une certaine pulsation, nommée **pulsation de résonance**.

Remarque

Lors de l’étude de u_R , la pulsation de résonance est $f_0 = 5 \text{ kHz}$, avec les valeurs numériques des composants précisées en introduction.

CHAPITRE 10 – RÉGIME SINUSOÏDAL

5.2 Interprétation

On se propose de calculer l'expression complexe $\underline{u_R}$ avec la formule du diviseur de tension, appliquée au schéma de la figure 10.14 :

$$\frac{\underline{u_R}}{\underline{e}} = \frac{\underline{Z_R}}{\underline{Z_R} + \underline{Z_L} + \underline{Z_C}} = \frac{R}{R + jL\omega + \frac{1}{jC\omega}} = \frac{1}{1 + j\left(\frac{L}{R}\omega - \frac{1}{RC\omega}\right)}.$$

Les signaux complexes associés à $u_R(t) = u_{R0} \cos(\omega t + \varphi)$ et $e(t) = E_0 \cos(\omega t)$ sont $\underline{u_R}(t) = U_{R0} \exp(j(\omega t + \varphi))$ et $\underline{e}(t) = E_0 \exp(j\omega t)$. Le module de leur rapport est donc : $\left| \frac{\underline{u_R}}{\underline{e}} \right| = \frac{U_{R0}}{E_0}$. Ainsi, U_{R0} sera-t-il d'autant plus important que ce module sera grand. Le module passe par un maximum quand la partie imaginaire du dénominateur est minimale, c'est-à-dire quand : $\frac{L}{R}\omega - \frac{1}{RC\omega} = 0$, soit en :

$$\omega = \frac{1}{\sqrt{LC}} = \omega_0.$$

Il y a donc résonance en ω_0 . On retrouve la valeur de la pulsation caractéristique d'un système du deuxième ordre, introduite dans le chapitre précédent. On en déduit une méthode expérimentale de mesure de la pulsation caractéristique :

La résonance en courant se produit à la pulsation caractéristique ω_0 du circuit.

De plus, le rapport $\left| \frac{\underline{u_R}}{\underline{e}} \right|$ dépend de la valeur du facteur de qualité Q . Sur le graphe de la figure 10.17, on observe que plus R est faible, à L et C constants, c'est-à-dire plus Q est important, plus la résonance est aigue.

Moins le système est amorti (ξ faible ou $Q = \frac{1}{2\xi}$ important), plus l'acuité de la résonance est forte.

5.3 Complément : interprétation graphique du facteur de qualité

L'affirmation que plus le facteur de qualité Q est important, plus l'acuité de la résonance est importante, c'est-à-dire plus le pic est étroit autour de ω_0 , peut être expérimentalement observée. En effet, les mesures expérimentales, pour plusieurs valeurs de R , donc de Q , mènent systématiquement à :

$$Q = \frac{\omega_0}{\Delta\omega},$$

où $\Delta\omega$ est la bande de fréquence indiquée sur la figure 10.18.

On comprend donc que plus Q est grand, plus la bande de fréquence $\Delta\omega$ est faible, c'est-à-dire plus la résonance est aigue.

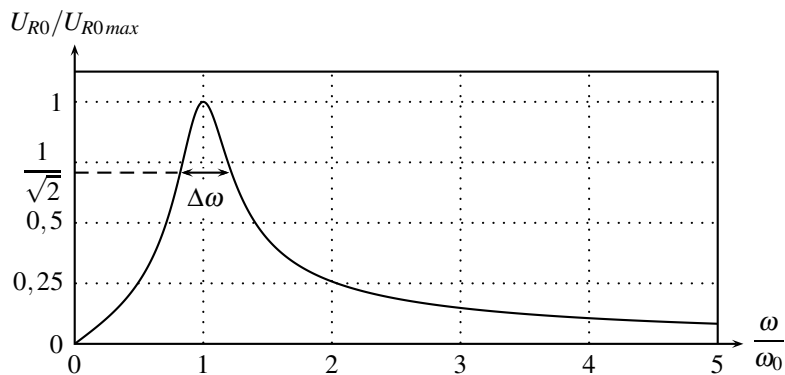


Figure 10.18 – Rapports de l’amplitude de u_R à sa valeur maximale ($Q = 2,5$).

5.4 Remarque expérimentale

À la fréquence de résonance, l’amplitude de la tension u_R ne vaut pas exactement celle du générateur (voir figure 10.16 pour $f = 5\text{ kHz}$), mais lui est inférieure. Pour modéliser précisément la réalité expérimentale, il convient de prendre en compte la résistance r de la bobine, qui est alors modélisée par un inductance L en série avec la résistance r .

L’association série du condensateur de capacité C et de l’inductance L est équivalente en $\omega_0 = \frac{1}{\sqrt{LC}}$ à un fil ; en effet : $\underline{Z}_C + \underline{Z}_L = \frac{1}{jC\omega_0} + jL\omega_0 = \sqrt{\frac{L}{C}}(-j + j) = 0$.

Le circuit est donc équivalent, à la pulsation de résonance, à un diviseur de tension constitué uniquement des résistances R et r :

$$\frac{u_R}{e} = \frac{R}{R + r} < 1.$$

Avec les valeurs expérimentales de la figure 10.16, on évalue la résistance r de la bobine à :

$$r = R \left(\frac{E_0}{U_{R0}} - 1 \right) = 10\ \Omega.$$

5.5 Étude de u_C

On observe expérimentalement que l’amplitude U_{C0} de la tension aux bornes du condensateur dépend de la pulsation ω et donc de la fréquence ($\omega = 2\pi f$) du signal d’entrée, de même pour son déphasage φ . Pour certaines valeurs de f et de ξ , le signal de sortie présente un maximum. C’est le phénomène de **résonance**.

Toutefois, lorsqu’on trace le rapport de l’amplitude U_{C0} de u_C sur l’amplitude E_0 de l’entrée e , en échelles linéaires ou logarithmiques, on observe que le phénomène de résonance n’existe pas toujours.

L’allure des courbes de la figure 10.19 dépend encore de la valeur de ξ , ou de Q . Encore une fois, plus ξ est faible, ou Q grand, plus le phénomène de résonance est marqué.

CHAPITRE 10 – RÉGIME SINUSOÏDAL

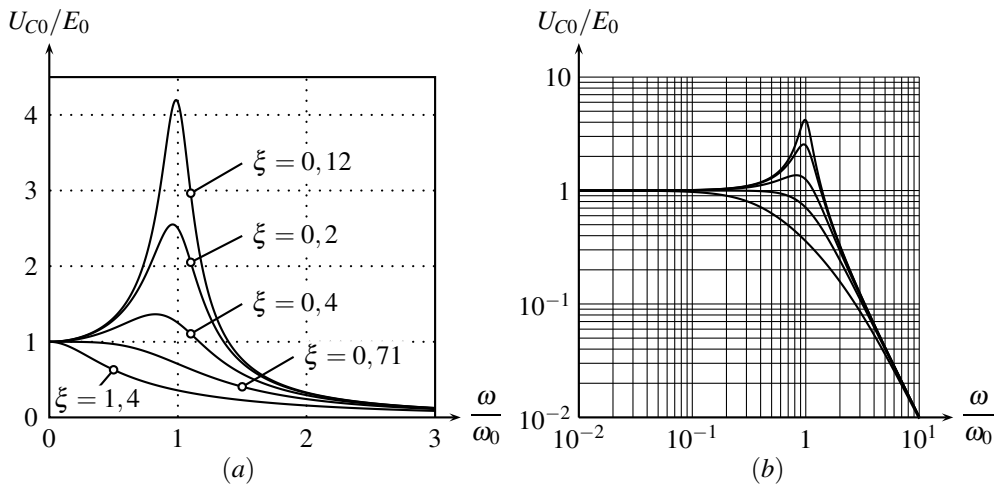


Figure 10.19 – Étude de la résonance aux bornes du condensateur : (a) échelles linéaires, (b) échelles logarithmiques.

La résonance aux bornes du condensateur ne se produit que pour des systèmes faiblement amortis : $\xi < \frac{\sqrt{2}}{2} \simeq 0,71$ ou $Q = \frac{1}{2\xi} > \frac{\sqrt{2}}{2} \simeq 0,71$.

Moins le système est amorti (ξ faible ou $Q = \frac{1}{2\xi}$ important), plus l'acuité de la résonance est forte.

L'allure du déphasage φ dépend aussi de ξ , ou de Q , comme le montre la figure 10.20 en échelles semi-logarithmiques où on trace φ en fonction de $\log(\omega/\omega_0)$.

Le déphasage part de 0 pour $\omega \ll \omega_0$, passe par $-\pi/2$ en $\omega = \omega_0$, ce qui est une propriété importante qui permet de mesurer ω_0 , puis converge vers $-\pi$ pour $\omega \gg \omega_0$.

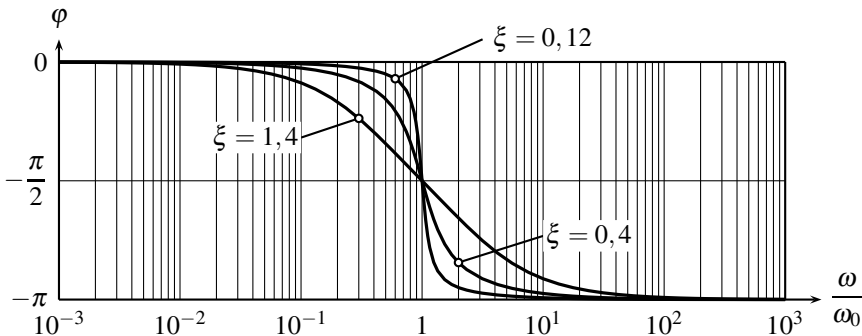


Figure 10.20 – Étude de la phase aux bornes du condensateur.

5.6 Interprétation

Le circuit est modifié afin d'étudier u_C . En effet, il ne peut exister qu'une seule masse dans le circuit. Il convient donc que la masse, imposée par l'alimentation et les moyens d'observation, soit identique pour la mesure de e et celle de u_C .

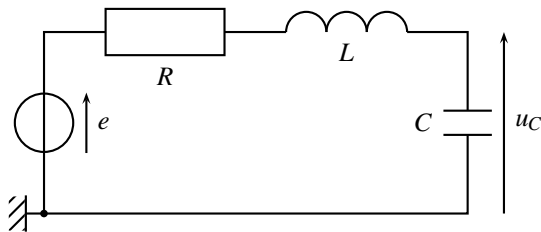


Figure 10.21 – Circuit d'étude de la tension u_C .

La formule du diviseur de tension mène, en notation complexe, à :

$$\frac{\underline{u_C}}{\underline{e}} = \frac{\underline{Z_C}}{\underline{Z_C} + \underline{Z_L} + \underline{Z_R}} = \frac{\frac{1}{jC\omega}}{\frac{1}{jC\omega} + R + jL\omega} = \frac{1}{1 + jRC\omega + LC(j\omega)^2},$$

qui se met sous la forme, avec les notations du chapitre précédent, $\omega_0 = \frac{1}{\sqrt{LC}}$ et $\xi = \frac{R}{2} \sqrt{\frac{C}{L}}$:

$$\frac{\underline{u_C}}{\underline{e}} = \frac{1}{1 + 2j\xi \frac{\omega}{\omega_0} + \left(j \frac{\omega}{\omega_0}\right)^2}.$$

En notant U_{C0} et E_0 les amplitudes des tensions, le module de cette expression est :

$$\left| \frac{\underline{u_C}}{\underline{e}} \right| = \frac{U_{C0}}{E_0} = \frac{1}{\sqrt{\left(1 - \frac{\omega^2}{\omega_0^2}\right)^2 + \left(2\xi \frac{\omega}{\omega_0}\right)^2}}.$$

À quelle condition U_{C0} passe-t-elle par un maximum ? Attendu que le numérateur est une constante, l'amplitude est maximale quand le dénominateur est minimum. Il est minimum quand son carré l'est aussi. Avec la variable $X = \frac{\omega^2}{\omega_0^2}$, le carré du dénominateur est :

$$D(X) = (1 - X)^2 + 4\xi^2 X.$$

Un extremum est obtenu quand la dérivée est nulle :

$$\frac{dD}{dX} = -2(1 - X) + 4\xi^2 = 0 \quad \text{implique} \quad X = 1 - 2\xi^2.$$

CHAPITRE 10 – RÉGIME SINUSOÏDAL

L'extremum de l'amplitude est alors obtenu pour la pulsation de résonance ω_r :

$$\omega_r = \omega_0 \sqrt{1 - 2\xi^2} < \omega_0.$$

ω_r n'existe que si $1 - 2\xi^2 > 0$, c'est-à-dire $\xi < \frac{\sqrt{2}}{2}$. On retrouve le point fondamental : il ne peut y avoir résonance aux bornes du condensateur que si le système est peu amorti.

5.7 Mesures expérimentales

Comment mesurer ξ (ou Q) et ω_0 à partir de graphes expérimentaux sur la tension u_C ?

a) Coefficient d'amortissement ou facteur de qualité

Si l'on observe une résonance, les coordonnées du maximum renseignent sur la valeur de ξ ou de Q . En effet, la résonance se produit pour la pulsation $\omega_r = \omega_0 \sqrt{1 - 2\xi^2}$. L'amplitude S_0 du signal de sortie y vaut :

$$S_0(\omega_r) = \frac{E_0}{\sqrt{\left(1 - \frac{\omega_r^2}{\omega_0^2}\right)^2 + \left(2\xi \frac{\omega_r}{\omega_0}\right)^2}}$$
$$S_0(\omega_r) = \frac{E_0}{\sqrt{4\xi^4 + 4\xi^2(1 - 2\xi^2)}} = \frac{E_0}{\sqrt{4\xi^2 - 4\xi^4}} = \frac{E_0}{2\xi\sqrt{1 - \xi^2}}.$$

La mesure de $S_0(\omega_r)$, alliée à la connaissance de E_0 , mène ainsi à la valeur de ξ donc de $Q = \frac{1}{2\xi}$.

Sans résonance, on ne peut qu'affirmer $Q < \frac{\sqrt{2}}{2}$ ou $\xi > \frac{\sqrt{2}}{2}$. Il faut alors déterminer ω_0 avec la méthode du paragraphe suivant, puis extraire la valeur de ξ ou de Q , en partant de la valeur de l'amplitude pour une ou plusieurs pulsations.

b) Pulsation caractéristique

La mesure expérimentale de ω_0 s'effectue très simplement. Pour des systèmes du deuxième ordre, la phase passe par l'angle moitié en ω_0 . On entend par angle moitié la moyenne entre $\varphi(0)$ et $\varphi(\infty)$. La lecture directe de la courbe de phase renseigne immédiatement sur ω_0 .

Il est possible d'utiliser le diagramme d'amplitude pour des résonances très aigües. Dans ce cas, le système est très peu amorti, $\xi \ll 1$ ou $Q \gg 1$ et la résonance a lieu en $\omega_r = \omega_0 \sqrt{1 - 2\xi^2} \simeq \omega_0$.

Les diagrammes d'amplitude et de phase sont donc complémentaires. L'amplitude permet d'immédiatement voir l'existence d'une résonance et les valeurs de ξ et Q ; la phase la valeur de ω_0 .

6 Fonction de transfert

6.1 Fonction de transfert harmonique

La **fonction de transfert**, ou **transmittance**, d'un système, est le lien entre le signal de sortie $\underline{s}(t)$ et celui d'entrée $\underline{e}(t)$, en régime permanent sinusoïdal à la pulsation ω , imposée par l'entrée :

$$\underline{H}(j\omega) = \frac{\underline{s}(t)}{\underline{e}(t)}.$$

Si le système étudié est le circuit RC , présenté au paragraphe 2.2 :

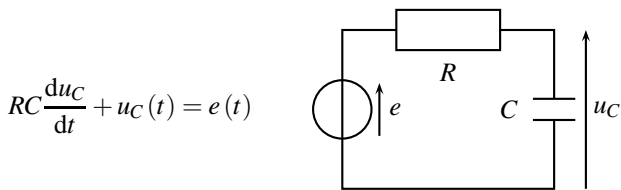


Figure 10.22 – Système du premier ordre étudié et équation différentielle.

l'équation différentielle devient, en régime sinusoïdal permanent :

$$\underline{u}_C(jRC\omega + 1) = \underline{e}.$$

On définit alors la fonction de transfert harmonique $\underline{H}(j\omega)$ par :

$$\underline{H}(j\omega) = \frac{\underline{u}_C}{\underline{e}} = \frac{1}{jRC\omega + 1}.$$

L'expression de la transmittance pour le circuit RC peut être facilement établie avec le formalisme des impédances. Avec la formule du diviseur de tension :

$$\frac{\underline{u}_C}{\underline{e}} = \frac{\underline{Z}_C}{\underline{Z}_R + \underline{Z}_C} = \frac{\frac{1}{jC\omega}}{R + \frac{1}{jC\omega}} = \frac{1}{jRC\omega + 1}.$$

On retrouve ainsi très rapidement le lien entre $\underline{u}_C$ et $\underline{e}$ en régime permanent sinusoïdal, sans avoir besoin d'établir préalablement l'équation différentielle.

6.2 Lien entre équation différentielle et transmittance

Le passage de l'équation différentielle à la transmittance a été établi dans un exemple au paragraphe 3.2. L'opération inverse, qui consiste à passer de la transmittance à l'équation différentielle est tout aussi aisée. Par exemple, on cherche à déterminer l'équation différentielle qui relie la tension de sortie $s(t)$ à celle d'entrée $e(t)$ dans le montage suivant :

CHAPITRE 10 – RÉGIME SINUSOÏDAL

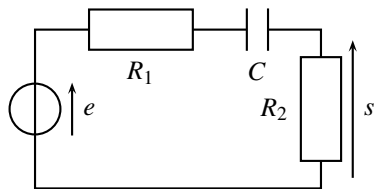


Figure 10.23 – Système du premier ordre étudié.

On se place formellement en régime permanent sinusoïdal pour établir la transmittance. La formule du diviseur de tension mène à :

$$\frac{\underline{s}}{\underline{e}} = \frac{Z_{R2}}{Z_{R1} + \underline{Z}_C + Z_{R2}} = \frac{R_2}{R_1 + \frac{1}{jC\omega} + R_2} = \frac{jR_2C\omega}{jR_1C\omega + 1 + jR_2C\omega} = \frac{jR_2C\omega}{1 + j(R_1 + R_2)C\omega}.$$

Le résultat est présenté sous la forme d’un polynôme en $j\omega$ au numérateur, sur un polynôme en $j\omega$ au dénominateur afin d’écrire simplement le lien entre $\underline{e}$ et $\underline{s}$:

$$j(R_1 + R_2)C\omega \underline{s} + \underline{s} = jR_2C\omega \underline{e}.$$

Puis on utilise la correspondance entre $j\omega$ et la dérivée d’un signal en régime permanent sinusoïdal :

$$j\omega \underline{s} \longleftrightarrow \frac{ds}{dt} \quad \text{et} \quad (j\omega)^2 \underline{s} \longleftrightarrow \frac{d^2s}{dt^2}.$$

Ainsi :

$$(R_1 + R_2)C \frac{ds}{dt} + s = R_2C \frac{de}{dt}.$$

6.3 Lien avec la transmittance de Laplace

On définit dans le cours de SII une fonction de transfert, issu d’une transformation intégrale nommée transformée de Laplace. Cette opération ne relève pas du présent cours, mais l’identité des résultats doit être soulignée. Sur l’exemple du paragraphe précédent :

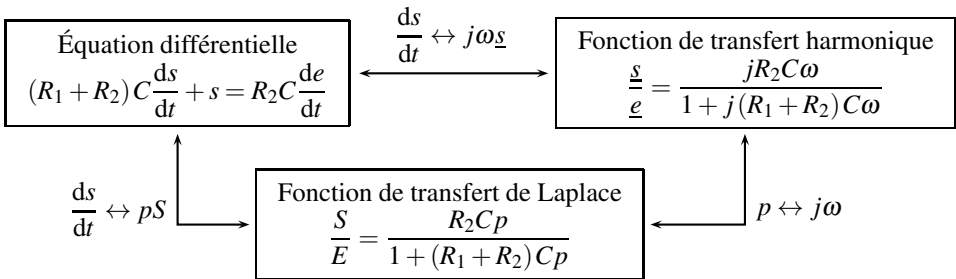


Figure 10.24 – Liens entre équation différentielle, transmittance harmonique et transmittance de Laplace.

Remarque

L'équivalence entre la dérivée et la transformée de Laplace est $\frac{ds}{dt} \leftrightarrow pS(p) - s(0^+)$.
L'usage est toutefois de considérer les conditions initiales nulles pour écrire la transmittance.

7 Complément : déphasage et rapport des amplitudes dans un système du deuxième ordre

7.1 Position du problème

On étudie le circuit RLC (voir figure 10.25) décrit par les valeurs de ses composants :

$R = 1,0 \text{ k}\Omega, \quad L = 20 \text{ mH} \quad \text{et} \quad C = 51 \text{ nF},$

ou ses paramètres canoniques :

$\omega_0 = \frac{1}{\sqrt{LC}} = 2\pi \times 5,0.10^3 \text{ rad.s}^{-1} \quad \text{et} \quad \xi = \frac{R}{2} \sqrt{\frac{C}{L}} = 0,80.$

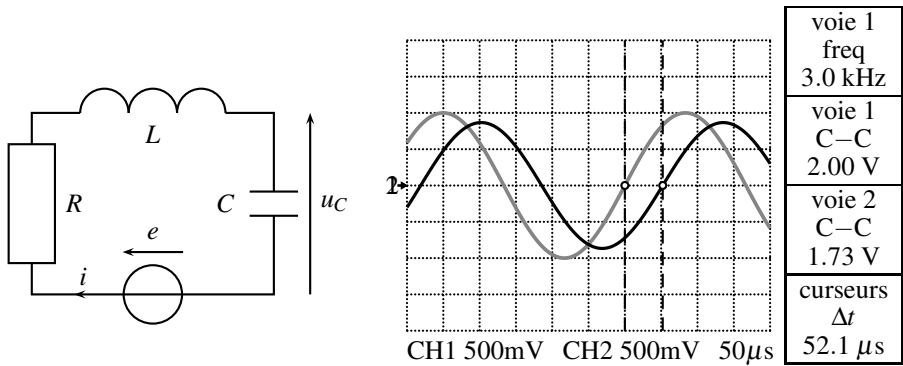


Figure 10.25 – Montage du deuxième ordre étudié et formes d'ondes des tensions.
En gris, la voie 1, en noir, la voie 2.

La tension $e(t)$, sur la voie 1 en gris, et la tension $u_C(t)$, sur la voie 2 en noir. Deux curseurs verticaux marquent les dates successives pour lesquelles les deux tensions passent par 0 avec des pentes de même signe, afin de mesurer le déphasage entre les signaux.

La tension $u_C(t)$ est en retard sur $e(t)$. Les curseurs nous permettent de définir le déphasage entre les deux tensions :

$\varphi = 2\pi \frac{\Delta t}{T} = 2\pi f \Delta t = 0,98 \text{ rad} = 56,3^\circ.$

CHAPITRE 10 – RÉGIME SINUSOÏDAL

7.2 Méthode des vecteurs de Fresnel (MPSI)

On cherche à prévoir *a priori* le déphasage entre les deux signaux $e(t)$ et $u_C(t)$ ainsi que le rapport de leurs amplitudes en régime permanent sinusoïdal. L'équation différentielle, liant $e(t)$ à $u_C(t)$, qui régit l'évolution du circuit est :

$$LC \frac{d^2 u_C}{dt^2} + RC \frac{du_C}{dt} + u_C(t) = e(t),$$

ou, sous forme canonique :

$$\frac{1}{\omega_0^2} \frac{d^2 u_C}{dt^2} + \frac{2\xi}{\omega_0} \frac{du_C}{dt} + u_C(t) = e(t).$$

On associe au signal $u_C(t)$ le vecteur de Fresnel $\vec{U}_C$, de norme U_{C0} , à sa dérivée $\frac{du_C}{dt}$, le vecteur $\vec{DU}_C$ de norme ωU_{C0} , tourné de $\frac{\pi}{2}$ et à sa dérivée seconde $\frac{d^2 u_C}{dt^2}$, le vecteur $\vec{D_2U}_C$, de norme $\omega \times \omega U_{C0}$, tourné de $\frac{\pi}{2}$ par rapport à $\vec{DU}_C$, c'est-à-dire de π par rapport à $\vec{U}_C$. L'équation différentielle du système se résume alors à :

$$LC \times \vec{D_2U}_C + RC \times \vec{DU}_C + \vec{U}_C = \vec{E}.$$

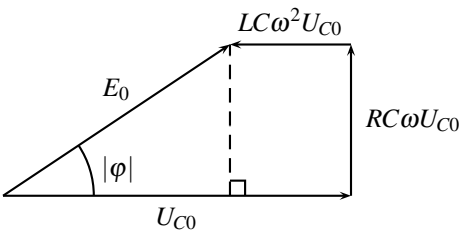


Figure 10.26 – Vecteurs de Fresnel du circuit RLC.

Remarque : avec les valeurs numériques, $LC\omega^2 = 0,36 < 1$, donc le vecteur de Fresnel $\vec{D_2U}_C$, associé à $\frac{d^2 u_C}{dt^2}$, est de norme inférieure à celle de $\vec{U}_C$, comme sur le schéma.

On en déduit, dans le triangle rectangle :

- la valeur absolue du **déphasage** entre les tensions $e(t)$ et $u_C(t)$:

$$\tan |\varphi| = \frac{RC\omega U_{C0}}{U_{C0} - LC\omega^2 U_{C0}} = \frac{RC\omega}{1 - LC\omega^2}.$$

$u_C(t)$ est en retard sur $e(t)$ de $|\varphi| = 0,985 \text{ rad} = 56,4^\circ$. L'écart de $0,1^\circ$ avec l'expérience n'est pas significatif attendu la précision des mesures.

- le **rapport des amplitudes** avec le théorème de Pythagore :

$$E_0^2 = (U_{C0} - LC\omega^2 U_{C0})^2 + (RC\omega U_{C0})^2,$$

soit :

$$\frac{U_{C0}}{E_0} = \frac{1}{\sqrt{(1 - LC\omega^2)^2 + (RC\omega)^2}}.$$

$\frac{U_{C0}}{E_0} = 0,87$, ce que confirme les relevés expérimentaux.

7.3 Méthode complexe

Le signal complexe $\underline{u}_C$ associé à $u_C(t) = U_{C0} \cos(\omega t + \varphi)$ est :

$$\underline{u}_C = U_{C0} \exp(j\varphi) \exp(j\omega t) = \underline{U}_{C0} \exp(j\omega t).$$

Le signal complexe associé à la dérivée seconde est alors $(j\omega)^2 \underline{u}_C$:

$$\frac{d^2 \underline{u}_C}{dt^2} = \frac{d^2}{dt^2} (\underline{U}_{C0} \exp(j\omega t)) = (j\omega)^2 \underline{U}_{C0} \exp(j\omega t) = (j\omega)^2 \underline{u}_C.$$

Ainsi, l'équation différentielle $LC \frac{d^2 u_C}{dt^2} + RC \frac{du_C}{dt} + u_C(t) = e(t)$ devient :

$$LC (j\omega)^2 \underline{u}_C + RC (j\omega) \underline{u}_C + \underline{u}_C = \underline{e},$$

soit :

$$\underline{u}_C = \frac{\underline{e}}{1 - LC\omega^2 + jRC\omega}.$$

Cette dernière relation aurait pu être obtenue avec la formule du diviseur de tension. Elle permet de calculer le déphasage et le rapport des amplitudes entre $u_C(t)$ et $e(t)$. Encore une fois, on choisit arbitrairement un déphasage nul pour $e(t)$, possible en décalant l'origine des temps. Les signaux complexes sont alors :

$$\begin{cases} e(t) = E_0 \cos(\omega t) & \text{donc } \underline{e} = E_0 \exp(j\omega t), \\ u_C(t) = U_{C0} \cos(\omega t + \varphi) & \text{donc } \underline{u}_C = U_{C0} \exp(j(\omega t + \varphi)). \end{cases}$$

On en déduit :

- le **déphasage** entre les tension $e(t)$ et $u_C(t)$ via l'argument :

$$\arg(\underline{u}_C) = \arg(\underline{e}) - \arg(1 - LC\omega^2 + jRC\omega).$$

Ainsi :

$$\omega t + \varphi = \omega t - \arctan\left(\frac{RC\omega}{1 - LC\omega^2}\right),$$

et :

$$\varphi = -\arctan\left(\frac{RC\omega}{1 - LC\omega^2}\right).$$

CHAPITRE 10 – RÉGIME SINUSOÏDAL

- le **rapport des amplitudes** *via* le module :

$$|u_C| = \left| \frac{e}{1 - LC\omega^2 + jRC\omega} \right| = \frac{|e|}{|1 - LC\omega^2 + jRC\omega|}.$$

Ainsi :

$$U_{C0} = \frac{E_0}{\sqrt{(1 - LC\omega^2)^2 + (RC\omega)^2}}.$$

Ces résultats sont identiques à ceux de la méthode des vecteurs de Fresnel.

SYNTHÈSE

SAVOIRS

- notion de résonance
- impédance d’une résistance, d’une bobine et d’un condensateur
- influence de Q ou de ξ sur la résonance

SAVOIR-FAIRE

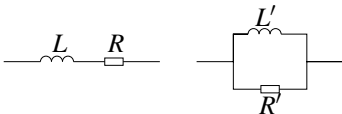
- mesurer un déphasage et un rapport d’amplitude entre deux signaux
- utiliser la construction de Fresnel pour étudier le régime permanent sinusoïdal (MPSI)
- utiliser les complexes pour étudier le régime permanent sinusoïdal
- relier l’existence et l’acuité de la résonance d’un deuxième ordre à ξ ou Q
- déterminer ω_0 , ξ et Q à partir de graphes d’amplitude et de phase
- établir l’impédance d’une résistance, d’un condensateur, d’une bobine, en régime harmonique
- calculer une impédance équivalente pour une association série ou parallèle

MOTS-CLÉS

- | | | |
|-------------------------------|--------------------------|-----------------------------|
| • régime permanent sinusoïdal | • facteur de qualité Q | • méthode de Fresnel (MPSI) |
| • résonance | • impédance | |
| | • fonction de transfert | |

S'ENTRAÎNER

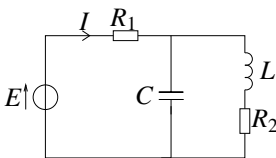
10.1 Équivalence entre deux dipôles (★)



On travaille en régime sinusoïdal de pulsation ω .

1. Pour quelles valeurs de R' et L' , a-t-on l'équivalence des montages série et parallèle représentés ci-dessus ?
2. Pour quelle(s) valeur(s) de la fréquence a-t-on $\frac{L}{R} = \frac{L'}{R'}$?

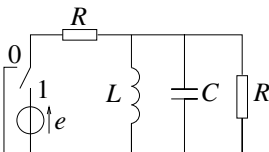
10.2 Courant et tension en phase (★)



On veut que le courant parcourant R_1 soit en phase avec la tension aux bornes du générateur. Quelle est la condition à vérifier pour qu'il en soit ainsi en régime sinusoïdal ?

10.3 Résonance d'un circuit R, L, C parallèle (★)

On considère le circuit suivant :



où e est une tension sinusoïdale de pulsation ω .

1. Donner l'expression complexe de la tension s aux bornes de l'association en parallèle R, L, C .
2. Établir qu'il y a un phénomène de résonance pour la tension s . On précisera la pulsation à laquelle ce phénomène se produit.
3. Que peut-on dire du déphasage à la résonance de la tension s ?
4. Comparer cette résonance avec la résonance en intensité d'un circuit R, L, C série.

CHAPITRE 10 – RÉGIME SINUSOÏDAL

10.4 Recherche d'un circuit résonant (★)

On dispose d'une inductance $L = 40,0 \text{ mH}$, d'une résistance R variable entre $1,00 \text{ k}\Omega$ et $100 \text{ k}\Omega$ d'une capacité variable entre $10,0 \text{ nF}$ et $1,00 \text{ }\mu\text{F}$.

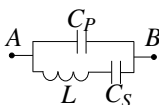
On veut réaliser un circuit résonant à la fréquence $f_0 = 2,00 \text{ kHz}$ de gain 1 à cette fréquence.

1. Proposer des réalisations possibles d'un tel montage.
2. Calculer la facteur de qualité Q de chacun d'eux.
3. Préciser les valeurs de R et C à choisir pour obtenir la bonne fréquence de résonance et la meilleure sélectivité de résonance.
4. Quelle est la « meilleure » solution quant à la sélectivité ?

APPROFONDIR

10.5 Étude de l'impédance d'un quartz piezoélectrique (★)

Le schéma électrique simplifié d'un quartz est donné sur la figure ci-dessous. On néglige sa résistance R . On donne $L = 500 \text{ mH}$; $C_S = 0,08 \text{ pF}$ et $C_P = 8 \text{ pF}$.



1. a. Calculer l'impédance complexe du quartz vue entre les bornes A et B et la mettre

sous la forme : $Z_{AB} = \left(-\frac{j}{\alpha\omega}\right) \frac{1 - \frac{\omega^2}{\omega_r^2}}{1 - \frac{\omega^2}{\omega_a^2}}$, où j est le nombre complexe tel que $j^2 = -1$, et α , ω_r et ω_a sont à déterminer. Montrer aussi que $\omega_a^2 > \omega_r^2$.

- b. Calculer les valeurs numériques des fréquences f_a et f_r correspondant aux pulsations ω_a et ω_r .

2. a. Étudier le comportement inductif (partie imaginaire positive) ou capacitif (partie imaginaire négative) du quartz en fonction de la fréquence. Représenter l'argument de Z_{AB} en fonction de ω .

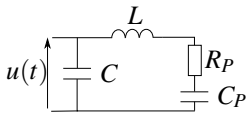
b. Tracer l'allure de $Z_{AB} = |Z_{AB}|$, module de l'impédance complexe du quartz, en fonction de la fréquence.

- c. Comment est modifiée cette courbe si on tient compte de la résistance de la bobine ?

10.6 Adaptation d'impédance, d'après Polytechnique 2004 (★)

Un dipôle électrocinétique linéaire passif est, en régime sinusoïdal permanent, caractérisé par son impédance complexe $\underline{Z} = R + jX$.

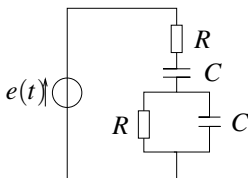
Le système étudié (réacteur à plasma) est modélisé par un circuit série $R_P - C_P$. On veut diminuer au maximum la partie imaginaire (appelée *partie réactive*) de cette impédance $\underline{Z}_P$. Pour cela, on réalise le circuit de la figure suivante.



- 1. Exprimer l'admittance totale $\underline{Y}$ du dipôle de la figure précédente. Déterminer l'expression de C qui annule la partie réactive de $\underline{Z}_P$.
- 2. La condition précédente étant réalisée, déterminer l'expression de l'impédance $\underline{Z}$ totale du dipôle, notée alors R_1 .

10.7 Étude d'un circuit en régime sinusoïdal (★★)

- 1. On considère le dipôle constitué d'une résistance R en parallèle avec une capacité C . Déterminer la résistance R' et la capacité C' qui, en série, ont la même impédance que ce dipôle pour une pulsation ω donnée de la tension appliquée.
- 2. Tracer sur un même graphique les courbes représentatives de $\frac{R}{R'}$ et $\frac{C}{C'}$ en fonction du rapport $\frac{\omega}{\omega_0}$ où $\omega_0 = \frac{1}{RC}$.
- 3. On considère désormais le dispositif constitué de l'association en série des dipôles formés d'une résistance R et d'une capacité C respectivement associées en série et en parallèle. On alimente l'ensemble par une tension sinusoïdale $e(t)$ d'amplitude E , de pulsation ω et de phase initiale nulle.

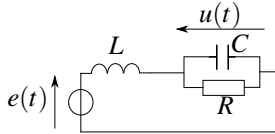


- Déterminer le rapport $\frac{u}{e}$ où $u(t)$ désigne la tension aux bornes de l'association en parallèle de R et C . Même question pour $v(t)$ la tension aux bornes de l'association en série de R et C .
- 4. Exprimer l'amplitude de u .
 - 5. Tracer et justifier l'allure du rapport entre l'amplitude de u et celle de e en fonction de la pulsation ω . On donnera toutes les valeurs particulières.
 - 6. Quelles sont les limites de variation du déphasage ? On pourra tracer un tableau asymptotique simplifiant la fonction de transfert.
 - 7. Déterminer l'intensité complexe circulant dans le générateur.
 - 8. On branche les deux bornes extrêmes d'un potentiomètre en parallèle avec le générateur et on relie la borne intermédiaire *via* un voltmètre à la connexion entre les associations en série et en parallèle de R et C . On souhaite que le voltmètre indique une tension nulle. Pourquoi choisit-on de se placer à la pulsation ω_0 ?
 - 9. Déterminer la position qu'il faut donner au potentiomètre pour que le voltmètre indique une tension nulle.

CHAPITRE 10 – RÉGIME SINUSOÏDAL

10.8 Étude d'une résonance (1) (★★)

On considère le circuit ci-dessous alimenté par une source de tension sinusoïdale de f.é.m. $e(t) = E\sqrt{2}\cos(\omega t)$.



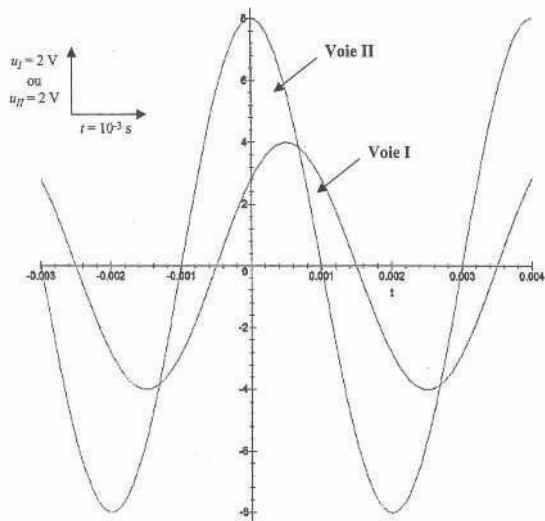
1. En utilisant un diviseur de tension, établir l'expression complexe de $\underline{u}$ en fonction de E , L , R , C et ω .
2. Étude de la réponse en fréquence du circuit : étudier le module de $\underline{u}$ en fonction de ω . Tracer l'allure des courbes.

10.9 Quelques résonances, d'après Centrale TSI 2003 (★)

1. Résonance série :
 - a. Écrire l'impédance $\underline{Z}'$ d'un circuit composé d'une résistance R , d'une inductance L et d'une capacité C placées en série.
 - b. Exprimer le module Z' de l'impédance en fonction de $\omega_0 = \frac{1}{\sqrt{LC}}$ et de $Q = \frac{L\omega_0}{R}$.
 - c. Déterminer le retard de phase φ de l'intensité i parcourant ce circuit lorsqu'on l'alimente par une tension sinusoïdale d'amplitude E et de pulsation ω .
 - d. Tracer en le justifiant le rapport $\frac{Z'}{R}$ en fonction de $\frac{\omega}{\omega_0}$.
 - e. En déduire que le module de l'intensité I passe par un maximum $I_{\max}$ dont on donnera la valeur.
 - f. Tracer le rapport $\frac{I}{I_{\max}}$ en fonction de $\frac{\omega}{\omega_0}$.
 - g. Même question pour φ .
 - h. Que représentent les paramètres ω_0 et Q ?
 - i. Dans quel domaine varie la phase φ pour une pulsation appartenant à la bande de pulsation $\Delta\omega$?
2. Résonance parallèle :
 - a. Écrire l'impédance $\underline{Z}$ du nouveau circuit composé de l'association en parallèle d'un condensateur de capacité C et d'une bobine d'inductance L et de résistance R .
 - b. Exprimer $\underline{Z}$ en fonction de Q , ω , ω_0 , $\underline{Z}'$, C et R .
 - c. Montrer que pour $Q \gg 1$, on peut écrire $\underline{Z} = \frac{R^2 Q^2}{\underline{Z}'}$.
 - d. Quelle est la valeur de $\underline{Z}$ lorsque la pulsation tend vers ω_0 ?
 - e. Que peut-on en conclure ?
 - f. Pour $\omega = \omega_0$, donner l'expression du courant $\underline{i}_1$ traversant le condensateur.
 - g. Même question pour le courant $\underline{i}_2$ traversant la bobine.

10.10 Analyse expérimentale d'un circuit R, L, C série, d'après CCP DEUG 2005 (*)

Soit un circuit R, L, C série alimenté par une tension sinusoïdale dont on regarde la tension e aux bornes du générateur en voie 1 et s aux bornes de R en voie 2.



1. Sélectivité d'un filtre passe-bande :
 - a. Établir l'expression de l'impédance complexe $\underline{Z}$ du circuit R, L, C série.
 - b. Exprimer la fonction de transfert du montage en fonction de R, L, C et ω .
 - c. On pose $\omega_0 = \frac{1}{\sqrt{LC}}$, $x = \frac{\omega}{\omega_0}$ et $Q = \frac{L\omega_0}{R}$. Donner l'expression de la fonction de transfert en fonction de ces nouveaux paramètres.
 - d. Montrer que, pour toute valeur de Q , le gain admet un même maximum $G_{\max}$. On précisera la pulsation correspondante.
 - e. Déterminer la pulsation pour laquelle le déphasage est nul.
2. Détermination des paramètres de la bobine :

Pour $R = 20 \, \Omega$ et $C = 10 \, \mu\text{F}$, on a les oscillogrammes donnés en annexe.

 - a. Déterminer graphiquement les valeurs de la période T , de la pulsation ω , de l'amplitude de l'intensité I_M , de l'amplitude de la tension délivrée par le générateur U_M et de la norme de l'impédance Z du circuit.
 - b. Quelle voie est en retard sur l'autre ?
 - c. Déterminer le déphasage φ entre les deux tensions.
 - d. Montrer que les valeurs expérimentales de Z , φ et R sont incohérentes dans l'hypothèse d'une bobine idéale modélisée par une seule inductance L .
 - e. On prend comme modèle de la bobine une inductance L en série avec une résistance r . Déterminer la valeur de r à partir des observations expérimentales.

CHAPITRE 10 – RÉGIME SINUSOÏDAL

f. En déduire la valeur de L .

10.11 Étude du modèle simplifié d'un moteur, d'après Banque PT 2004 (★)

Un moteur est alimenté par une tension sinusoïdale d'amplitude complexe $\underline{V}$; on note $\underline{I}$ l'amplitude complexe de l'intensité du courant passant dans le circuit.

Le schéma électrique équivalent est représenté sur la figure (10.27). La source de tension alimentant le moteur est supposée parfaite.

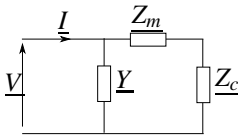


Figure 10.27

$\underline{Y}$ représente l'admittance électrique complexe intrinsèque du moteur. On admettra que le comportement mécanique du moteur se traduit du point de vue électrique par une impédance complexe $\underline{Z}_m$ (dite impédance motionnelle) et par $\underline{Z}_c$, dont la valeur est fonction de la charge mécanique du moteur.

Le dipôle d'admittance $\underline{Y}$ est constitué d'une résistance R_0 en parallèle avec un condensateur C_0 . L'impédance $\underline{Z}_m$ correspond à un circuit RLC série. L'impédance $\underline{Z}_c$ correspond à une résistance R_c .

1. Représenter la figure (10.27) en remplaçant les éléments $\underline{Y}$, $\underline{Z}_m$ et $\underline{Z}_c$ par les inductances, résistances et condensateurs qui leur correspondent.

2. Déterminer la pulsation de résonance en courant ω_s du circuit série constitué de $\underline{Z}_m$ et $\underline{Z}_c$. Pour la suite, on prendra les valeurs numériques $R_0 = 18\text{k}\Omega$, $C_0 = 8\text{ nF}$, $R = 50\ \Omega$, $L = 0,1\text{ H}$, $C = 0,2\text{ nF}$ et $R_c = 50\ \Omega$.

3. On note $\underline{Y}_p$ l'admittance complexe du circuit constitué de l'ensemble des éléments $\underline{Z}_m$, $\underline{Z}_c$ et $\underline{Y}$. Les figures (10.28) et (10.29) représentent l'évolution du module Y_p de l'admittance complexe $\underline{Y}_p$ en fonction de la pulsation réduite $x = \frac{\omega}{\omega_s}$. La figure (10.29) est un agrandissement de la figure (10.28) autour de $x = 1$.

À partir des figures (10.28) et (10.29), déterminer la valeur numérique de la fréquence(en Hertz) qui permet d'obtenir une résonance en courant.

4. Comparer $Y = |\underline{Y}|$ et le module de l'admittance complexe $\underline{Y}' = \frac{1}{\underline{Z}_m + \underline{Z}_c}$ lorsque $\omega = \omega_s$ et interpréter le résultat de la question précédente.

5. Une modification du poids de la charge transportée par le moteur provoque une variation de la résistance R_c , qui reste toutefois de l'ordre de grandeur de la dizaine d'ohms ; cette modification a-t-elle un effet significatif sur la valeur de la résonance en courant ? Justifier la réponse.

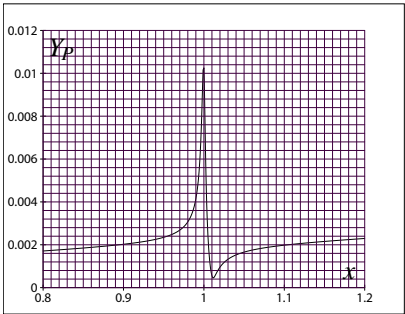


Figure 10.28

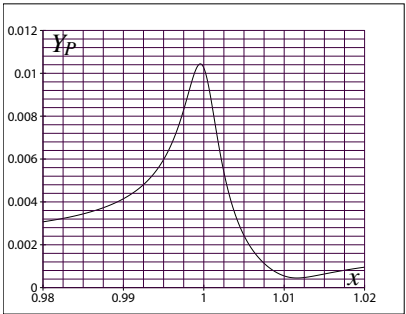


Figure 10.29

CORRIGÉS

10.1 Équivalence entre deux dipôles

1. En série, on a $\underline{Z}' = R' + jL'\omega$ et en parallèle, $\underline{Z} = \frac{RjL\omega}{R + jL\omega} = \frac{RL^2\omega^2}{R^2 + L^2\omega^2} + j\frac{R^2L\omega}{R^2 + L^2\omega^2}$. En identifiant les parties réelles et imaginaires, on a l'équivalence des deux dipôles pour : $R' = \frac{RL^2\omega^2}{R^2 + L^2\omega^2}$ et $L' = \frac{R^2L}{R^2 + L^2\omega^2}$. On peut noter que l'équivalence n'est valable que pour une fréquence précise. On verra dans le chapitre sur l'instrumentation électrique qu'il s'agit du modèle bande étroite d'une bobine.

2. Le rapport $\frac{L'}{R'}$ vaut, d'après la question précédente : $\frac{L'}{R'} = \frac{R^2L(R^2 + L^2\omega^2)}{(R^2 + L^2\omega^2)RL^2\omega^2} = \frac{R}{L\omega^2}$. On aura $\frac{L'}{R'} = \frac{L}{R}$ si $\frac{R}{L\omega^2} = \frac{L}{R}$, soit pour $\omega = \frac{R}{L}$.

10.2 Courant et tension en phase

On calcule l'impédance équivalente à l'ensemble L , R_2 et C soit $\underline{Z} = \frac{\frac{1}{jC\omega}(R_2 + jL\omega)}{\frac{1}{jC\omega} + R_2 + jL\omega} = \frac{R_2 + jL\omega}{1 + jR_2C\omega - LC\omega^2}$. On en déduit l'intensité $\underline{I}$ par une loi des mailles $\underline{I} = \frac{\underline{E}}{R_1 + \underline{Z}}$. Intensité et tension sont en phase si l'argument de $\underline{Z}$ est nul soit $R_2^2C = L(1 - LC\omega^2)$.

10.3 Résonance d'un circuit R, L, C parallèle

1. L'association en parallèle de R , L et C admet pour impédance $\underline{Z} = \frac{1}{\frac{1}{R} + \frac{1}{jL\omega} + jC\omega} = \frac{jRL\omega}{R + jL\omega - RCL\omega^2}$. On obtient alors l'expression de la tension s en utilisant la relation du pont diviseur de tension soit $\underline{s} = \frac{\underline{Z}}{\underline{Z} + r}\underline{e} = \frac{jRL\omega}{rR + jL\omega(r + R) - rRCL\omega^2}\underline{e}$.

2. On a résonance si l'amplitude passe par un maximum. Or l'amplitude s'obtient en déterminant le module soit $|\underline{s}| = \frac{|\underline{e}|}{\sqrt{\left(1 + \frac{r}{R}\right)^2 + r^2\left(C\omega - \frac{1}{L\omega}\right)^2}}$. Le dénominateur est toujours plus grand que $1 + \frac{r}{R}$ et prend cette valeur pour $\omega_0 = \frac{1}{\sqrt{LC}}$. On a donc résonance pour ω_0 .

3. Pour $\omega = \omega_0$, on a $\underline{s} = \frac{R\underline{e}}{r + R}$ et il n'y a pas de déphasage.

4. On obtient la même pulsation de résonance que pour le circuit R, L, C série mais un facteur de qualité variable en fonction de la valeur de r .

10.4 Recherche d'un circuit résonant

1. On aura un montage du type pont diviseur de tension avec deux impédances $\underline{Z}_1$ et $\underline{Z}_2$. Par conséquent, on a $\underline{s} = \frac{\underline{Z}_2}{\underline{Z}_1 + \underline{Z}_2} \underline{e}$. On veut un gain de 1 à la résonance, ce qui offre deux possibilités : $\underline{Z}_1(\omega_0) = 0$ ou $\underline{Z}_2(\omega_0)$ infinie.

Pour le premier cas, on associe en série L et C soit $\underline{Z}_1 = j \left(L\omega - \frac{1}{C\omega} \right)$ qui s'annule pour $\omega_0 = \frac{1}{\sqrt{LC}}$ soit $C = \frac{1}{L\omega_0^2} = 158 \text{ nF}$. Dans ce cas, on a $\underline{Z}_2 = R$.

Pour le deuxième cas, $\underline{Z}_2$ s'obtient par une association en parallèle de L et C soit $\underline{Z}_2 = \frac{1}{j \left(C\omega - \frac{1}{L\omega} \right)}$ qui tend vers l'infini pour $\omega = \omega_0$. Dans ce cas, $\underline{Z}_1 = R$.

2. $Q = \frac{1}{2\xi}$ où ξ est le coefficient d'amortissement. On le calcule en écrivant la fonction de transfert, en l'écrivant sous forme canonique, en trouvant ξ , puis Q .

Pour le premier cas, on trouve $Q = \frac{1}{R} \sqrt{\frac{L}{C}}$. Pour le second $Q = R \sqrt{\frac{C}{L}}$.

La sélectivité est meilleure pour le second montage car on peut augmenter la valeur de R tandis que diminuer sa valeur est difficile.

3. On a déjà déterminé la valeur $C = 158 \text{ nF}$.

Pour le premier cas, on prend la valeur la plus faible pour R soit $1 \text{ k}\Omega$ pour ne pas être gêné par les résistances internes des sources de tension. On en déduit $Q = 0,50$.

Pour le second cas, on choisit la valeur la plus grande possible pour R soit $100 \text{ k}\Omega$ pour ne pas être gêné par les impédances d'entrée des appareils de mesure. Dans ce cas, $Q = 199$.

4. Le montage le plus sélectif est celui qui présente le pic le plus haut et le plus étroit, c'est à dire de coefficient d'amortissement ξ le plus faible, ou de facteur de qualité $Q = \frac{1}{2\xi}$ le plus fort. Le deuxième montage est donc meilleur d'après ce qui précède.

10.5 Étude de l'impédance d'un quartz piezoélectrique, d'après ENSTIM 2004

1. a. On a donc une impédance (C_p) en parallèle sur deux impédances en série (L et C_s).

On calcule l'admittance de l'ensemble : $\underline{Y} = jC_p\omega + \frac{1}{jL\omega + \frac{1}{jC_s\omega}} = jC_p\omega + \frac{jC_s\omega}{1 - LC_s\omega^2}$. En

réduisant au même dénominateur puis en factorisant, on obtient : $\underline{Y} = j(C_p + C_s)\omega \frac{1 - \frac{LC_pC_s\omega^2}{C_p + C_s}}{1 - LC_s\omega^2}$.

En inversant, on peut écrire $\underline{Z}$ sous la forme de l'énoncé : $\underline{Z}_{AB} = \frac{1}{\underline{Y}} = \left(-\frac{j}{\alpha\omega} \right) \frac{1 - \frac{\omega_r^2}{\omega_a^2}}{1 - \frac{\omega^2}{\omega_a^2}}$, avec

$\alpha = C_p + C_s$, $\omega_r = 1/\sqrt{LC_s}$ et $\omega_a = 1/\sqrt{LC_e}$ où l'on a posé $C_e = C_pC_s/(C_p + C_s)$.

Comme $C_e < C_s$, alors $\omega_a > \omega_r$.

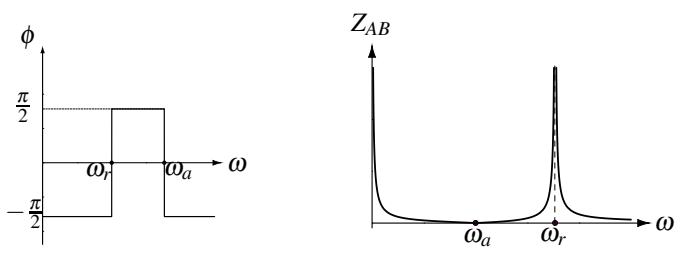
CHAPITRE 10 – RÉGIME SINUSOÏDAL

b. La relation entre pulsation et fréquence étant $\omega = 2\pi f$, l'application numérique donne : $f_r = 796$ Hz et $f_a = 800$ kHz.

2. a. L'impédance $\underline{Z}_{AB}$ est un imaginaire pur, donc la phase ϕ est $\pm\pi/2$. Si $\phi = \pi/2$ le comportement est inductif (comme une bobine), sinon il est capacitif (comme un condensa-

teur). Pour cette étude, il suffit d'étudier le signe du rapport $-\frac{1 - \frac{\omega^2}{\omega_r^2}}{1 - \frac{\omega^2}{\omega_a^2}}$. On a représenté ϕ en

fonction de ω :



Le comportement est donc capacitif pour $\omega \in [0, \omega_r]$ ou $\omega \in [\omega_a, \infty]$ et inductif pour $\omega \in [\omega_r, \omega_a]$.

b. On a tracé au-dessus l'allure du module Z_{AB} en fonction de ω .

c. Si l'on tient compte de la résistance de la bobine, les valeurs de ω_a et ω_r seront légèrement différentes. Il n'y aura pas de pulsation qui annule totalement $|Z|$ ou qui le rende infini.

10.6 Adaptation d'impédance, d'après Polytechnique 2004

1. L'impédance de L , R_p et C_p en série est : $\underline{Z}_1 = jL\omega + R_p + \frac{1}{jC\omega}$. Cette impédance est en parallèle avec le condensateur C soit : $\underline{Y} = jC\omega + \frac{1}{\underline{Z}_1} = jC\omega + \frac{1}{jL\omega + R_p + \frac{1}{jC\omega}}$

$\frac{R_p - j\left(L\omega - \frac{1}{C_p\omega}\right)}{R_p^2 + \left(L\omega - \frac{1}{C_p\omega}\right)^2}$. On veut que la partie imaginaire de $\underline{Y}$ soit nulle, ce qui entraîne :

$$C\omega = \frac{\left(L\omega - \frac{1}{C_p\omega}\right)}{R_p^2 + \left(L\omega - \frac{1}{C_p\omega}\right)^2} \text{ d'où } C = \frac{\left(L - \frac{1}{C_p\omega^2}\right)}{R_p^2 + \left(L\omega - \frac{1}{C_p\omega}\right)^2}.$$

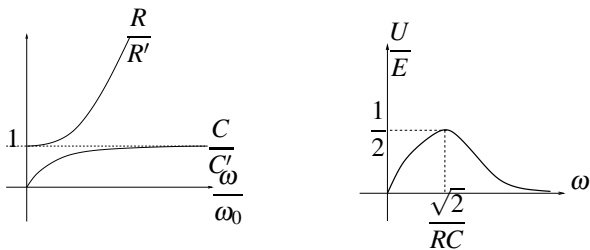
2. On obtient alors : $R_1 = \frac{1}{\underline{Y}} = \frac{R_p^2 + \left(L\omega - \frac{1}{C_p\omega}\right)^2}{R_p} = R_p + \frac{1}{R_p} \left(L\omega - \frac{1}{C_p\omega}\right)^2$.

10.7 Étude d'un circuit en régime sinusoïdal

1. Les impédances s'écrivent $\underline{Z} = \frac{1}{jC\omega + \frac{1}{R}} = \frac{R}{1 + jRC\omega}$ et $\underline{Z}' = R' + \frac{1}{jC\omega}$. On obtient en

identifiant $R' = \frac{R}{1 + (RC\omega)^2}$ et $C' = C \frac{1 + (RC\omega)^2}{(RC\omega)^2}$.

2. On en déduit $\frac{R}{R'} = 1 + \left(\frac{\omega}{\omega_0}\right)^2$ et $\frac{C}{C'} = \frac{\left(\frac{\omega}{\omega_0}\right)^2}{1 + \left(\frac{\omega}{\omega_0}\right)^2}$. Les allures sont les suivantes :



3. On utilise la relation du pont diviseur de tension et on obtient $\frac{u}{e} = \frac{jRC\omega}{2 + 2jRC\omega - R^2C^2\omega^2}$

ainsi que $\frac{v}{e} = \frac{(1 + jRC\omega)^2}{2 + 2jRC\omega - R^2C^2\omega^2}$.

4. Pour l'amplitude, on prend le module soit $U = \frac{RC\omega}{\sqrt{(2 - R^2C^2\omega^2)^2 + 4R^2C^2\omega^2}} E$.

5. On pose $x = RC\omega$ et on étudie la fonction $f(x) = \frac{x}{\sqrt{1 + x^4}}$ dont la dérivée vaut $f'(x) = \frac{4 - x^4}{(4 + x^4)^{3/2}}$. La fonction f est croissante de 0 à $\frac{1}{2}$ pour $\omega < \frac{\sqrt{2}}{RC}$ et décroissante ensuite jusqu'à 0.

6. Avec un tableau asymptotique qui simplifie l'expression de la transmittance, suivant $\omega \ll \omega_0 = \frac{\sqrt{2}}{RC}$ ou $\omega \gg \omega_0$:

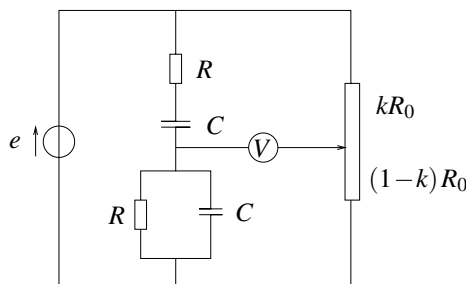
	ω_0	
$2 + 2jRC\omega - R^2C^2\omega^2$	1	$-R^2C^2\omega^2$
$H_{asympt}(j\omega)$	$\frac{jRC\omega}{2}$	$\frac{1}{jRC\omega}$
pente	+20 dB/dec	-20 dB/dec
phase	$+\pi/2$	$-\pi/2$

La phase varie bien de $+\frac{\pi}{2}$ à $-\frac{\pi}{2}$.

CHAPITRE 10 – RÉGIME SINUSOÏDAL

7. L'intensité complexe est $\underline{i} = \frac{\underline{e}}{\underline{Z}_{tot}} = \frac{jC\omega(1 + jRC\omega)}{1 - R^2C^2\omega^2 + 3jRC\omega}\underline{e}$.

8. Il n'y a pas de déphasage entre u et e pour $\omega = \omega_0$, ce qui sera aussi le cas aux bornes de $(1 - k)R_0$. C'est la raison pour laquelle on choisit cette pulsation.



9. On veut également que les modules des tensions u et celle aux bornes de $(1 - k)R_0$ soient égales soit $(1 - k)E = \frac{E}{\sqrt{5}}$. On en déduit $k = \frac{\sqrt{5} - 1}{\sqrt{5}}$.

10.8 Étude d'une résonance (1)

1. On note $\underline{Z}_{eq}$ l'impédance équivalente à R et C en parallèle et $\underline{Y}_{eq} = \frac{1}{R} + jC\omega$ l'admittance.

Le diviseur de tension entre $\underline{u}$ et $\underline{e}$ donne : $\underline{u} = \frac{\underline{Z}_{eq}}{jL\omega + \underline{Z}_{eq}}\underline{e} = \frac{1}{jL\omega \underline{Y}_{eq} + 1}\underline{e} = \frac{E\sqrt{2}e^{j\omega t}}{1 - LC\omega^2 + \frac{jL\omega}{R}}$.

2. On pose $\underline{u} = U\sqrt{2}e^{j(\omega t + \phi)}$. Avec l'expression précédente, on obtient :

$$U = \frac{E}{\sqrt{(1 - LC\omega^2)^2 + \left(\frac{L\omega}{R}\right)^2}} = \frac{E}{\sqrt{f(\omega)}}. \text{ Étant donné que le numérateur est constant, } U \text{ a}$$

les variations inverse de $f(\omega)$: on étudie donc cette dernière fonction.

$$\frac{df}{d\omega} = \frac{d}{d\omega} \left(1 + (LC)^2\omega^4 - 2LC\omega^2 + \frac{L^2\omega^2}{R^2} \right) = 4(LC)^2\omega^3 - 4LC\omega + 2\frac{L^2}{R^2}\omega$$

La dérivée de f est nulle pour $\omega = 0$ et pour $\omega_r = \frac{1}{\sqrt{LC}}\sqrt{1 - \frac{L}{2R^2C}}$ si $R > \sqrt{\frac{L}{2C}}$.

Dans le cas $R > \sqrt{\frac{L}{2C}}$, on peut vérifier, avec le signe de la dérivée, que $\omega = 0$ correspond à un maximum pour $f(\omega)$ donc à un minimum pour U tandis que ω_r correspond à un minimum pour $f(\omega)$ donc à un maximum pour U .

Dans le cas $R < \sqrt{\frac{L}{2C}}$, on peut vérifier, avec le signe de la dérivée, que $\omega = 0$ correspond à un minimum pour $f(\omega)$ donc à un maximum pour U .

Dans les deux cas, $U(\omega = 0) = E$ et $U(\omega \rightarrow \infty) = 0$.

On a tracé $U(\omega)$ sur la figure (10.30) pour deux cas : cas (1) pour $R > \sqrt{\frac{L}{2C}}$ et cas (2) pour $R < \sqrt{\frac{L}{2C}}$.

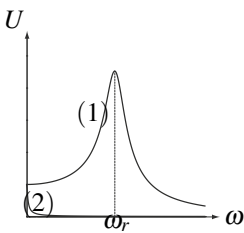
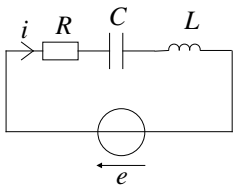


Figure 10.30

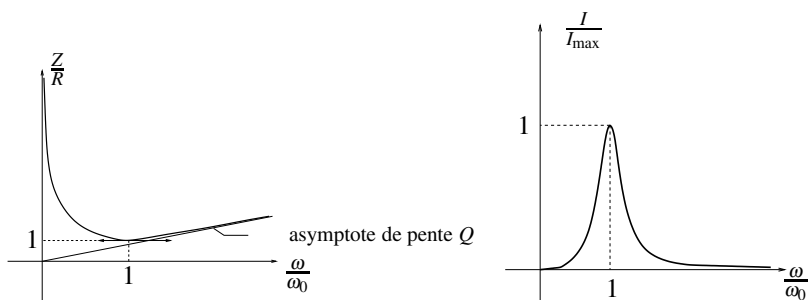
10.9 Quelques résonances, d'après Centrale TSI 2003

1. Le circuit considéré est dessiné page suivante.



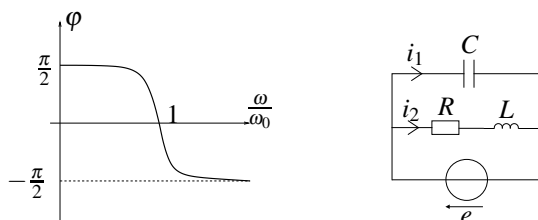
- a. L'impédance du circuit s'écrit $\underline{Z}' = R + j \left(L\omega - \frac{1}{C\omega} \right)$.
- b. Son module vaut $Z' = \sqrt{R^2 + \left(L\omega - \frac{1}{C\omega} \right)^2} = R \sqrt{1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2}$.
- c. La loi des mailles s'écrit $\underline{e} = \underline{Z}'i$. Le déphasage cherché est l'opposé de la phase de $\underline{Z}'$ soit $\varphi = -\text{Arctan}Q \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right) \in \left[-\frac{\pi}{2}, \frac{\pi}{2} \right]$.
- d. On remarque que pour toute valeur de la pulsation ω , on a $\frac{Z}{R} \geq 1$. Le cas d'égalité $\frac{Z}{R} = 1$ est obtenu pour $\omega = \omega_0$ qui correspond donc à un minimum. On en déduit que $\frac{Z}{R}$ décroît de $+\infty$ à 1 pour $\omega = \omega_0$ et croît ensuite vers $+\infty$. Le comportement à l'infini est la même que celle de $Q \frac{\omega}{\omega_0}$: on a donc une asymptote oblique.

CHAPITRE 10 – RÉGIME SINUSOÏDAL



e. L'intensité s'écrit $I = \frac{E}{Z}$ donc les variations sont inverses de celles de Z : I passe par un maximum pour $\omega = \omega_0$ et ce maximum vaut $\frac{E}{R}$ (cf dessin *supra*).

f. La fonction $f(\omega) = -L\omega + \frac{1}{C\omega}$ admet pour dérivée $f'(\omega) = -L - \frac{1}{C\omega^2} < 0$ donc la fonction f est décroissante comme $\varphi = \tan f(\omega)$. Les limites sont $\frac{\pi}{2}$ en 0 et $-\frac{\pi}{2}$ en $+\infty$. On en déduit l'allure du graphe page suivante.



g. La quantité Q est le facteur de qualité et ω_0 la pulsation propre ou pulsation de résonance.

h. Du fait que la fonction φ est décroissante, il suffit calculer sa valeur pour ω_1 et ω_2 pour obtenir l'intervalle demandé soit $\left[-\frac{\pi}{4}, +\frac{\pi}{4}\right]$.

2. Le nouveau circuit est dessiné juste au dessus.

a. L'impédance s'écrit $\underline{Z} = \frac{R + jL\omega}{1 + jRC\omega - LC\omega^2}$.

b. On peut mettre cette impédance sous la forme $\underline{Z} = \frac{R \left(1 + jQ\frac{\omega}{\omega_0}\right)}{jC\omega\underline{Z}'}$.

c. Dans le cas où $Q \gg 1$, on obtient $\underline{Z} \approx \frac{R^2 Q^2}{\underline{Z}'}$ en négligeant 1 devant $jQ\frac{\omega}{\omega_0}$ et en utilisant $L\omega_0 = \frac{1}{C\omega_0}$.

d. $\underline{Z}'$ tend vers R pour ω tendant vers ω_0 donc $\underline{Z}$ tend vers RQ^2 .

e. L'impédance est équivalente à une résistance comme dans le cas de l'association en série mais la valeur a changé.

f. Le courant demandé s'écrit $\underline{i}_1 = jC\omega \underline{e}$ soit après calculs $i_1 = -\frac{E}{RQ} \sin \omega t$.

g. Par la loi des nœuds, on en déduit $\underline{i}_2 = \underline{i} - \underline{i}_1$ et $i_2 = \frac{E}{RQ} \sqrt{1 + \frac{1}{Q^2}} \cos(\omega_0 t - \text{Arctan} Q)$.

10.10 Analyse expérimentale d'un circuit R, L, C série, d'après CCP DEUG 2005

1. a. L'impédance de l'association en série de R, L et C s'écrit $\underline{Z} = R + j \left(L\omega - \frac{1}{C\omega} \right)$.
b. En appliquant la relation du pont diviseur de tension, on obtient

$$\underline{H} = \frac{\underline{R}}{\underline{Z}} = \frac{R}{R + j \left(L\omega - \frac{1}{C\omega} \right)}$$

c. En introduisant les quantités proposées, on obtient $\underline{H} = \frac{1}{1 + jQ \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)}$.

d. Le gain s'obtient en prenant le module de la fonction de transfert soit :

$$G = \frac{1}{\sqrt{1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2}}$$

Or $1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2 \geq 1$ et on a le cas d'égalité pour $\omega = \omega_0$. Dans ce cas, le gain G est maximal et sa valeur est $G_{\max} = 1$.

e. Le déphasage est nul lorsque la partie imaginaire est nulle soit $\omega = \omega_0$.

2. a. La période lue sur le graphique fourni est $T = 4$ ms, on en déduit la pulsation $\omega = \frac{2\pi}{T} = 1,57 \cdot 10^3 \text{ rad.s}^{-1}$. L'amplitude de l'intensité est $I_M = 0,2$ A, celle de la tension $U_M = 8$ V et on en déduit $Z = \frac{U_M}{I_M} = 40 \Omega$.

b. La voie I est en retard sur la voie II.

c. Le déphasage lu sur la courbe est $\varphi = 2\pi \frac{2,5}{20} = \frac{\pi}{4}$.

d. À partir de la valeur de Z , on en déduit $L = \frac{1}{\omega} \left(\frac{1}{C\omega} + \sqrt{Z^2 - R^2} \right) = 62,6$ mH d'après

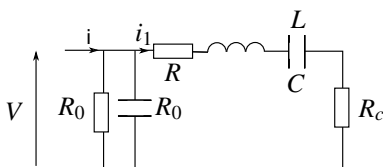
la relation établie dans la première partie. On a alors $\varphi = \text{Arctan} \frac{L\omega - \frac{1}{C\omega}}{R}$ soit numériquement $\varphi = 1,73 \neq 0,79$, ce qui est incohérent.

CHAPITRE 10 – RÉGIME SINUSOÏDAL

- e. On remplace R par $R + r$ et on obtient $Z = \sqrt{(R+r)^2 + \left(L\omega - \frac{1}{C\omega}\right)^2}$ et $\tan \varphi = \frac{L\omega - \frac{1}{C\omega}}{R+r}$. On en déduit l'expression $r = \frac{Z}{\sqrt{1 + \tan^2 \varphi}} - R$ soit numériquement $r = 8,3 \, \Omega$.
- f. On obtient aussi $L \frac{1}{\omega} \left(\frac{1}{C\omega} + (R+r) \tan \varphi \right) = 58,6 \, \text{mH}$.

10.11 Étude du modèle simplifié d'un moteur, d'après Banque PT 2004

1. D'après les indications de l'énoncé, on obtient le circuit de la figure suivante



2. On s'intéresse donc à un circuit RLC série constitué d'une résistance $R + R_c$, d'une bobine L et d'un condensateur C . La loi d'Ohm donne pour les valeurs efficaces $I_1 = V/Z$ où Z est le module de $\underline{Z} = R + R_c + j \left(L\omega - \frac{1}{C\omega} \right)$. Puisque V est constante, l'intensité I est maximale

si $Z = \sqrt{(R + R_c)^2 + \left(L\omega - \frac{1}{C\omega} \right)^2}$ est minimale donc si $\left(L\omega - \frac{1}{C\omega} \right) = 0$ ou $\omega = 1/\sqrt{LC}$.

3. La résonance en courant correspond à une impédance totale minimale donc une admittance totale maximale. La figure donne Y_p maximale pour x légèrement inférieur à 1 soit ω légèrement inférieure à ω_s .

L'application numérique donne $\omega = 36\,000 \, \text{rad.s}^{-1}$ soit $f = \omega/2\pi = 5,6 \, \text{kHz}$.

4. L'admittance $\underline{Y}$ est formée de R_0 en parallèle avec C_0 soit $Y = \sqrt{\frac{1}{R_0^2} + (C_0\omega)^2}$. Pour

$\omega = \omega_s$, la valeur de l'admittance Y est $2,9 \cdot 10^{-4} \, \text{S}$. On rappelle que la siemens S est l'inverse de l'ohm : $1 \, \text{S} = 1 \, \Omega^{-1}$. D'après les questions précédentes, lorsque $\omega = \omega_s$, l'impédance $\underline{Z}_m + \underline{Z}_c$ est égale à $R + R_c$, donc $Y' = 1/(R + R_c)$ et $Y' = 0,01 \, \text{S}$. On remarque donc que $Y \ll Y'$, or l'admittance totale est $\underline{Y}_{\text{tot}} = \underline{Y} + \underline{Y}'$ donc $\underline{Y}_{\text{tot}} \simeq \underline{Y}'$. La fréquence de résonance est quasiment celle de l'admittance $\underline{Y}'$.

5. Si $R_c \simeq 10 \, \Omega$, alors $Y' \simeq 1,6 \cdot 10^{-2} \, \text{S}$. L'approximation précédente est encore mieux vérifiée, donc la fréquence de résonance est quasiment inchangée.

Analyse fréquentielle d'un système linéaire

11

Dans ce chapitre, on s'intéresse à la fonction de transfert d'un montage électrique. On verra comment elle se représente en pratique et les différents types de fonction de transfert qui existent.

1 Représentation graphique de la fonction de transfert

La fonction de transfert d'un système, $\underline{H}(j\omega)$, indique comment un système modifie l'amplitude et la phase d'un signal en fonction de sa pulsation ω , en régime permanent sinusoïdal.

Pour saisir d'un coup l'influence d'un système sur un signal, on utilise la représentation graphique de la fonction de transfert, nommée **diagramme de Bode**¹, qui contient deux graphiques : l'un représente la variation de $|\underline{H}(j\omega)|$ en fonction de ω , l'autre celle de $\arg(\underline{H}(j\omega))$.

1.1 Le gain en décibel

Le rapport des amplitudes des signaux de sortie et d'entrée, exprimé de façon logarithmique, est en **décibel**². Il est nommé **gain en décibel** G_{dB} du système à la pulsation ω :

$$G_{dB}(\omega) = 20 \log |\underline{H}(j\omega)| = 20 \log \left| \frac{\underline{s}(t)}{\underline{e}(t)} \right|.$$



Le log opère sur un réel positif, ainsi est-ce bien le module de $\underline{H}$ qui intervient dans la formule.

1. inventés par Hendrik Wade Bode, 1905 – 1982, ingénieur américain dont les travaux présentèrent des avancées majeures en automatique et en télécommunications.
2. nommé en l'honneur d'Alexandre Graham Bell, 1847 – 1922, professeur de phonologie à l'université de Boston, qui déposa le brevet du téléphone en 1876.

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D’UN SYSTÈME LINÉAIRE

Le lecteur latiniste s’étonnera du préfixe déci utilisé dans le décibel, alors que le log est multiplié par 20. Le **décibel** fut à l’origine introduit pour exprimer des rapports de puissance. En notant $\mathcal{P}_E$ la puissance du signal d’entrée et $\mathcal{P}_S$ celle du signal de sortie, le rapport de puissance s’écrit en décibel avec le facteur 10 caractéristique du préfixe déci :

$$\left(\frac{\mathcal{P}_S}{\mathcal{P}_E} \right)_{dB} = 10 \log \frac{\mathcal{P}_S}{\mathcal{P}_E}.$$

Or la puissance d’un signal est proportionnelle à la valeur efficace au carré de ce signal. La constante de proportionnalité s’élimine lorsqu’on passe au rapport :

$$\left(\frac{\mathcal{P}_S}{\mathcal{P}_E} \right)_{dB} = 10 \log \frac{S_{eff}^2}{E_{eff}^2} = 20 \log \frac{S_{eff}}{E_{eff}}.$$

Pour un rapport de tensions, le log est donc multiplié par 20 dans une expression en décibel.

1.2 La phase

Il s’agit simplement de :

$$\varphi(\omega) = \arg(\underline{H}(j\omega)).$$

1.3 Relevés expérimentaux

Avec la notation complexe des signaux d’entrée et de sortie du système, $\underline{e}(t) = E_0 \exp(j\omega t)$ et $\underline{s}(t) = S_0 \exp(j\varphi) \exp(j\omega t)$, la fonction de transfert s’écrit :

$$\underline{H}(j\omega) = \frac{S_0}{E_0} \exp(j\varphi).$$

Le gain du système est donc simplement lié au rapport des amplitudes des signaux :

$$G_{dB}(\omega) = 20 \log \frac{S_0}{E_0} = 20 \log \frac{S_{eff}}{E_{eff}} = 20 \log \frac{S_{CC}}{E_{CC}},$$

où S_{CC} représente l’amplitude crête à crête du signal, double de l’amplitude.

Quant à la phase, il s’agit du déphasage de la sortie par rapport à l’entrée.

Exemple

Sur l’exemple du circuit RC présenté au paragraphe 2.2 du précédent chapitre, page 339, on relève les valeurs :

$$G_{dB}(\omega) = 20 \log \frac{32,9 \cdot 10^{-3}}{1,0} = -29,7 \text{ dB} \quad \text{et} \quad \varphi = -88^\circ.$$

Le signe de φ s’explique par le retard de la sortie sur l’entrée.

1.4 Exemples de diagramme de Bode

On trace, pour le système RC du premier ordre dont on a calculé la fonction de transfert à la page 355, $G_{dB}(\omega)$ et $\varphi(\omega)$ en fonction du produit $RC\omega$, en échelle semilogarithmique, car seul l'axe des abscisses est gradué logarithmiquement : $\underline{H}(j\omega) = \frac{1}{1 + jRC\omega}$ donc :

$$\varphi(\omega) = \arg\left(\frac{1}{1 + jRC\omega}\right),$$

et :

$$G_{dB}(\omega) = 20\log\left(\frac{1}{\sqrt{1 + (RC\omega)^2}}\right) = -20\log\left(\sqrt{1 + (RC\omega)^2}\right) = -10\log\left(1 + (RC\omega)^2\right).$$

La variable $RC\omega$ est choisie pour l'axe des abscisses car c'est ce produit qui intervient dans la fonction de transfert. Attendu que le produit RC s'exprime en s et ω en rad.s^{-1} , $RC\omega$ est en rad, c'est à dire qu'il est sans dimension. Les courbes de gain et de phase sont les suivantes :

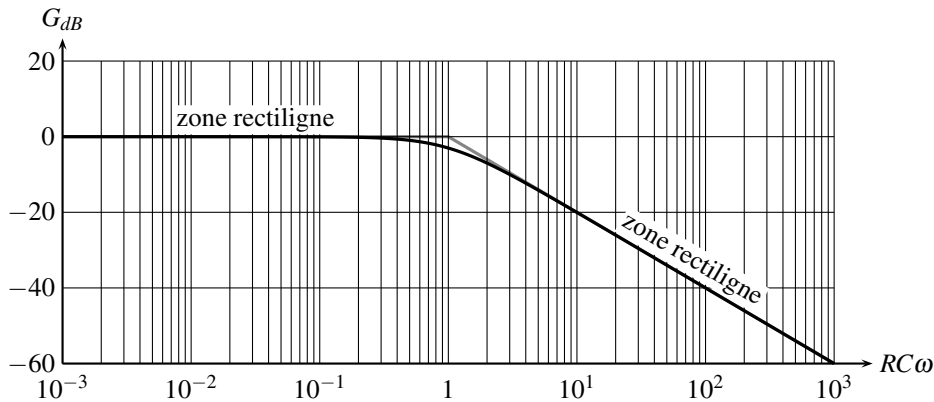


Figure 11.1 – Gain du système (les asymptotes sont en gris).

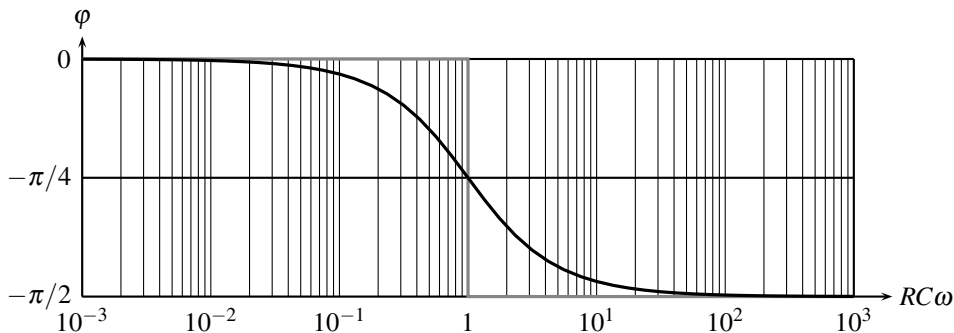


Figure 11.2 – Phase du système (les asymptotes sont en gris).

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D’UN SYSTÈME LINÉAIRE

On observe des zones rectilignes dans le diagramme dans Bode, qui sont interprétables simplement avec l’expression de la fonction de transfert. Elles correspondent pour le gain $G(\omega) = |\underline{H}(j\omega)|$, en échelles logarithmiques, à des lois de puissance en $G(\omega) = \alpha\omega^n$ (voir annexe mathématique page 1081).

Afin d’interpréter simplement ces zones rectilignes, on se propose :

- de simplifier la fonction de transfert ;
- d’interpréter le graphe du gain en décibel en fonction de ω ;
- d’interpréter de même le graphe de la phase en fonction de ω .

Simplification de la fonction de transfert Elle s’écrit $\underline{H}(j\omega) = \frac{1}{1 + jRC\omega}$. Le dénominateur est la somme de deux termes, 1 et $jRC\omega$. À quelle condition l’un est-il prédominant devant l’autre en module ? 1 est beaucoup plus petit que $RC\omega = |jRC\omega|$ dès que :

$$1 \ll RC\omega \quad \text{soit} \quad \omega \gg \frac{1}{RC}, \quad \text{en pratique} \quad \omega > \frac{10}{RC}.$$

De même 1 est beaucoup plus grand que $RC\omega = |jRC\omega|$ dès que :

$$1 \gg RC\omega \quad \text{soit} \quad \omega \ll \frac{1}{RC}, \quad \text{en pratique} \quad \omega < \frac{1}{10RC}.$$

Dès lors, $\frac{1}{RC}$ apparaît comme la pulsation caractéristique du système, qui délimite deux comportements limites ou asymptotiques :

domaine de pulsation	$\omega \ll \frac{1}{RC}$	$\frac{1}{RC} \ll \omega$
$\underline{H}(j\omega) \simeq$	1	$\frac{1}{jRC\omega}$

Les deux expressions limites de $\underline{H}(j\omega)$ sont appelées **fonctions de transfert asymptotiques**.

Interprétation des zones rectilignes de la courbe de gain Pour $\omega \ll \frac{1}{RC}$, on observe un segment horizontal. En effet, la fonction de transfert y vaut $\underline{H}(j\omega) \simeq 1$ et $20\log(1) = 0$. Le gain en décibel reste constamment nul.

Pour $\omega \gg \frac{1}{RC}$, on observe un segment oblique. Le gain en décibel vaut dans cet intervalle de pulsation :

$$G_{dB}(\omega) = 20\log\left|\frac{1}{jRC\omega}\right| = 20\log\left(\frac{1}{RC\omega}\right) = -20\log(RC\omega).$$

Que vaut-il une **décade** plus loin, c’est à dire quand la pulsation est multipliée par 10 ?

$$\begin{aligned} G_{dB}(10\omega) &= -20\log(RC \times 10\omega) = -20\log(10) - 20\log(RC\omega) \\ G_{dB}(10\omega) &= G_{dB}(\omega) - 20. \end{aligned}$$

Le gain diminue de 20 dB en une décade, la pente du segment oblique vaut donc **−20 dB/décade**. Une telle valeur de la pente est matérialisée sur le diagramme de Bode par l’expression −20 dB/décade en toutes lettres ou le nombre −1, sous-entendu pour −1 × 20 dB/décade. Une diminution de −20 dB/décade correspond à une amplitude du signal de sortie divisée par 10 lorsque la fréquence est multipliée par 10.

En résumé :

domaine de pulsation	$\omega \ll \frac{1}{RC}$	$\frac{1}{RC} \ll \omega$
$\underline{H}(j\omega) \simeq$	1	$\frac{1}{jRC\omega}$
pente	nulle	−20 dB/décade

Les deux zones rectilignes du diagramme de Bode sont formées des asymptotes, qui se croisent en la pulsation caractéristique du filtre. En effet, quand les deux asymptotes se croisent :

$$1 = \left| \frac{1}{jRC\omega} \right| \quad \text{implique} \quad 1 = \frac{1}{RC\omega}, \quad \text{soit} \quad \omega = \omega_0 = \frac{1}{RC}.$$

Interprétation des zones rectilignes de la courbe de phase Pour $\omega \ll \frac{1}{RC}$, on observe une phase $\varphi = 0$. En effet, la fonction de transfert y vaut $\underline{H}(j\omega) \simeq 1$ et $\varphi = \arg(1) = 0$.

Pour $\omega \gg \frac{1}{RC}$, on observe une phase $\varphi = -\frac{\pi}{2}$. La fonction de transfert vaut alors $\underline{H}(j\omega) \simeq \frac{1}{jRC\omega}$ et ainsi :

$$\varphi = \arg\left(\frac{1}{jRC\omega}\right) = \arg(1) - \arg(jRC\omega) = 0 - \frac{\pi}{2} = -\frac{\pi}{2}.$$

En résumé :

domaine de pulsation	$\omega \ll \frac{1}{RC}$	$\frac{1}{RC} \ll \omega$
$\underline{H}(j\omega) \simeq$	1	$\frac{1}{jRC\omega}$
$\varphi(\omega) \simeq$	0	$-\frac{\pi}{2}$

1.5 Tracé à la calculatrice

Comment tracer simplement un diagramme en échelles semilogarithmiques ? Il suffit de se souvenir qu’on trace le gain et la phase en fonction de $\log \omega$. On convertit alors les échelles linéaires en échelles logarithmiques avec la transformation $\omega = 10^x$, qui est la réciproque de $x = \log \omega$. On trace donc à la calculatrice, par exemple pour le gain, en fonction de x :

$$x \mapsto 20 \log \frac{1}{\sqrt{1 + (RC10^x)^2}}.$$

Lorsqu’on lit $x = 2$ sur l’axe des absisses, il faut ainsi comprendre $\omega = 10^2 \text{ rad.s}^{-1}$.

2 Différents types de fonction de transfert

Il existe un très grand nombre de filtres, seuls les plus simples seront abordés.

2.1 Filtres du premier ordre

a) Passe-bas et intégrateur

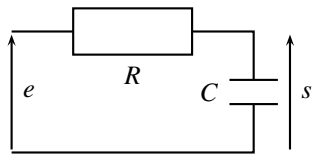


Figure 11.3 – Passe-bas du premier ordre.

La fonction de transfert du filtre étudié est, avec $RC = \tau$:

$$\underline{H}(j\omega) = \frac{1}{1 + j\tau\omega},$$

ou, avec la notation de Laplace vue en cours de SII ($p = j\omega$) :

$$H(p) = \frac{1}{1 + \tau p}.$$

La fonction de transfert se simplifie pour $\omega \ll \frac{1}{\tau}$ et $\omega \gg \frac{1}{\tau}$. Dans le cas où $\omega \gg \frac{1}{\tau}$:

$$\underline{H}(j\omega) \simeq \frac{1}{j\tau\omega},$$

et le signal d’entrée est intégré. En effet :

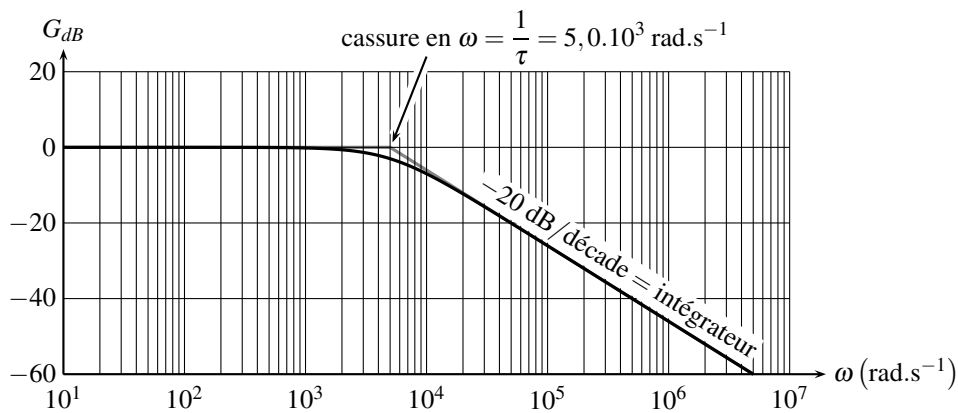
$$\frac{\underline{s}(t)}{\underline{e}(t)} = \underline{H}(j\omega) = \frac{1}{j\tau\omega} \quad \text{donc} \quad j\tau\omega \underline{s}(t) = \underline{e}(t),$$

soit en notation temporelle :

$$\tau \frac{ds}{dt} = e(t) \quad \text{ou} \quad s(t) = \frac{1}{\tau} \int e(t) dt.$$

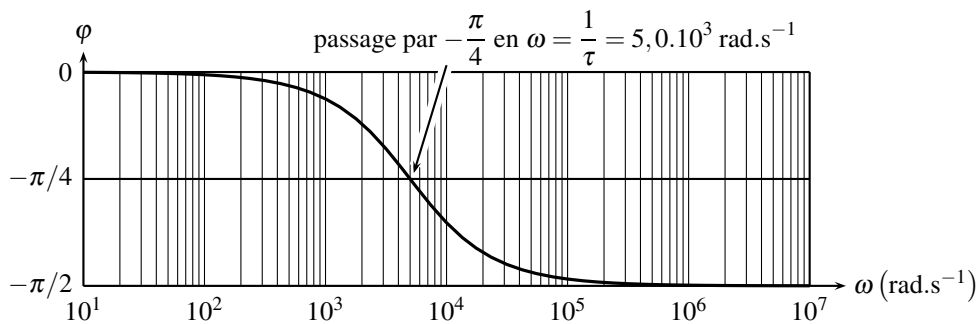
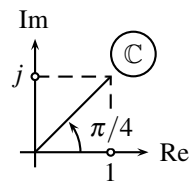
Cette équivalence a été vue dans le précédent chapitre, paragraphe 3.2. Ainsi, la sortie est l’intégrale de l’entrée, à condition que ω soit très supérieure à $\frac{1}{\tau}$, par exemple $\omega > \frac{10}{\tau}$.

Un filtre passe-bas se comporte en **intégrateur** pour tout signal de pulsation très supérieure à la pulsation caractéristique du filtre.



De plus, la phase vaut en $\omega = \frac{1}{\tau}$ vaut $-\frac{\pi}{4}$:

$$\underline{H}\left(\frac{j}{\tau}\right) = \frac{1}{1+j} \quad \text{donc} \quad \varphi\left(\frac{1}{\tau}\right) = \arg(1) - \arg(1+j) = 0 - \frac{\pi}{4}.$$



b) Passe-haut et dérivateur

L'établissement de la fonction de transfert s'effectue avec la formule du diviseur de tension. Pour le circuit RC (voir figure 11.6) :

$$\underline{H}(j\omega) = \frac{\underline{s}}{\underline{e}} = \frac{\underline{Z}_R}{\underline{Z}_R + \underline{Z}_C} = \frac{R}{R + \frac{1}{jC\omega}} = \frac{jRC\omega}{1 + jRC\omega},$$

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D’UN SYSTÈME LINÉAIRE

qui s’écrit sous forme canonique, en notant $\tau = RC$ la constante de temps du système :

$$\underline{H}(j\omega) = \frac{j\tau\omega}{1 + j\tau\omega},$$

ou en notation de Laplace

$$H(p) = \frac{\tau p}{1 + \tau p}.$$

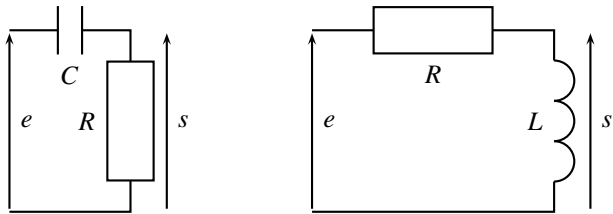


Figure 11.6 – Passe-haut du premier ordre.

De même, pour le circuit RL : $\underline{H}(j\omega) = \frac{\underline{Z}_L}{\underline{Z}_R + \underline{Z}_L} = \frac{jL\omega}{R + jL\omega}$, qui adopte la même forme avec la constante de temps $\tau = \frac{L}{R}$.

Le dénominateur de la fonction de transfert est la somme de deux termes ; comme pour un passe-bas, 1 est prédominant en module devant $j\tau\omega$ si $\omega \ll \frac{1}{\tau}$ et $j\tau\omega$ est prédominant en module devant 1 si $\omega \gg \frac{1}{\tau}$. $\frac{1}{\tau}$ est la pulsation caractéristique du système. Deux comportements asymptotiques ou limites apparaissent :

domaines de pulsation	$\omega \ll \frac{1}{\tau}$	$\frac{1}{\tau} \ll \omega$
numérateur =	$j\tau\omega$	$j\tau\omega$
dénominateur $\simeq$	1	$j\tau\omega$
$\underline{H}(j\omega) \simeq$	$j\tau\omega$	1

La phase est immédiatement interprétable dans le plan complexe :

$$\arg(j\tau\omega) = +\frac{\pi}{2} \left(\text{pour } \omega \ll \frac{1}{\tau} \right)$$
$$\arg(1) = 0 \left(\text{pour } \frac{1}{\tau} \ll \omega \right)$$

Quant au gain en décibel, il est nul pour $\omega \gg \frac{1}{\tau}$ car $20\log(1) = 0$ et il vaut pour $\omega \ll \frac{1}{\tau}$:

$$G_{dB}(\omega) = 20\log|j\tau\omega| = 20\log(\tau\omega),$$

il devient une décade plus loin :

$$G_{dB}(10\omega) = 20\log(10\tau\omega) = 20\log(\tau\omega) + 20\log(10).$$

Attendu que $\log(10) = 1$, le gain a augmenté de +20 dB en une décade, d'où une pente de +20 dB/décade. En résumé :

domaines de pulsation	$\omega \ll \frac{1}{\tau}$	$\frac{1}{\tau} \ll \omega$
$\underline{H}(j\omega) \simeq$	$j\tau\omega$	1
$\varphi(\omega) \simeq$	$\frac{\pi}{2}$	0
pente	+20 dB/décade	nulle

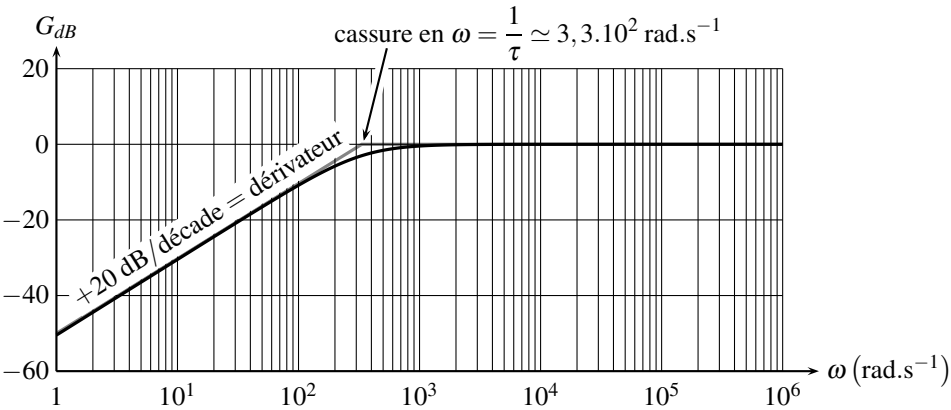


Figure 11.7 – Gain d'un passe-haut du premier ordre dans le cas où $\tau = 3,0.10^{-3}$ s (les asymptotes sont en gris).

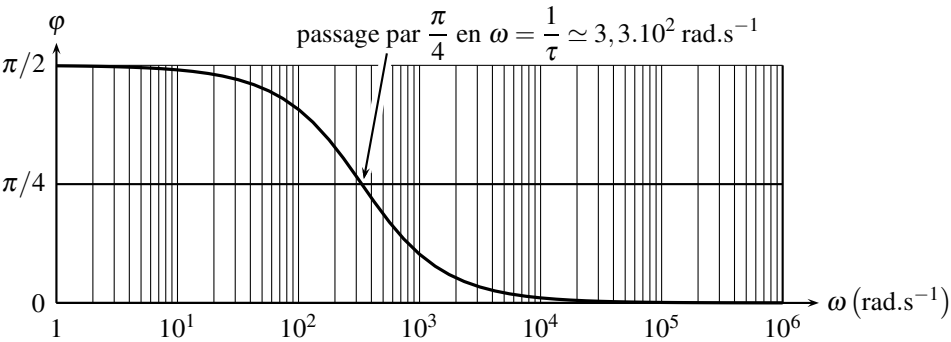


Figure 11.8 – Phase d'un passe-haut du premier ordre dans le cas où $\tau = 3,0.10^{-3}$ s.

La phase vaut $\frac{\pi}{4}$ en $\omega = \frac{1}{\tau}$: $\underline{H}\left(\frac{j}{\tau}\right) = \frac{j}{1+j}$; $\varphi\left(\frac{1}{\tau}\right) = \arg(j) - \arg(1+j) = \frac{\pi}{2} - \frac{\pi}{4} = \frac{\pi}{4}$.

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D’UN SYSTÈME LINÉAIRE

Le filtre se comporte en **dérivateur** pour $\omega \ll \frac{1}{\tau}$, soit $\omega \ll \frac{1}{10\tau}$. En effet, dans cette zone de fréquence :

$$\frac{s}{e} = \underline{H}(j\omega) \simeq j\tau\omega \quad \text{donc} \quad \underline{s} = j\tau\omega \underline{e},$$

soit en notation temporelle $s(t) = \tau \frac{de}{dt}$.

Un filtre passe-haut se comporte en **dérivateur** pour un signal dont la pulsation est très inférieure à celle du filtre.

c) Complément : autre exemple de filtre passe-bas

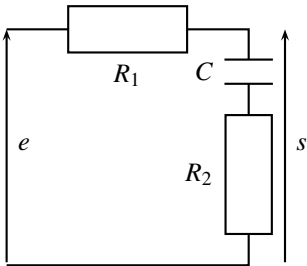


Figure 11.9 – Passe-bas du premier ordre.

On reconnaît un diviseur de tension :

$$\underline{H}(j\omega) = \frac{s}{e} = \frac{\underline{Z}_C + \underline{Z}_{R_2}}{\underline{Z}_C + \underline{Z}_{R_2} + \underline{Z}_{R_1}} = \frac{\frac{1}{jC\omega} + R_2}{\frac{1}{jC\omega} + R_2 + R_1} = \frac{1 + jR_2C\omega}{1 + j(R_1 + R_2)C\omega}.$$

Ce système présente deux constantes de temps, $\tau_n = R_2C$ au numérateur et $\tau_d = (R_1 + R_2)C$ au dénominateur. Le numérateur et le dénominateur sont simplifiables, en fonction des valeurs de ω par rapport à $\frac{1}{\tau_n}$ et $\frac{1}{\tau_d}$.

domaines de pulsation	$\omega \ll \frac{1}{(R_1 + R_2)C}$	$\frac{1}{(R_1 + R_2)C} \ll \omega \ll \frac{1}{R_2C}$	$\frac{1}{R_2C} \ll \omega$
numérateur $\simeq$	1	1	$jR_2C\omega$
dénominateur $\simeq$	1	$j(R_1 + R_2)C\omega$	$j(R_1 + R_2)C\omega$
$\underline{H}(j\omega) \simeq$	1	$\frac{1}{j(R_1 + R_2)C\omega}$	$\frac{R_1}{R_1 + R_2}$
$\varphi(\omega) \simeq$	0	$-\frac{\pi}{2}$	0
pente	nulle	-20 dB/décade	nulle

Apparaissent alors trois comportements limites, ou asymptotiques, qui correspondent aux trois zones rectilignes du diagramme de Bode (voir figures 11.10 et 11.11). Comme le montre le tableau de la page suivante, ce filtre se comporte en **intégrateur** dans la zone $\frac{1}{(R_1 + R_2)C} \ll \omega \ll \frac{1}{R_2C}$ car la fonction de transfert est en $\frac{1}{j\tau\omega}$. Toutefois, ce comportement reste peu visible si les deux pulsations limites sont trop proches l’une de l’autre.

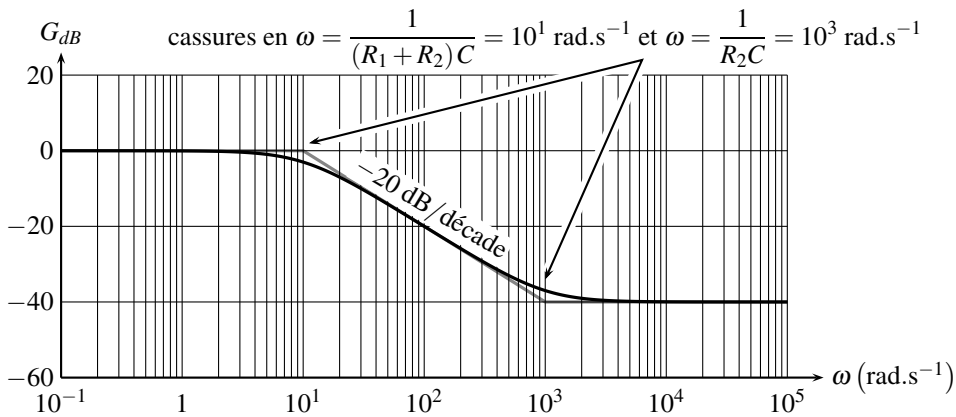


Figure 11.10 – Gain du système avec $R_2C = 10^{-3} \text{ rad.s}^{-1}$ et $(R_1 + R_2)C = 10^{-1} \text{ rad.s}^{-1}$ (les asymptotes sont en gris).

Sur la figure 11.11, attendu l’étroitesse de la zone asymptotique à $-\frac{\pi}{2}$, la phase ne parvient pas à son expression asymptotique.

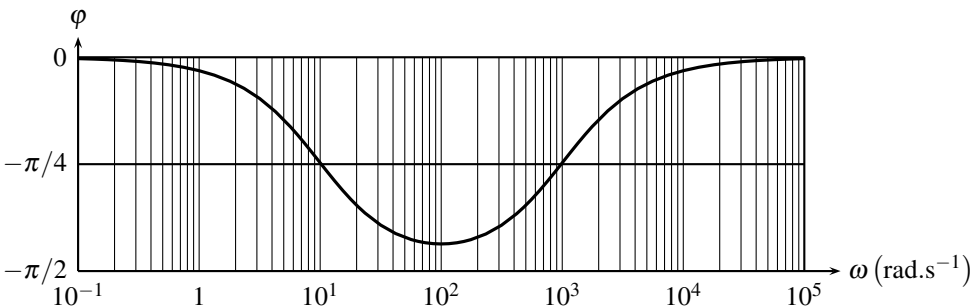


Figure 11.11 – Phase du système avec $R_2C = 10^{-3} \text{ rad.s}^{-1}$ et $(R_1 + R_2)C = 10^{-1} \text{ rad.s}^{-1}$.

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D’UN SYSTÈME LINÉAIRE

d) Complément : déphaseur

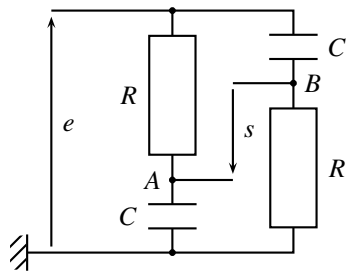


Figure 11.12 – Passe-tout déphaseur du premier ordre.

La sortie est prise entre les nœuds A et B dont les potentiels sont fournis par la formule du diviseur de tensions :

$$\underline{V}_A = \frac{\underline{Z}_C}{\underline{Z}_C + \underline{Z}_R} = \frac{\frac{1}{jC\omega}}{\frac{1}{jC\omega} + R} \underline{e}(t) = \frac{1}{1 + jRC\omega} \underline{e}(t),$$
$$\underline{V}_B = \frac{\underline{Z}_R}{\underline{Z}_R + \underline{Z}_C} = \frac{R}{R + \frac{1}{jC\omega}} \underline{e}(t) = \frac{jRC\omega}{1 + jRC\omega} \underline{e}(t).$$

Ainsi $\underline{s}(t) = \underline{U}_{BA} = \underline{V}_A - \underline{V}_B$:

$$\frac{\underline{s}}{\underline{e}} = \frac{1}{1 + jRC\omega} - \frac{jRC\omega}{1 + jRC\omega} = \frac{1 - jRC\omega}{1 + jRC\omega},$$

qu’on écrit sous forme canonique avec $\tau = RC$, constante de temps du système :

$$\underline{H}(j\omega) = \frac{1 - j\tau\omega}{1 + j\tau\omega}, \text{ ou en notation de Laplace } H(p) = \frac{1 - \tau p}{1 + \tau p}.$$

Le gain en décibel du système est toujours nul, quelque soit la pulsation, d’où le nom de **passe-tout** : $G_{dB} = 0$. En effet, $\underline{H}$ est le rapport de deux complexes conjugués, donc de même module ; leur rapport vaut 1 et leur log est nul.

La phase, quant à elle, varie. Le numérateur et le dénominateur présentent tous deux la même constante de temps τ , ce qui permet d’asymptotiquement simplifier la fonction de transfert :

domaines de pulsation	$\omega \ll \frac{1}{\tau}$	$\frac{1}{\tau} \ll \omega$
numérateur $\simeq$	1	$-j\tau\omega$
dénominateur $\simeq$	1	$j\tau\omega$
$\underline{H}(j\omega) \simeq$	1	-1
$\varphi(\omega) \simeq$	0	$-\pi$

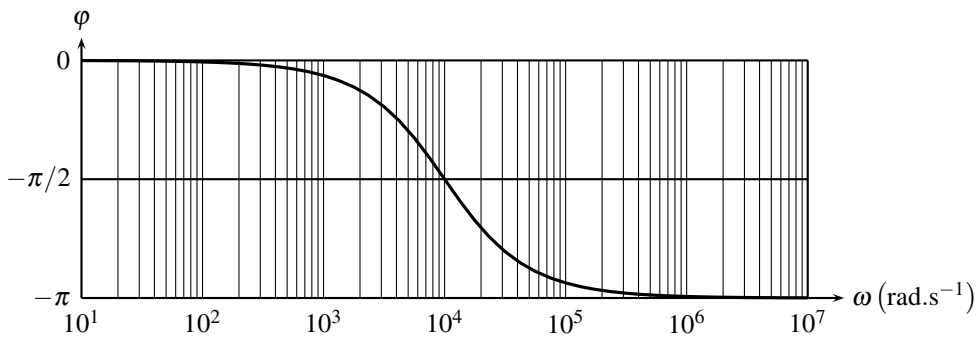


Figure 11.13 – Phase d'un passe-tout déphaseur du premier ordre dans le cas où $\tau = 10^{-4}\text{s}$.

2.2 Filtres du deuxième ordre

a) Passe-bas

Fonction de transfert Avec la formule du diviseur de tension :

$$\underline{H}(j\omega) = \frac{\underline{s}}{\underline{e}} = \frac{\underline{Z_C}}{\underline{Z_C} + \underline{Z_R} + \underline{Z_L}} = \frac{\frac{1}{jC\omega}}{\frac{1}{jC\omega} + R + jL\omega} = \frac{1}{1 + jRC\omega + LC(j\omega)^2}.$$

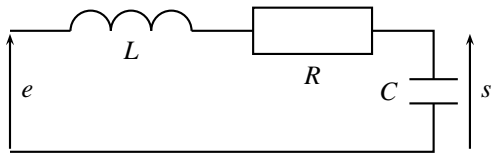


Figure 11.14 – Passe-bas du deuxième ordre.

Avec $LC = \frac{1}{\omega_0^2}$ et $RC = \frac{2\xi}{\omega_0} = \frac{1}{Q\omega_0}$, la fonction de transfert se met sous forme canonique :

$$\underline{H}(j\omega) = \frac{1}{1 + 2\xi \frac{j\omega}{\omega_0} + \left(\frac{j\omega}{\omega_0}\right)^2}$$

ou

$$H(p) = \frac{1}{1 + 2\xi \frac{p}{\omega_0} + \frac{p^2}{\omega_0^2}}.$$

Asymptotes Le dénominateur, somme de trois termes, est simplifiable suivant les valeurs de ω par rapport à la pulsation caractéristique ω_0 du système. Si ω est très inférieur à ω_0 , alors les termes en $\frac{\omega}{\omega_0}$ et $\left(\frac{\omega}{\omega_0}\right)^2$ sont négligeables en module devant 1 ; et si ω est très supérieur à ω_0 , alors le terme prédominant en module est le terme de plus haut degré en $\left(\frac{\omega}{\omega_0}\right)^2$:

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D’UN SYSTÈME LINÉAIRE

	$\omega \ll \omega_0$	$\omega_0 \ll \omega$
dénominateur $\simeq$	1	$\left(\frac{j\omega}{\omega_0}\right)^2$
$\underline{H}(j\omega) \simeq$	1	$\frac{\omega_0^2}{(j\omega)^2}$

Deux zones rectilignes apparaissent dès lors sur le diagramme de Bode. Pour $\omega \ll \omega_0$, la phase est nulle car $\arg(1) = 0$; pour $\omega \gg \omega_0$:

$$\arg\left(\frac{\omega_0}{j\omega}\right)^2 = 2 \arg(\omega_0) - 2 \arg(j\omega) = 0 - 2 \times \frac{\pi}{2} = -\pi.$$

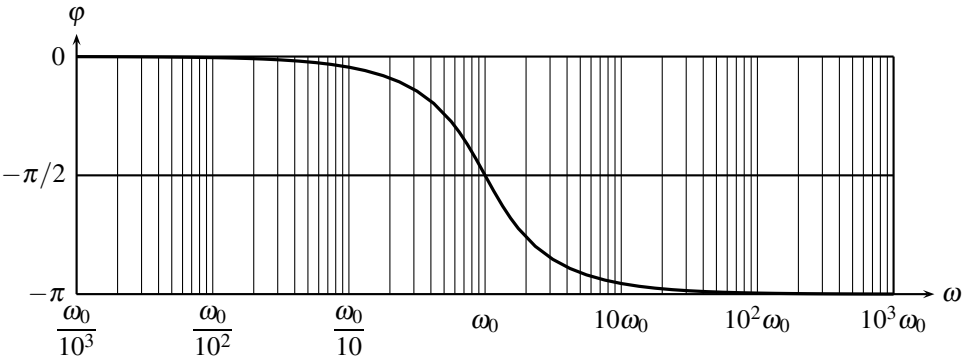


Figure 11.15 – Phase d’un passe-bas du deuxième ordre (tracé pour $\xi = 0,70$).

Quant au gain en décibel, il est nul pour $\omega \ll \omega_0$; pour $\omega \gg \omega_0$:

$$G_{dB}(\omega) = 20 \log \left| -\frac{\omega_0^2}{\omega^2} \right| = 40 \log \left(\frac{\omega_0}{\omega} \right),$$

il devient une décade plus loin :

$$G_{dB}(10\omega) = 40 \log \left(\frac{\omega_0}{10\omega} \right) = 40 \log \left(\frac{\omega_0}{\omega} \right) - 40 \log(10) = G_{dB}(\omega) - 40.$$

Le gain a diminué de 40 dB en une décade, d’où une pente de **−40 dB/décade**. Une diminution de −40 dB/décade correspond à une amplitude du signal de sortie divisée par 100 lorsque la fréquence est multipliée par 10.

Résonance On observe la possibilité d’une **résonance** si le système est peu amorti, soit $\xi < \sqrt{2}/2$ ou $Q > \sqrt{2}/2$, comme au chapitre sur les régimes harmoniques. Plus ξ est faible, c’est à dire moins le circuit est amorti, plus la résonance est aigue.

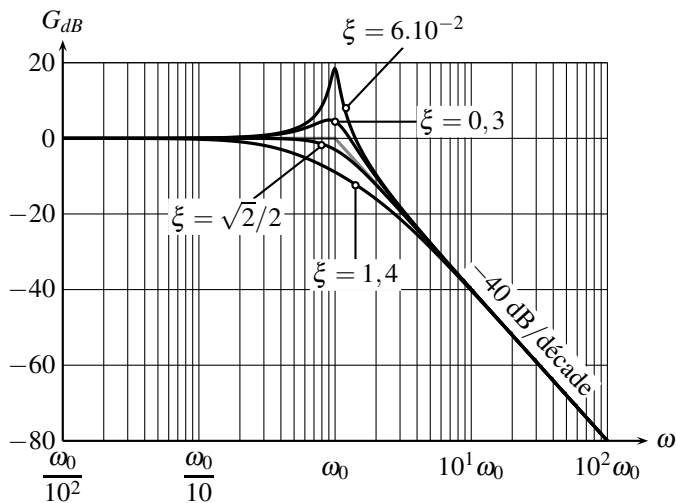


Figure 11.16 – Gain d'un passe-bas du deuxième ordre (les asymptotes sont en gris).

b) Passe-bande

Fonction de tranfert Avec la formule du diviseur de tension :

$$\underline{H}(j\omega) = \frac{\underline{s}}{\underline{e}} = \frac{\underline{Z}_R}{\underline{Z}_R + \underline{Z}_L + \underline{Z}_C} = \frac{R}{R + jL\omega + \frac{1}{jC\omega}} = \frac{jRC\omega}{1 + jRC\omega + LC(j\omega)^2}.$$

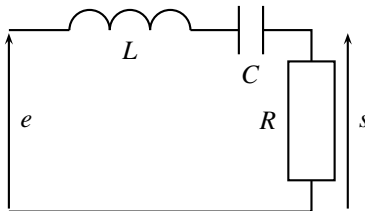


Figure 11.17 – Passe-bande du deuxième ordre.

Avec les paramètres caractéristiques d'un deuxième ordre, $LC = \frac{1}{\omega_0^2}$ et $RC = \frac{2\xi}{\omega_0} = \frac{1}{Q\omega_0}$:

$$\underline{H}(j\omega) = \frac{2\xi \frac{j\omega}{\omega_0}}{1 + 2\xi \frac{j\omega}{\omega_0} + \left(\frac{j\omega}{\omega_0}\right)^2}$$

ou

$$H(p) = \frac{2\xi \frac{p}{\omega_0}}{1 + 2\xi \frac{p}{\omega_0} + \frac{p^2}{\omega_0^2}}.$$

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D’UN SYSTÈME LINÉAIRE

Remarque

Le caractère passe-bande est dû à la mise en série de la bobine et du condensateur. En effet, l’impédance de ce dipôle est : $Z = jL\omega + \frac{1}{jC\omega}$, qui est nulle en $\omega = \omega_0 = \frac{1}{\sqrt{LC}}$:
 $Z(\omega_0) = j\frac{L}{\sqrt{LC}} + \frac{\sqrt{LC}}{jC} = j\sqrt{\frac{L}{C}} - j\sqrt{\frac{L}{C}} = 0$. Ainsi, en $\omega = \omega_0$, l’ensemble {bobine + condensateur} est équivalent à un fil de résistance nulle qui relie directement la sortie à l’entrée. La pulsation ω_0 passe donc sans aucune atténuation ni déphasage à travers le filtre.

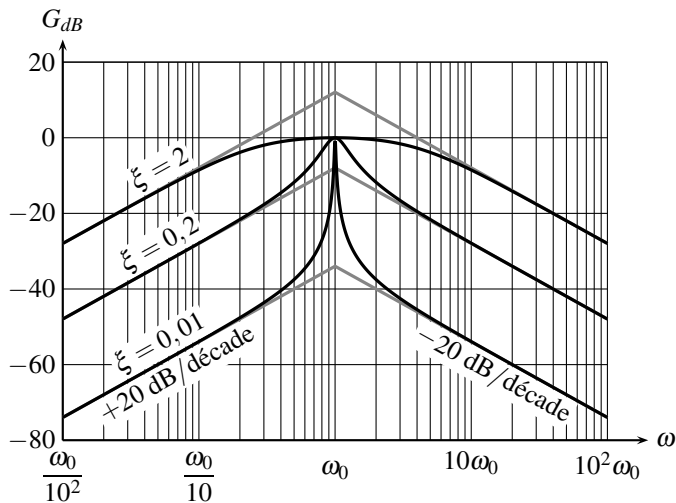


Figure 11.18 – Gain d’un passe-bande du deuxième ordre (les expressions asymptotiques sont en gris).

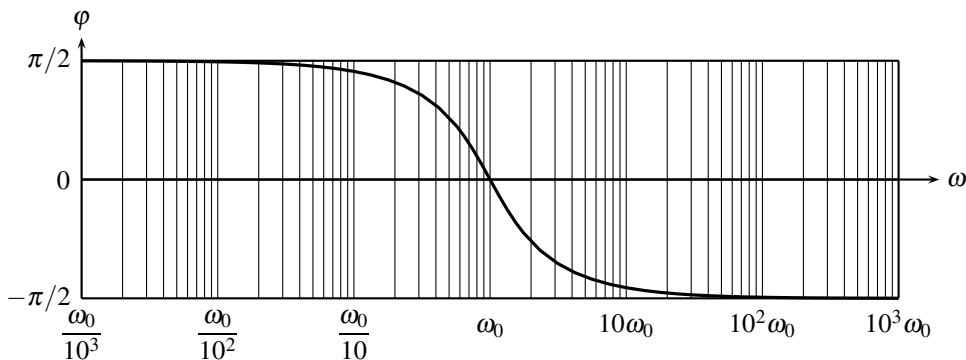


Figure 11.19 – Phase d’un passe-bande du deuxième ordre (tracé pour $\xi = 0,70$).

Asymptotes Le dénominateur, somme de trois termes, est simplifiable suivant les valeurs de ω par rapport à la pulsation caractéristique ω_0 du système. Si ω est très inférieur à ω_0 , alors 1 prédomine et si ω est très supérieur à ω_0 , alors $\left(\frac{\omega}{\omega_0}\right)^2$ prédomine. Deux zones asymptotiques apparaissent :

domaines de pulsation	$\omega \ll \omega_0$	$\omega_0 \ll \omega$
numérateur =	$2\xi \frac{j\omega}{\omega_0}$	$2\xi \frac{j\omega}{\omega_0}$
dénominateur $\simeq$	1	$\left(\frac{j\omega}{\omega_0}\right)^2$
$\underline{H}(j\omega) \simeq$	$2\xi \frac{j\omega}{\omega_0}$	$\frac{2\xi \omega_0}{j\omega}$

Pour la phase, si $\omega \ll \omega_0$ alors : $\arg(\underline{H}(j\omega)) = \arg\left(2\xi \frac{j\omega}{\omega_0}\right) = +\frac{\pi}{2}$;

et si $\omega \gg \omega_0$: $\arg(\underline{H}(j\omega)) = \arg\left(\frac{2\xi \omega_0}{j\omega}\right) = \arg(2\xi \omega_0) - \arg(j\omega) = 0 - \frac{\pi}{2}$.

Quant au gain en décibel, le terme $j\omega$ impose une pente de +20 dB/décade pour $\omega \ll \omega_0$ et $\frac{1}{j\omega}$ une pente de -20 dB/décade pour $\omega \gg \omega_0$.

domaines de pulsation	$\omega \ll \omega_0$	$\omega_0 \ll \omega$
$\varphi(\omega) \simeq$	$+\frac{\pi}{2}$	$-\frac{\pi}{2}$
pente	+20 dB/décade	-20 dB/décade

Le passe-bande présente un caractère **dérivateur** pour $\omega \ll \omega_0$ car l'expression asymptotique de la fonction de transfert y est proportionnelle à $j\omega$, et **intégrateur** pour $\omega \gg \omega_0$ à cause du terme en $\frac{1}{j\omega}$.

Résonance Le gain est maximum en $\omega = \omega_0$ et est nul, ainsi que la phase :

$$\underline{H}(j\omega_0) = \frac{2\xi \frac{j\omega_0}{\omega_0}}{1 + 2\xi \frac{j\omega_0}{\omega_0} + \left(\frac{j\omega_0}{\omega_0}\right)^2} = 1 \quad \text{donc} \quad \begin{cases} G_{dB}(\omega_0) = 0 \\ \varphi(\omega_0) = 0. \end{cases}$$

De plus, le gain asymptotique en ω_0 , c'est à dire le gain des asymptotes en ω_0 , vaut :

$$20 \log \left| 2\xi \frac{j\omega_0}{\omega_0} \right| = 20 \log(2\xi) = -20 \log(Q).$$

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D’UN SYSTÈME LINÉAIRE

Moins le filtre est amorti (ξ faible ou Q important), plus les asymptotes se croisent bas sur le diagramme de Bode en gain.

Moins le filtre est amorti (ξ faible ou Q grand), plus il est sélectif, c’est à dire qu’il ne laisse passer d’une bande de fréquence resserrée autour de ω_0 . Cette constatation légitime le nom de **facteur de qualité** donné à Q .

c) Complément : passe-haut

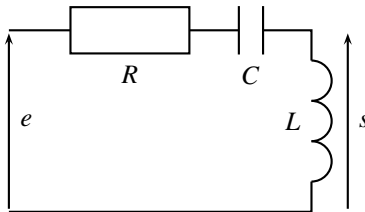


Figure 11.20 – Passe-haut du deuxième ordre.

La formule du diviseur de tension implique :

$$\underline{H}(j\omega) = \frac{\underline{s}}{\underline{e}} = \frac{\underline{Z_L}}{\underline{Z_R} + \underline{Z_L} + \underline{Z_C}} = \frac{jL\omega}{R + jL\omega + \frac{1}{jC\omega}} = \frac{LC(j\omega)^2}{1 + jRC\omega + LC(j\omega)^2}.$$

Avec les paramètres caractéristiques d’un deuxième ordre, ω_0 et ξ (ou $Q = 1/2\xi$) :

$$\underline{H}(j\omega) = \frac{\left(\frac{j\omega}{\omega_0}\right)^2}{1 + 2\xi\frac{j\omega}{\omega_0} + \left(\frac{j\omega}{\omega_0}\right)^2}, \quad \text{ou} \quad H(p) = \frac{\frac{p^2}{\omega_0^2}}{1 + 2\xi\frac{p}{\omega_0} + \frac{p^2}{\omega_0^2}}.$$

Le dénominateur d’un deuxième ordre est simplifiable :

domaines de pulsation	$\omega \ll \omega_0$	$\omega_0 \ll \omega$
numérateur =	$\left(\frac{j\omega}{\omega_0}\right)^2$	$\left(\frac{j\omega}{\omega_0}\right)^2$
dénominateur $\simeq$	1	$\left(\frac{j\omega}{\omega_0}\right)^2$
$\underline{H}(j\omega) \simeq$	$\left(\frac{j\omega}{\omega_0}\right)^2$	1

La phase est nulle pour $\omega \gg \omega_0$ car $\arg(1) = 0$. Pour $\omega \ll \omega_0$:

$$\arg\left(\frac{j\omega}{\omega_0}\right)^2 = 2\arg\left(\frac{j\omega}{\omega_0}\right) = 2 \times \frac{\pi}{2} = \pi.$$

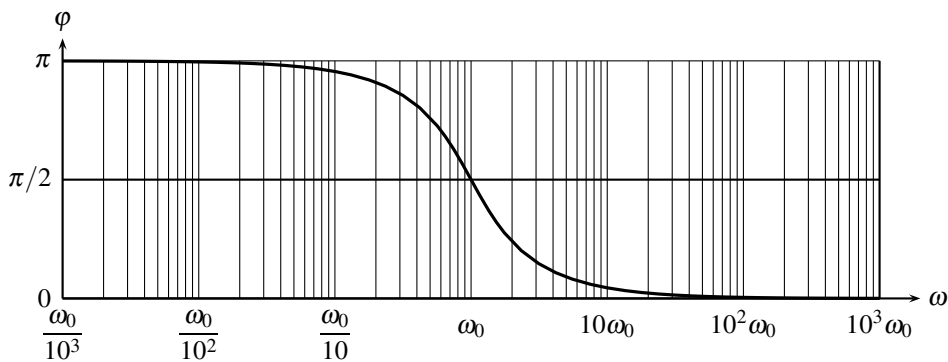


Figure 11.21 – Phase d’un passe-haut du deuxième ordre (tracé pour $\xi = 0,70$).

Quant au gain en décibel, il vaut pour $\omega \gg \omega_0$:

$$G_{dB}(\omega) = 20 \log \left| -\frac{\omega^2}{\omega_0^2} \right| = 20 \log (\omega^2) - 20 \log (\omega_0^2) = 40 \log (\omega) - 20 \log (\omega_0^2),$$

il devient une décade plus loin :

$$\begin{aligned} G_{dB}(10\omega) &= 40 \log (\omega) - 20 \log (\omega_0^2) = 40 \log (\omega) + 40 \log (10) - 20 \log (\omega_0^2) \\ G_{dB}(10\omega) &= G_{dB}(\omega) + 40. \end{aligned}$$

Le gain a augmenté de 40 dB en une décade, d’où une pente de **+40 dB/décade**.

Un phénomène de **résonance** est visible dès que $\xi < \sqrt{2}/2$ ou $Q > \sqrt{2}/2$. Plus ξ est faible (ou Q important), moins le système est amorti, plus la résonance est aigue.

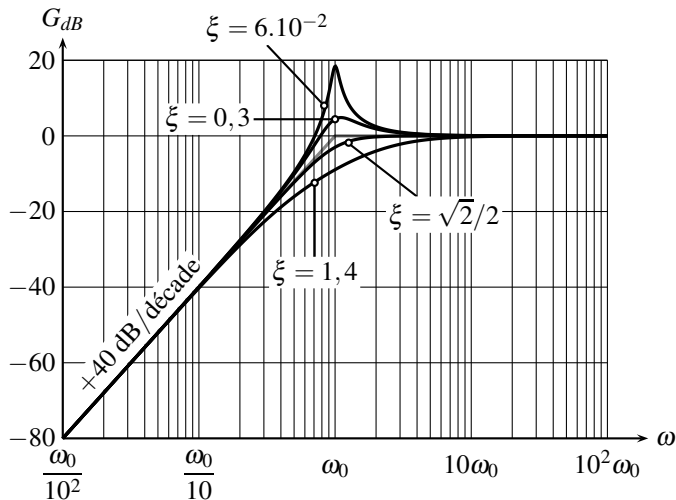


Figure 11.22 – Gain d’un passe-haut du deuxième ordre (les asymptotes sont en gris).

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D’UN SYSTÈME LINÉAIRE

d) Complément : coupe-bande

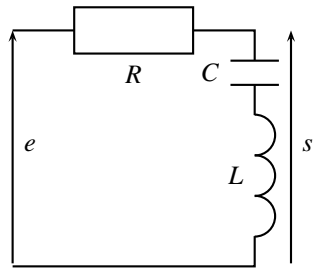


Figure 11.23 – Coupe-bande du deuxième ordre.

Avec la formule du diviseur de tension :

$$\underline{H}(j\omega) = \frac{s}{e} = \frac{\underline{Z}_L + \underline{Z}_C}{\underline{Z}_R + \underline{Z}_L + \underline{Z}_C} = \frac{jL\omega + \frac{1}{jC\omega}}{R + jL\omega + \frac{1}{jC\omega}} = \frac{1 + LC(j\omega)^2}{1 + jRC\omega + LC(j\omega)^2}.$$

Le numérateur de la fonction de transfert s’annule pour :

$$1 - LC\omega^2 = 0 \quad \text{soit} \quad \omega = \frac{1}{\sqrt{LC}} = \omega_0,$$

ce qui entraîne un zéro de transmission à cette pulsation.

Le caractère coupe-bande est dû à la mise en série de la bobine et du condensateur à la sortie.

L’impédance de ce dipôle s’annule en $\omega_0 = \frac{1}{\sqrt{LC}}$:

$$\underline{Z}(\omega_0) = jL\omega_0 + \frac{1}{jC\omega_0} = 0.$$

Ainsi, en $\omega = \omega_0$, l’ensemble {bobine + condensateur} est équivalent à un fil de résistance nulle qui court-circuite la sortie. La pulsation ω_0 est donc complètement éliminée par le filtre.

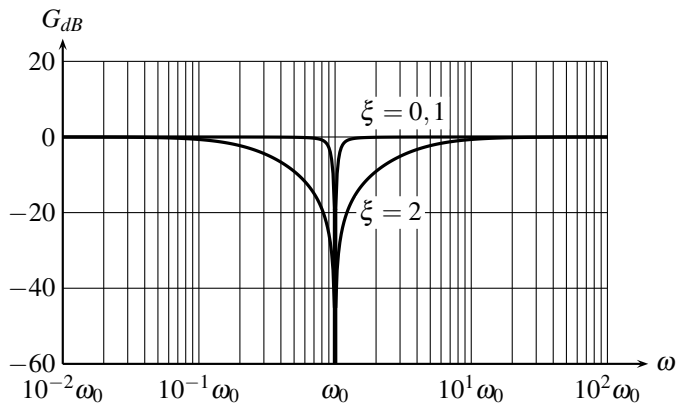


Figure 11.24 – Gain d’un coupe-bande du deuxième ordre.

Moins le filtre est amorti (ξ faible ou Q grand), plus il est sélectif, c'est à dire que la bande de fréquence éliminée est resserrée autour de ω_0 . Cette constatation légitime le nom de facteur de qualité donné à Q .

2.3 Récapitulatif

Comment interpréter systématiquement les zones rectilignes des diagrammes de Bode ? Il suffit :

- d'identifier les pulsations caractéristiques des numérateur et dénominateur,
- de préciser chaque intervalle de pulsation entre ces pulsations caractéristiques,
- de simplifier numérateur et le dénominateur de la fonction de transfert en ne gardant que le terme prédominant, dans chaque intervalle de pulsation,
- d'interpréter l'expression asymptotique ainsi obtenue en se souvenant du tableau de correspondances, appelé relations de Bayard-Bode :

$\underline{H}_{asymptotique}$	phase	pente	comportement	réalisation pratique
$j\tau\omega$ ou τp	$+\frac{\pi}{2}$	+20 dB/décade	dérivateur	passé-haut premier ordre passé-bande deuxième ordre
$\frac{1}{j\tau\omega}$ ou $\frac{1}{\tau p}$	$-\frac{\pi}{2}$	-20 dB/décade	intégrateur	passé-bas premier ordre passé-bande deuxième ordre

$\underline{H}_{asymptotique}$	phase	pente	réalisation pratique
$(j\tau\omega)^2$ ou $(\tau p)^2$	$+\pi$	+40 dB/décade	passé-haut deuxième ordre
$\frac{1}{(j\tau\omega)^2}$ ou $\frac{1}{(\tau p)^2}$	$-\pi$	-40 dB/décade	passé-bas deuxième ordre

$\underline{H}_{asymptotique}$	phase	pente
1	0	nulle
-1	$-\pi$	nulle

3 Gabarit (PTSI)

3.1 Amplitude

La fonction d'un système est de récupérer le signal utile puis de le manipuler pour en retirer une information précise. Par exemple un téléphone portable capte une seule communication au sein d'un grand nombre de signaux électromagnétiques ; puis il traite l'information contenue dans ce signal pour la transformer, *in fine*, en signal sonore. Le système devrait donc idéalement éliminer tous les signaux inutiles pour transmettre, sans atténuation ni déphasage, le signal utile, puis traiter l'information. Il est alors nommé **filtre** car il filtre certains signaux. Les systèmes réels ne présentent toutefois pas de telles performances. En particulier, le signal qu'ils laissent passer est légèrement atténué et déphasé ; de plus, le passage entre les

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D’UN SYSTÈME LINÉAIRE

fréquences transmises et les fréquences absorbées n’est pas brutal, mais progressif.

On est donc amené à définir un gabarit du filtre à construire. Un filtre **passé-bas** a pour fonction de laisser passer toutes les fréquences inférieures à f_p (indice p pour passante) et d’éliminer toutes les fréquences supérieures à f_a (indice a pour atténuée). On définit alors :

- la dernière **fréquence passante** f_p (le filtre doit laisser passer les fréquences $f < f_p$),
- le **gain minimum** G_{sup} , en dB, pour les fréquences qui passent à travers le filtre (il ne doit pas être trop faible, car ces fréquences doivent sortir du filtre sans être atténuées),
- la première **fréquence atténuée** f_a (toutes les fréquences $f > f_a$ doivent être éliminées par le filtre),
- le **gain maximum** G_{inf} pour les fréquences atténuées. Le gain du filtre à ces fréquences doit être inférieur à G_{inf} pour être sûr qu’elles soient éliminées.

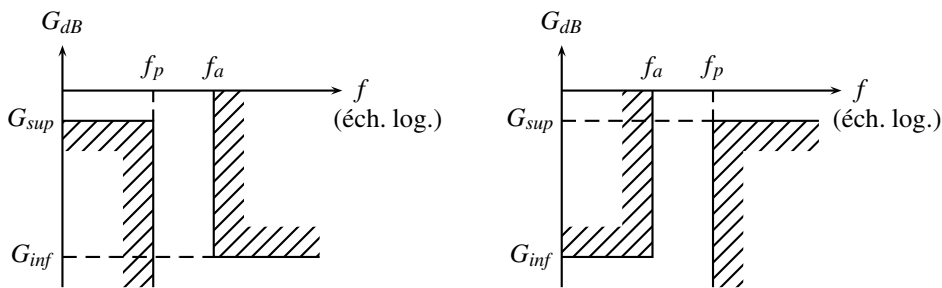


Figure 11.25 – Gabarit (a) d’un passe-bas (b) d’un passe-haut.

Par exemple, si on choisit $G_{inf} = -60$ dB, alors :

$$20 \log \left| \frac{s}{e} \right| < -60 \quad \text{donc} \quad |s| < |e| \times 10^{-3}.$$

L’amplitude des fréquences supérieures à f_a sont au moins divisées par 1000 ; elles sont effectivement atténuées par le filtre.

Et si on décide que $G_{sup} = -3$ dB, en ne dépassant pas 0 dB pour les fréquences transmises :

$$0 > 20 \log \left| \frac{s}{e} \right| > -3 \quad \text{donc} \quad |e| > |s| > |e| \times 0,71.$$

L’amplitude des fréquences inférieures à f_p diminue peu.

La plage de fréquences utiles, qui passent à travers le filtre, ici l’intervalle $[0, f_p]$, est nommée **bande passante** à G_{sup} . Dans cette définition, la valeur la plus souvent retenue pour G_{sup} est -3 dB, mais on rencontre aussi, plus rarement, -6 dB.

En résumé, une **passé-bas** transmet les basses fréquences et atténue les hautes. C’est le contraire pour un **passé-haut** (voir figure 11.25), qui atténue les basses fréquences, inférieures à f_a , et laisse passer les hautes, supérieures à f_p . On définit de même pour ce filtre le gain maximum G_{inf} pour les fréquences atténuées, le gain minimum G_{sup} pour les fréquences transmises, la bande passante $[f_p, +\infty[$ à G_{sup} .

Une autre donnée importante qui intervient dans le gabarit du filtre, est la régularité de sa réponse dans sa bande passante. Le gain y varie-t-il rapidement ? Y a-t-il une **résonance** ? Ce phénomène n'existe que pour les systèmes au moins d'ordre deux. Elle est à éviter si on ne souhaite pas amplifier une zone autour de la fréquence de résonance. On augmentera alors le coefficient d'amortissement ξ du filtre afin de l'éliminer.

Il convient de compléter la taxinomie des filtres par le **passe-bande** et le **coupe-bande**, aussi nommé **réjecteur de bande** (voir figure 11.26).

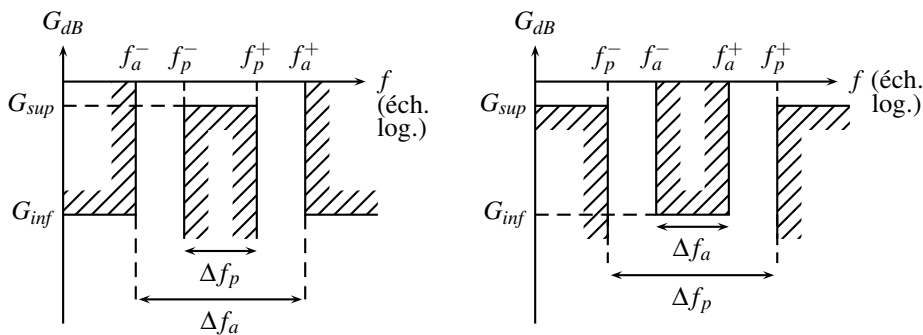


Figure 11.26 – Gabarit (c) d'un passe-bande (d) d'un coupe-bande.

Le **passe-bande** ne laisse passer qu'une bande de fréquence et atténue toutes les autres. Les fréquences comprises dans $[f_p^-, f_p^+]$ passent à travers le filtre, avec un gain G_{sup} , qui ne doit pas être trop faible pour qu'elles ne soient pas atténuées. Les fréquences inférieures à f_a^- ou f_a^+ sont atténuées, avec un gain inférieur à G_{inf} qu'on prendra suffisamment négatif.

Quant au **coupe-bande**, il atténue les signaux dont les fréquences sont comprises entre f_a^- et f_a^+ et laisse passer les autres.

3.2 Phase

La paragraphe précédent s'intéressait à l'atténuation en amplitude d'une partie du signal, pour ne laisser passer à travers le système que la partie utile du signal. Mais tout filtre introduit un certain déphasage, qui dépend de la fréquence, y compris dans sa bande passante. Ce déphasage est équivalent à un décalage temporel. En effet, pour un signal sinusoïdal $s(t)$:

$$s(t) = S_0 \cos(\omega t + \psi) = S_0 \cos\left(\omega\left(t + \frac{\psi}{\omega}\right)\right) = S_0 \cos(\omega(t + \Delta t)),$$

le déphasage ψ est équivalent au décalage $\Delta t = \frac{\psi}{\omega}$ introduit sur le signal lors du passage dans le filtre.

Dès lors, si un signal, composé de plusieurs fréquences, passe à travers un filtre, chacune de ses harmoniques est déphasée différemment et arrive en sortie du filtre avec un retard différent.

On ne sait pas construire de filtre analogique simple qui présente simultanément un filtrage idéal en amplitude sans déformer le signal utile par déphasage. On construit donc en pratique

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D’UN SYSTÈME LINÉAIRE

les filtres en ne tenant compte que du gabarit en gain, sans tenir compte de la déformation due au déphasage.

Toutefois, si le déphasage devient primordial, en particulier si le régime transitoire du filtre est important, on construira un filtre qui présente une grande régularité en déphasage, c’est-à-dire un retard Δt identique pour chaque fréquence, en s’accommodant de la piètre qualité du filtrage.

3.3 Exemple de réalisation d’un gabarit

Tous les filtres servent à assurer une fonction exposée dans un cahier des charges. Il convient de choisir le bon filtre qui entre dans le gabarit souhaité.

On cherche, par exemple, à effectuer le filtrage passe-bas suivant (voir figure 11.27) :

- les pulsations jusqu’à $1,2 \cdot 10^4 \text{ rad.s}^{-1}$ doivent passer à travers le filtre, sans être amplifiées, et ne pas être atténuées de plus de 3 dB,
- les pulsations supérieures à $3,9 \cdot 10^4 \text{ rad.s}^{-1}$ doivent être filtrées et atténuées d’au moins 15 dB.

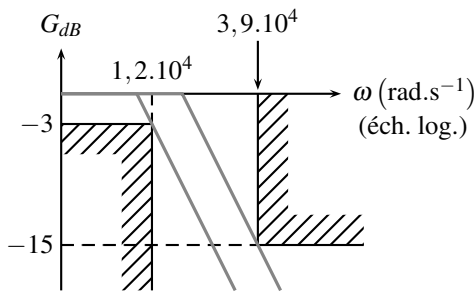


Figure 11.27 – Gabarit du filtre et gains asymptotiques limites en gris.

Quel est l’ordre du filtre à utiliser ? On le détermine avec la pente minimale dans la zone de transition, ici $[1,2 \cdot 10^4 \text{ rad.s}^{-1} ; 3,9 \cdot 10^4 \text{ rad.s}^{-1}]$:

- combien de décade y a-t-il dans la zone de transition ? La fréquence entre le début et la fin de cette zone est multipliée par :

$$\frac{3,9 \cdot 10^4}{1,2 \cdot 10^4} = 3,25.$$

Un décade correspond à la fréquence multipliée par 10 ; il y a donc 0,325 décade dans la zone de transition.

- L’atténuation dans la zone de transition est au minimum de :

$$A = -15 - (-3) = -12 \text{ dB}.$$

- La pente minimale vaut donc :

$$p = \frac{-12}{0,325} = -36,9 \text{ dB/décade}.$$

- un premier ordre a une pente insuffisante de -20 dB/décade, on utilise donc un deuxième ordre de pente -40 dB/décade.

Quelle est la pulsation caractéristique ? Plusieurs choix sont possibles. Sur le gabarit, sont tracés en gris les deux diagrammes asymptotiques limites du deuxième ordre qui répondent au cahier des charges. ω_0 est calculé avec l'expression asymptotique de la fonction de transfert et les coordonnées d'un point du gabarit :

$$\underline{H}(j\omega) = \frac{1}{1 + 2\xi \frac{j\omega}{\omega_0} + \left(\frac{j\omega}{\omega_0}\right)^2} \simeq \left(\frac{\omega_0}{j\omega}\right)^2 \quad \text{pour } \omega \gg \omega_0.$$

Si l'on choisit le filtre qui coupe le plus tôt, à savoir qui passe par $\omega = 1,2 \cdot 10^4 \text{ rad.s}^{-1}$ et -3 dB :

$$20 \log \left| \left(\frac{\omega_0}{j\omega} \right)^2 \right| = 40 \log \left(\frac{\omega_0}{\omega} \right) = -3 \quad \text{implique } \omega_0 = 1,0 \cdot 10^4 \text{ rad.s}^{-1}.$$

Dans l'autre cas limite qui coupe le plus tard et passe par $\omega = 3,9 \cdot 10^4 \text{ rad.s}^{-1}$ et -15 dB :

$$40 \log \left(\frac{\omega_0}{\omega} \right) = -15 \quad \text{implique } \omega_0 = 1,6 \cdot 10^4 \text{ rad.s}^{-1}.$$

La pulsation caractéristique sera choisie entre ces deux valeurs limites.

3.4 Cahier des charges

Quel type de filtre choisir et quel gabarit imposer ? La réponse dépend de l'information cherchée dans un ou plusieurs signaux. Si on souhaite :

- récupérer une fréquence particulière, par exemple pour choisir une station de radio parmi d'autres, on utilise un **passe-bande**, centré sur la fréquence utile,
- éliminer une fréquence, on utilise un **coupe-bande**, centré sur la fréquence à éliminer,
- le dériver, on utilise un **passe-haut** bien choisi,
- l'intégrer, on utilise un **passe-bas**.

Et, ce qui est étudié en détail dans le prochain chapitre, si on souhaite :

- éliminer la valeur moyenne du signal, on utilise un **passe-haut**,
- recueillir l'information sur la forme générale du signal, on utilise un **passe-bas**,
- recueillir l'information relative aux détails du signal, on utilise un **passe-haut**.
- moyenner un signal, on utilise un **passe-bas** très sélectif.

Exemple

L'émetteur qui assure la diffusion de France-Inter sur les grandes ondes, transmet deux signaux autour de $f_0 = 162 \text{ kHz}$, comme on le voit sur la figure 11.28 :

- un signal analogique qui contient des informations sonores. Il occupe une bande de fréquence comprise entre 50 Hz et 5 kHz de chaque côté de f_0 , donc les intervalles $[112 \text{ kHz} ; 161,5 \text{ kHz}]$ et $[162,5 \text{ kHz} ; 212 \text{ kHz}]$;
- un signal numérique qui porte des informations horaires (date et heure légale). C'est ce signal qui permet de synchroniser les horloges des gares ou les dispositifs radio-

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D’UN SYSTÈME LINÉAIRE

commandés des particuliers. Il occupe une bande de fréquence de largeur 80 Hz, centrée en f_0 , soit l’intervalle [161,6 kHz ; 162,4 kHz].

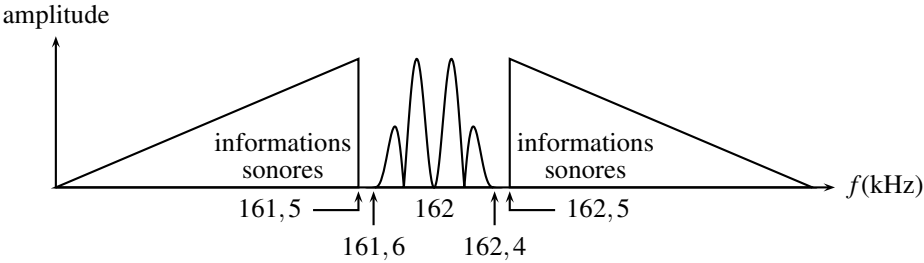


Figure 11.28 – Spectre du signal France-Inter (échelle horizontale non respectée pour des raisons de clarté)

Comment récupérer le signal horaire ? On utilise un passe-bande, qui élimine le signal analogique avec $f_a^- = 161,5$ kHz et $f_a^+ = 162,5$ kHz, mais laisse passer le signal numérique avec $f_p^- = 161,6$ kHz et $f_p^+ = 162,4$ kHz. On choisit $G_{sup} = -3$ dB afin que les signaux utiles ne soient que peu atténués et $G_{inf} = -40$ dB afin de diviser le signal analogique par 10^2 pour le filtrer. Les bandes de transition de ce filtre n’ont qu’une largeur de 0,1 kHz, ce qui est faible par rapport à $f_0 = 162$ kHz. Dans la pratique, un prétraitement du signal sera effectué avant son filtrage.

D’une manière générale, plus le filtre tend vers un filtre parfait (pas d’atténuation dans la bande passante, atténuation infinie dans la bande rejetée, passage net de la bande passante à la bande atténuée, c’est à dire bande de transition étroite, f_a proche de f_p par exemple), plus le filtre sera compliqué, construit avec de nombreux éléments.

SYNTHÈSE

SAVOIRS

- fonction de transfert
- diagramme de Bode
- notion de gabarit
- passe-bas du premier ordre
- passe-haut du premier ordre
- passe-bas du deuxième ordre
- passe-haut du deuxième ordre

SAVOIR-FAIRE

- utiliser les échelles logarithmiques
- interpréter les zones rectilignes des diagrammes de Bode d’après la fonction de transfert
- établir un gabarit en fonction du cahier des charges (PTSI)
- expliciter les conditions d’utilisation d’un filtre en dérivateur
- expliciter les conditions d’utilisation d’un filtre en intégrateur

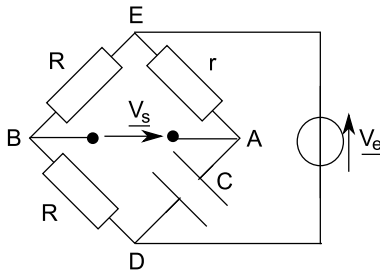
MOTS-CLÉS

- | | | |
|----------------------------|--------------|---------------|
| • fonction de transfert ou | • gabarit | • dérivateur |
| transmittance | • passe-bas | • intégrateur |
| • diagramme de Bode | • passe-haut | |

S’ENTRAÎNER

11.1 Fonction de transfert d’un pont, d’après ENAC 2005 (★)

On considère le circuit représenté sur le schéma de la figure suivante. Un pont dont les quatre branches sont constituées par trois résistors et un condensateur est alimenté par une source de tension sinusoïdale $v_e(t) = v_C - v_E = V_{E0} \cos(\omega t)$, de pulsation ω , connectée aux bornes de la diagonale ED . On désigne par $v_s(t) = v_A - v_B = V_{s0} \cos(\omega t + \phi_1)$ la tension de sortie recueillie aux bornes de la diagonale AB .



On définit la fonction de transfert $\underline{T}(j\omega)$ du circuit par le rapport de l’amplitude complexe $\underline{V}_s$ associée à la tension de sortie sur l’amplitude complexe $\underline{V}_e$, associée à la tension d’entrée.

1. Exprimer $\underline{T}(j\omega) = \underline{V}_s / \underline{V}_e$.
2. Calculer le module de $\underline{T}$. Simplifier $\underline{T}(j\omega)$ pour préciser comment varie sa phase ϕ_1 .
3. On donne $\omega = 1000 \text{ rad.s}^{-1}$, $C = 1 \text{ } \mu\text{F}$. Quelle valeur r_0 doit-on donner à r pour que $\phi_1 = -\pi/2$?

11.2 Circuit en régime sinusoïdal (★)

On considère un circuit composé d’une résistance R et d’un condensateur de capacité C en série aux bornes duquel on place un générateur basses fréquences délivrant une tension sinusoïdale d’amplitude E et de pulsation ω . On considère alors le filtre dont l’entrée est la tension du générateur basses fréquences et la sortie la tension aux bornes du condensateur.

1. Comment se comporte le condensateur à hautes et basses fréquences ? En déduire la nature du filtre.
2. Établir l’expression de la fonction de transfert du filtre et la mettre sous forme canonique.
3. Étudier les variations du gain avec la pulsation.
4. Même question pour le déphasage.
5. Donner à basses fréquences les valeurs limites du gain, du gain en décibels et du déphasage.
6. Même question à hautes fréquences.

APPROFONDIR

11.3 Étude expérimentale d'un circuit RC (★)

Soit un circuit comprenant en série un condensateur de capacité C variable , une bobine d'inductance L et de résistance R , un générateur idéal de tension basse fréquence d'amplitude E et de fréquence f fixe et un ampèremètre de résistance négligeable (sauf indication contraire).

1. Représenter le schéma du circuit.
2. Déterminer l'expression de l'intensité parcourant le circuit.
3. Établir qu'en faisant varier la valeur de la capacité, l'intensité passe par un maximum I_0 . On notera C_0 la valeur de la capacité correspondante.
4. On détermine les valeurs C_1 et C_2 de la capacité correspondant à une intensité $\frac{I_0}{\sqrt{2}}$. Expliciter les expressions de C_1 et C_2 en fonction de L , R et ω la pulsation.
5. On fait l'hypothèse que C_1 et C_2 sont peu différentes de C_0 . Expliciter l'approximation que cela permet de faire.
6. Montrer que sous cette hypothèse les valeurs C_1 et C_2 sont équidistantes de C_0 .
7. Que peut-on dire du déphasage pour $C = C_0$?
8. Même question pour C_1 et C_2 . Comment peut-on qualifier ces valeurs ?
9. Quelle est la meilleure façon de déterminer expérimentalement C_0 ?
10. Pourquoi est-il plus avantageux ici de mesurer C_1 et C_2 que C_0 ?
11. Exprimer le facteur de qualité Q en fonction de C_1 et C_2 .
12. Dédire de ce qui précède les expressions de L et R en fonction de C_1 , C_2 et de la fréquence f .
13. Que se passe-t-il si on ne néglige plus la résistance R_A de l'ampèremètre ? Préciser.
14. Application numérique : On donne $C_1 = 120 \text{ nF}$, $C_2 = 130 \text{ nF}$ et $f = 100 \text{ kHz}$. Déterminer le facteur de surtension et les valeurs de R et L en négligeant R_A puis lorsque $R_A = 0,1 \Omega$.
15. Les approximations faites sont-elles justifiées ?

CORRIGÉS

11.1 Fonction de transfert d’un pont, d’après ENAC 2005

1. La tension $\underline{V}_s$ est égale à $\underline{V}_A - \underline{V}_B$. Si l’on note $\underline{v}_1$ la tension $\underline{V}_A - \underline{V}_E$ et $\underline{v}_2$ la tension $\underline{V}_E - \underline{V}_B$, on peut écrire : $\underline{V}_s = \underline{v}_1 + \underline{v}_2$.

Comme il n’y pas de dipôle entre A et B, Les deux résistances R sont en série dans la branche EBD , et r et C sont en série dans la branche EAD . On peut utiliser la relation du diviseur de tension entre $\underline{v}_1$ et $\underline{V}_e$, et entre $\underline{v}_2$ et $\underline{V}_e$, soit : $\underline{v}_1 = -\frac{r}{r + \frac{1}{jC\omega}} \underline{V}_e = -\frac{jrc\omega}{1 + jrc\omega} \underline{V}_e$ et

$$\underline{v}_2 = \frac{R}{2R} \underline{V}_e = \frac{1}{2} \underline{V}_e.$$

On en déduit : $\underline{V}_s = \underline{v}_1 + \underline{v}_2 = \left(-\frac{jrc\omega}{1 + jrc\omega} + \frac{1}{2} \right) \underline{V}_e$ soit $\underline{T}(j\omega) = \frac{\underline{V}_s}{\underline{V}_e} = \frac{1}{2} \frac{1 - jrc\omega}{1 + jrc\omega}$.

Si l’on pose $T_0 = 1/2$ et $\omega_1 = 1/rC$, on peut écrire $\underline{T}$ sous la forme : $\underline{T} = T_0 \frac{1 - j\frac{\omega}{\omega_1}}{1 + j\frac{\omega}{\omega_1}}$.

2. Le module de $\underline{T}$ est égal à : $T(\omega) = |\underline{T}| = T_0 \frac{\sqrt{1 + \left(\frac{\omega}{\omega_1}\right)^2}}{\sqrt{1 + \left(\frac{\omega}{\omega_1}\right)^2}} = T_0$.

Une unique pulsation caractéristique ω_1 apparaît tant au numérateur qu’au dénominateur. Dressons un tableau asymptotique pour simplifier $\underline{T}$:

	ω_1	
$1 - j\frac{\omega}{\omega_1}$	1	$-j\frac{\omega}{\omega_1}$
$1 + j\frac{\omega}{\omega_1}$	1	$j\frac{\omega}{\omega_1}$
$H_{asympt}(j\omega)$	T_0	$-T_0$
pente	nulle	nulle
phase	0	$-\pi$

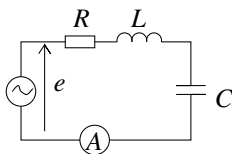
On retrouve que le gain est constant. La phase varie entre 0, pour $\omega \ll \omega_1$, et $-\pi$, pour $\omega \gg \omega_1$.

3. Pour avoir $\phi = -\pi/2$, il faut $r_0C\omega = 1/\sqrt{2}$ soit $r_0 = 707 \, \Omega$.

11.2 Circuit en régime sinusoïdal

- 1. À basses fréquences, la capacité C se comporte comme un interrupteur ouvert donc $s = e$. A hautes fréquences, le comportement est celui d'un fil donc $s = 0$. On en déduit qu'il s'agit d'un filtre passe-bas.
- 2. Par la relation du pont diviseur de tension, on en déduit la fonction de transfert $\underline{H} = \frac{1}{1 + jRC\omega} = \frac{H_0}{1 + j\frac{\omega}{\omega_0}}$ avec $H_0 = 1$ et $\omega_0 = \frac{1}{RC}$.
- 3. On obtient le gain $G = \frac{1}{\sqrt{1 + R^2C^2\omega^2}}$ qui est une fonction décroissante.
- 4. Le déphasage s'exprime par $\varphi = -\text{Arctan}RC\omega$ qui est une fonction décroissante.
- 5. Pour $\omega \ll \omega_0$, le gain G tend vers 1, G_{dB} vers 0 et φ vers 0.
- 6. Pour $\omega \gg \omega_0$, le gain G tend vers 0, G_{dB} vers $-\infty$ et φ vers $-\frac{\pi}{2}$.

11.3 Étude expérimentale d'un circuit RC



- 1.
- 2. On utilise la notation complexe et : $\underline{i} = \frac{\underline{e}}{\underline{Z}} = \frac{\underline{e}}{R + j(L\omega - \frac{1}{C\omega})}$. On en déduit le module : $I = \frac{E}{\sqrt{R^2 + (L\omega - \frac{1}{C\omega})^2}}$ et le déphasage φ tel que $\tan \varphi = \frac{-L\omega + \frac{1}{C\omega}}{R}$.
- 3. On considère ici I comme une fonction de C : I est maximale si $f(C) = R^2 + \left(L\omega - \frac{1}{C\omega}\right)^2$ est minimale. Or $f(C) > R^2$ pour toute valeur de C et $f(C) = R^2$ pour $C = C_0 = \frac{1}{L\omega^2}$. La fonction $I(C)$ présente bien un maximum.
- 4. Pour obtenir les valeurs de C_1 et de C_2 , il faut résoudre : $I = \frac{E}{\sqrt{R^2 + (L\omega - \frac{1}{C\omega})^2}} = \frac{I_0}{\sqrt{2}}$ avec $I_0 = \frac{E}{R}$. Cela conduit à l'équation : $R^2 + (L\omega - \frac{1}{C\omega})^2 = 2R^2$ soit $L\omega - \frac{1}{C\omega} = \pm R$. On en déduit : $C_1 = \frac{1}{L\omega^2 + R\omega}$ et $C_2 = \frac{1}{L\omega^2 - R\omega}$.
- 5. C_1 et C_2 sont supposés proches, ce qui sera possible si $L\omega^2 \ll R\omega$ soit $L\omega \ll R$.
- 6. Dans le cadre de cette approximation : $C_1 \approx \frac{1}{L\omega^2} \left(1 - \frac{R}{L\omega}\right)$ et $C_2 \approx \frac{1}{L\omega^2} \left(1 + \frac{R}{L\omega}\right)$ donc : $C_0 - C_1 = C_2 - C_0 = \frac{R}{L^2\omega^3}$.

CHAPITRE 11 – ANALYSE FRÉQUENTIELLE D'UN SYSTÈME LINÉAIRE

7. Lorsque $C = C_0$ alors $\tan \varphi = 0$, on en déduit : $\varphi = 0$: il n'y a pas de déphasage.
8. Comme $\frac{1}{C_1 \omega} = L\omega + R$ et $\frac{1}{C_2 \omega} = L\omega - R$, on en déduit : $\tan \varphi_1 = 1$ et $\tan \varphi_2 = -1$ soit $\varphi_1 = \frac{\pi}{4}$ et $\varphi_2 = -\frac{\pi}{4}$ puisque $\cos \varphi > 0$. On a donc des valeurs quadrantales.
9. Pour obtenir C_0 , on cherche la valeur maximale de I à l'ampèremètre ou à l'oscilloscope.
10. La mesure de C_0 est délicate car il faut chercher un maximum, ce qui conduit à des incertitudes importantes. D'autre part, pour déterminer R et L , il faut deux mesures, ce que donne la détermination de C_1 et de C_2 .
11. Comme $Q = \frac{L\omega}{R}$ et que $C_2 - C_1 = \frac{2RC_0}{L\omega} = \frac{2C_0}{Q}$, on en déduit : $Q = \frac{C_2 + C_1}{C_2 - C_1}$ puisque $C_0 = \frac{C_1 + C_2}{2}$.
12. De $C_0 = \frac{C_1 + C_2}{2} = \frac{1}{L\omega^2}$, on tire : $L = \frac{1}{2\pi^2(C_1 + C_2)f^2}$. Alors en utilisant l'expression de Q , on en déduit : $R = \frac{C_2 - C_1}{\pi f(C_1 + C_2)^2}$.
13. Si on tient compte de la résistance de l'ampèremètre, R devient $R + R_A$. Par conséquent, la valeur de L ne change pas et $R = \frac{C_2 - C_1}{\pi f(C_1 + C_2)^2} - R_A$.
14. On obtient : $Q = \frac{C_1 + C_2}{C_1 - C_2} = 25$, $R = \frac{C_2 - C_1}{\pi f(C_1 + C_2)^2} = 0,51 \, \Omega$ en négligeant la résistance de l'ampèremètre, $R = \frac{C_2 - C_1}{\pi f(C_1 + C_2)^2} - R_A = 0,41 \, \Omega$ en en tenant compte et $L = \frac{1}{2\pi^2(C_1 + C_2)f^2} = 20 \, \mu\text{H}$.
15. L'approximation $R = 0,5 \, \Omega \ll L\omega = 12,6 \, \Omega$ est justifiée. En revanche, R est du même ordre de grandeur que R_A et l'approximation $R_A \ll R$ ne peut être faite.

Filtrage linéaire

12

Tout signal physique est décomposable en une somme de signaux sinusoïdaux : c'est le point de départ de l'analyse de Fourier qui justifie l'importance de l'étude des réponses en régime permanent sinusoïdal. En utilisant les méthodes introduites dans les précédents chapitres, en découlent la notion de filtrage d'un signal par un système, ainsi que celle de l'adaptation d'un système en vue de satisfaire à un cahier des charges.

1 Réponse d'un système linéaire en régime permanent

1.1 Légitimité de l'étude harmonique

L'étude de la réponse d'un système à un signal sinusoïdal est récurrente. Or tous les signaux ne sont pas sinusoïdaux. Est-ce alors une perte de généralité ? Non, car les signaux sinusoïdaux sont les briques élémentaires avec lesquelles on fabrique n'importe quel autre signal, périodique ou non (les signaux apériodiques peuvent être étudiés par des sommes continues de sinusoïdes, nommées intégrales de Fourier). On étudie donc la réponse d'un système à une de ces briques élémentaires pour ensuite sommer les réponses.

1.2 Cadre de l'étude

Un système est un dispositif qui traite un ou plusieurs signaux d'entrée afin de réaliser la fonction voulue par son concepteur ; le système produit ainsi à son tour un ou plusieurs signaux de sortie. Par exemple, un téléphone portable est une suite de systèmes qui accomplissent successivement la transformation d'un signal sonore, la voix, en un signal électrique ; la numérisation du signal électrique en un autre signal électrique ; le codage du signal numérisé en un signal, toujours électrique, constitué de niveaux logiques ; finalement son émission, c'est-à-dire la transformation du signal électrique codé en un signal électromagnétique.

Les systèmes étudiés sont des **systèmes linéaires** : la réponse à une combinaison linéaire d'entrées est la combinaison linéaire des sorties associées à chaque entrée. Par exemple, si une personne *A* parle dans un téléphone, le récepteur entend *A*. Il entend de même *B* si *B* parle dans le même téléphone. Et si *A* et *B* sont côte à côte et parlent tous les deux en même temps, alors le récepteur entend les voix de *A* et de *B*. Ceci se formalise de la manière suivante :

CHAPITRE 12 – FILTRAGE LINÉAIRE

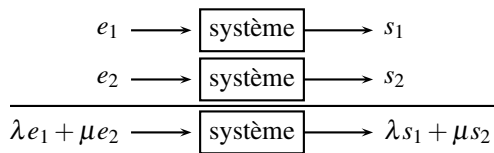


Figure 12.1 – Linéarité d'un système.

Les systèmes étudiés sont aussi **invariants dans le temps** : leurs propriétés ne changent pas au cours du temps. Ce n'est le cas qu'approximativement, car tous les systèmes vieillissent et doivent être remplacés lorsque leurs propriétés se sont dégradées. Ils restent toutefois invariants pendant une longue durée, leur durée de vie.

Les systèmes linéaires et invariants dans le temps sont souvent indiqués par leur acronyme anglais **LTI systems**, pour *linear time-invariant systems*.

1.3 Entrée sinusoïdale

On considère un système *RC* du premier ordre, alimenté par une entrée sinusoïdale $e(t) = E_0 \cos(\omega t)$, dont le déphasage est pris nul par un choix convenable de l'origine des temps. Le signal complexe associé à $e(t)$ est $\underline{e}(t)$:

$$\underline{e}(t) = E_0 \exp(j\omega t) \quad \text{ainsi} \quad e(t) = \text{Re}(\underline{e}(t)).$$

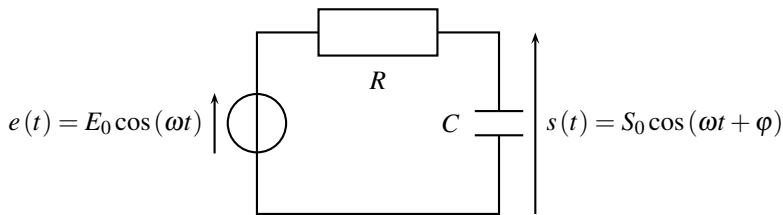


Figure 12.2 – Système du premier ordre étudié.

La fonction de transfert du système *RC*, en régime permanent sinusoïdal, a été vue dans les deux chapitres précédents :

$$\underline{H}(j\omega) = \frac{1}{1 + jRC\omega} = \frac{1}{1 + j\tau\omega},$$

avec $\tau = RC$. Quelle est l'expression de la sortie $s(t)$ du système en régime permanent sinusoïdal ? On cherche d'abord le signal complexe $\underline{s}(t)$:

$$\underline{s}(t) = \underline{H}(j\omega) \times \underline{e}(t) = \frac{1}{1 + j\tau\omega} E_0 \exp(j\omega t).$$

ω identiques



Le signal d'entrée oscille à la pulsation ω . Le système réagit linéaire donc à cette même pulsation ω . C'est donc la même pulsation ω qui apparaît dans $\exp(j\omega t)$, qui décrit $e(t)$ et dans $\underline{H}(j\omega)$, qui décrit le système.

Afin d'établir l'amplitude de $s(t)$ et son déphasage par rapport à $e(t)$, il faut écrire $\underline{s}(t)$ sous forme module $\times \exp(j \arg)$. De même pour la transmittance :

$$\begin{aligned}\frac{1}{1+j\tau\omega} &= \left| \frac{1}{1+j\tau\omega} \right| \exp(j(\arg(1) - \arg(1+j\tau\omega))) \\ &= \frac{1}{\sqrt{1+(\tau\omega)^2}} \exp(-j \arctan(\tau\omega)).\end{aligned}$$

Dès lors :

$$\underline{s}(t) = \frac{E_0}{\sqrt{1+(\tau\omega)^2}} \exp(j(\omega t - \arctan(\tau\omega))).$$

En prenant la partie réelle de $\underline{s}(t)$:

$$s(t) = \frac{E_0}{\sqrt{1+(\tau\omega)^2}} \cos(\omega t - \arctan(\tau\omega)).$$

Pour un signal d'entrée sinusoïdal, le signal de sortie est sinusoïdal de même pulsation ω ; on utilise $\underline{H}(j\omega)$ pour le calculer.

1.4 Entrée combinaison linéaire de fonctions sinusoïdales

Si l'entrée du système RC se compose maintenant de deux fonctions sinusoïdales de pulsations différentes, ω_1 et ω_2 :

$$e(t) = E_1 \cos(\omega_1 t) + E_2 \cos(\omega_2 t),$$

que vaut la sortie $s(t)$ en régime sinusoïdal permanent ?

La méthode est simple : l'entrée est la superposition de deux entrées, $e_1(t) = E_1 \cos(\omega_1 t)$ et $e_2(t) = E_2 \cos(\omega_2 t)$. $e_1(t)$ seul impose une sortie $s_1(t)$, $e_2(t)$ seul impose une sortie $s_2(t)$. Par linéarité, la sortie est alors la superposition des deux sorties :

$$s(t) = s_1(t) + s_2(t).$$

Que vaut $s_1(t)$? L'entrée $e_1(t)$ oscille à la pulsation ω_1 , le système réagit donc à cette même pulsation ω_1 et la sortie $s_1(t)$ oscille en régime sinusoïdal permanent à ω_1 . Le calcul est identique à celui du paragraphe précédent :

$$s_1(t) = \frac{E_1}{\sqrt{1+(\tau\omega_1)^2}} \cos(\omega_1 t - \arctan(\tau\omega_1)).$$

CHAPITRE 12 – FILTRAGE LINÉAIRE

De même pour $s_2(t)$, qui oscille à la pulsation ω_2 imposée par $e_2(t)$:

$$s_2(t) = \frac{E_2}{\sqrt{1 + (\tau\omega_2)^2}} \cos(\omega_2 t - \arctan(\tau\omega_2)) .$$

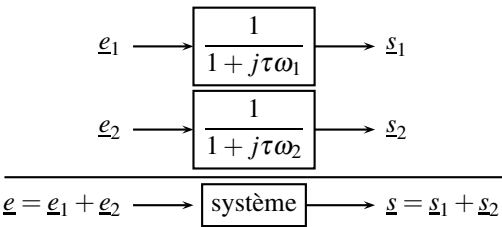


Figure 12.3 – Sortie associée à une entrée composée de deux sinusoïdes.

Finalement :

$$s(t) = \frac{E_1}{\sqrt{1 + (\tau\omega_1)^2}} \cos(\omega_1 t - \arctan(\tau\omega_1)) + \frac{E_2}{\sqrt{1 + (\tau\omega_2)^2}} \cos(\omega_2 t - \arctan(\tau\omega_2)) .$$

La sortie associée à la somme de deux signaux sinusoïdaux de pulsations ω_1 et ω_2 , est la somme des deux signaux de mêmes pulsations ω_1 et ω_2 et se calcule à l’aide de $\underline{H}(j\omega_1)$ et $\underline{H}(j\omega_2)$.

Remarque

$|\underline{H}(j\omega_1)|$ et $\arg(\underline{H}(j\omega_1))$ se lisent sur les deux courbes du diagramme de Bode.

2 Contenu spectral d’un signal

2.1 Décomposition de Fourier

La notion de spectre d’un signal périodique et de **série de Fourier**¹ est introduite dans le chapitre sur les ondes.

On examine ici le cas, très fréquent en électronique, d’un signal créneau $e(t)$, de 2 V d’amplitude, de période $T = 2$ s, de valeur moyenne $\langle e \rangle = 1$ V (aussi connu sous le terme anglais d’*offset*, directement réglable sur un GBF) :

1. Joseph Fourier, 1768 – 1830, élève à l’École Normale Supérieure, professeur à l’École Polytechnique, participe à la campagne d’Égypte de Napoléon Bonaparte puis fut préfet de l’Isère, où, en parallèle de son travail administratif, il étudia la diffusion thermique pour laquelle il introduisit ces séries trigonométriques.

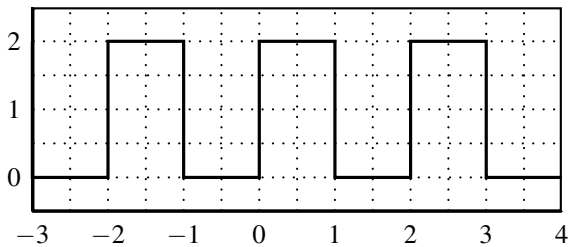


Figure 12.4 – Forme d’onde du signal créneau étudié.

Ce signal admet un développement en série de Fourier :

$$e(t) = 1 + \frac{4}{\pi} \sin\left(2\pi \frac{t}{T}\right) + \frac{4}{3\pi} \sin\left(3 \times 2\pi \frac{t}{T}\right) + \frac{4}{5\pi} \sin\left(5 \times 2\pi \frac{t}{T}\right) + \dots$$

dans lequel on reconnaît :

- la valeur moyenne $\langle e \rangle = 1 \text{ V}$,
- le fondamental, de même période T que le signal $e(t)$: $\frac{4}{\pi} \sin\left(2\pi \frac{t}{T}\right)$,
- l’harmonique de rang 3, de période $\frac{T}{3}$: $\frac{4}{3\pi} \sin\left(3 \times 2\pi \frac{t}{T}\right)$,
- l’harmonique de rang n impair, de période $\frac{T}{n}$: $\frac{4}{n\pi} \sin\left(n \times 2\pi \frac{t}{T}\right)$.

Le signal $e(t)$ ne contient aucune harmonique de rang pair, on peut écrire son développement en série de Fourier :

$$e(t) = 1 + \sum_{n=1}^{\infty} 2 \frac{1 + (-1)^{n+1}}{n\pi} \sin\left(n \times 2\pi \frac{t}{T}\right).$$

Le spectre de $e(t)$ a l’allure suivante :

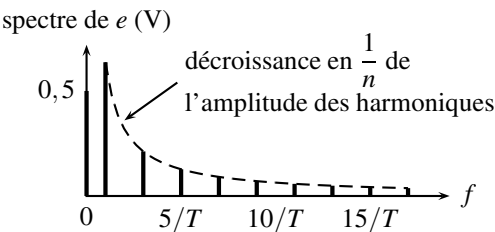


Figure 12.5 – Spectre d’un signal créneau de valeur moyenne non nulle.

Examinons l’influence de l’ajout successif des harmoniques sur le signal en dents de scie, en traçant, sur une demi-période, la somme partielle de Fourier :

$$\langle e \rangle + \sum_{n=1}^k E_n \cos\left(2\pi n \frac{t}{T} + \varphi_n\right) = 1 + \sum_{n=1}^k 2 \frac{1 + (-1)^{n+1}}{n\pi} \sin\left(n \times 2\pi \frac{t}{T}\right).$$

CHAPITRE 12 – FILTRAGE LINÉAIRE

On place tout d’abord le terme de pulsation nulle, la valeur moyenne, constante ; puis on lui ajoute le fondamental et l’harmonique de rang trois (on rappelle que l’harmonique de rang deux est nulle). Le signal original créneau est tracé en gris :

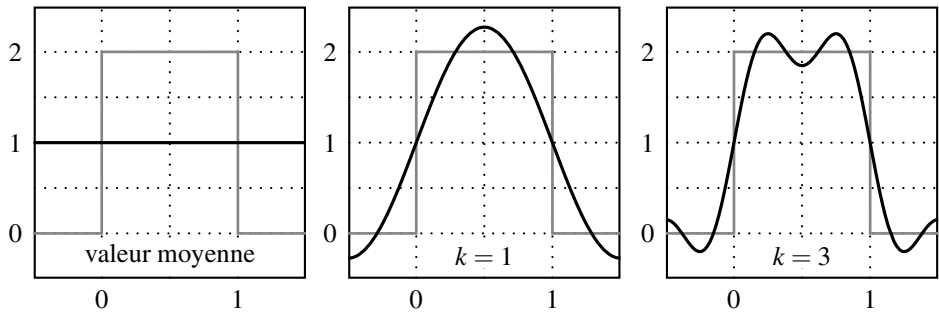


Figure 12.6 – Sommes partielles de Fourier.

On continue de sommer en ajoutant successivement les harmoniques de rang 5, 7 et 9.

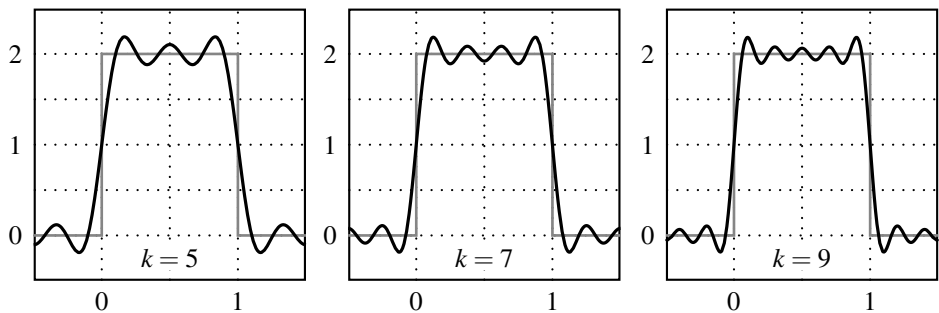


Figure 12.7 – Sommes partielles de Fourier.

En allant plus loin dans la somme de Fourier, avec l’ajout des harmoniques 11, 13 et 15, on se rapproche du signal en créneau.

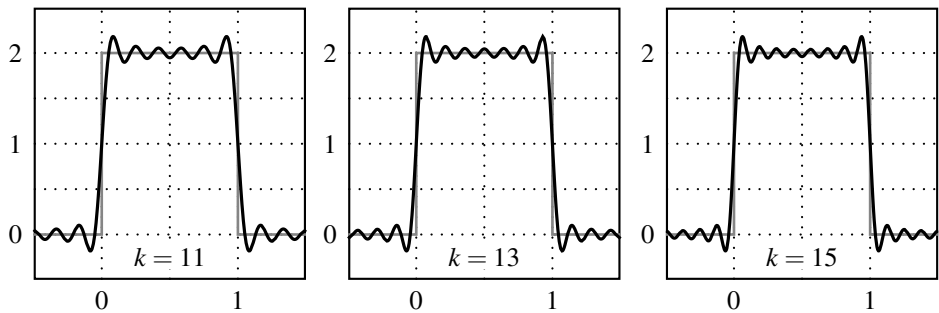


Figure 12.8 – Sommes partielles de Fourier.

En sommant pour aller jusqu'aux harmoniques 21 et 41, on voit que la somme de Fourier est très proche du signal en créneau. La somme des 400 premières harmoniques permet de reconstituer pratiquement le signal :

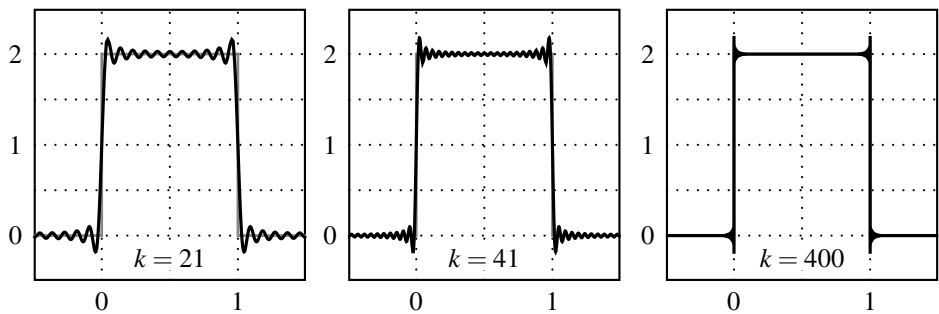


Figure 12.9 – Sommes partielles de Fourier.

On observe que si l'on se restreint aux premiers termes de la série de Fourier, on retrouve l'allure générale du signal : proche de zéro pour $t < 0$ et $t > 1$ s, proche de deux sur $[0, 1$ s], avec une variation rapide autour de 0 s et 1 s. Ces premiers termes sont suffisants pour décrire la forme générale du signal.

Les basses fréquences (valeur moyenne, fondamental et premières harmoniques) contiennent la forme générale du signal.

Est-ce à dire que les hautes fréquences ne représentent rien ? Absolument pas ! Pour converger finement vers le signal créneau, il a fallu ajouter des harmoniques de rang toujours plus élevé. Les hautes fréquences sont indispensables pour synthétiser les brusques variations du signal, comme en $t = 0$ s ou $t = 1$ s, ainsi que pour parfaire les petits détails.

Que les détails du signal soient contenus dans les hautes fréquences, cela se comprend aisément. Les hautes fréquences varient vite, avec une très faible période ; elles construisent la partie du signal la plus fine. Quelle que soit la précision demandée, elle sera atteinte en montant assez haut en fréquence, c'est-à-dire avec une période aussi petite que l'on souhaite. La nécessité de la présence des hautes fréquences dans un signal qui varie brusquement est visible quand on les isole. Plus on augmente le nombre d'harmoniques dans les hautes fréquences, plus la pente du signal reconstitué est importante en $t = 0$ s et $t = 1$ s. On peut tracer uniquement ces hautes fréquences, à savoir la somme :

$$\sum_K^\infty E_n \cos \left(n \times 2\pi \frac{t}{T} + \varphi_n \right),$$

dans trois cas, $K = 20$, $K = 50$ puis $K = 600$. On isole ainsi la partie haute fréquence du signal, contenue par les harmoniques de rang supérieur ou égal à K .

CHAPITRE 12 – FILTRAGE LINÉAIRE

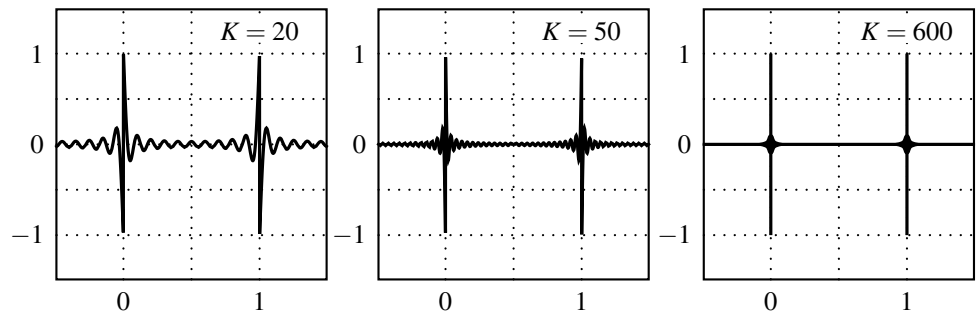


Figure 12.10 – Hautes fréquences contenues dans le signal triangle.

Les brusques variations en $t = 0$ s et $t = 1$ s se détachent très nettement lorsqu’on monte dans les hautes fréquences. Pour $K = 600$ par exemple, il ne reste plus qu’elles. Ce sont bien ces hautes fréquences qui contiennent cette discontinuité du signal.

Les hautes fréquences contiennent les brusques variations et les détails du signal.

2.2 Complément : phénomène de Gibbs

Les séries de Fourier tronquées présentent systématiquement un dépassement d’environ 9% lors d’une somme finie, nommé **phénomène de Gibbs**². Il est dû à l’absence dans la somme des hautes fréquences, indispensables pour modéliser correctement les détails fin. Par exemple pour le signal créneau déjà rencontré, on l’observe quand on somme les 400 premières harmoniques :

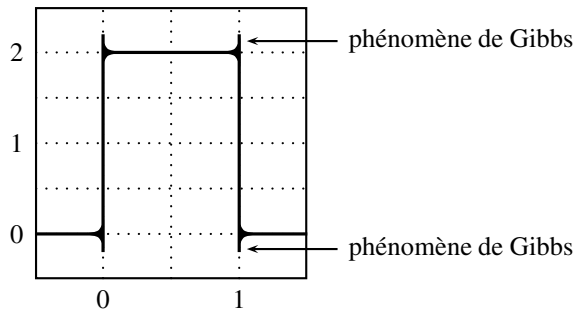


Figure 12.11 – Dépassement de Gibbs.

2. en l’honneur de Josiah Willard Gibbs, 1839 – 1903, physicien, chimiste et mathématicien américain, professeur à l’université de Yale, qui l’a le premier étudié. L’œuvre fondamentale en thermochimie de Gibbs relève du cours de seconde année.

3 Valeur efficace

3.1 Définition

La **valeur efficace** S_{eff} d'un signal $s(t)$, T -périodique, est définie par :

$$S_{eff}^2 = \langle s^2 \rangle = \frac{1}{T} \int_{t_0}^{t_0+T} s^2(t) dt,$$

où t_0 est une date quelconque, bien choisie afin de simplifier le calcul.

3.2 Cas d'un signal sinusoïdal

On étudie, par exemple, $s(t) = S_0 \cos(\omega t)$. Le déphasage de $s(t)$ est arbitrairement choisi nul. Ceci est toujours possible avec un décalage de l'origine des temps, c'est-à-dire le choix de la date $t = 0$. Alors :

$$S_{eff}^2 = \frac{1}{T} \int_0^T S_0^2 \cos^2(\omega t) dt.$$

Or $\cos^2(\omega_0 t) = \frac{1}{2} (1 + \cos(2\omega_0 t))$, donc :

$$S_{eff}^2 = \frac{S_0^2}{T} \int_0^T \left(\frac{1}{2} + \frac{\cos(2\omega t)}{2} \right) dt = \frac{S_0^2}{T} \left(\frac{1}{2} T + \left[\frac{\sin(2\omega t)}{4\omega} \right]_0^T \right).$$

Or $\sin(0) = 0$ et $\sin(2\omega T) = \sin(4\pi) = 0$. Ainsi : $S_{eff}^2 = \frac{S_0^2}{2}$, d'où :

$$S_{eff} = \frac{S_0}{\sqrt{2}}.$$

La valeur efficace S_{eff} d'un signal harmonique $s(t) = S_0 \cos(\omega t)$ est $S_{eff} = \frac{S_0}{\sqrt{2}}$. Alors :

$$s(t) = S_{eff} \sqrt{2} \cos(\omega t).$$

Les oscilloscopes numériques et les multimètres mesurent directement la valeur efficace et l'affichent avec la mention **RMS**, acronyme anglais signifiant *root mean square*. En effet, la valeur efficace est la racine carrée (*square root*) de la moyenne du carré du signal (*mean of the square*).

3.3 Cas d'un signal créneau

Si le signal $s(t)$ est un créneau T -périodique :

CHAPITRE 12 – FILTRAGE LINÉAIRE

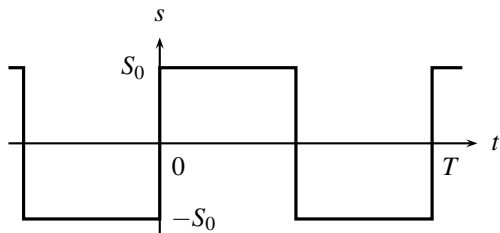


Figure 12.12 – Signal crêteau périodique.

Alors :

$$S_{eff}^2 = \frac{1}{T} \int_0^T S^2(t) dt = \frac{1}{T} \left(\int_0^{T/2} S_0^2 dt + \int_{T/2}^T (-S_0)^2 dt \right).$$

Ainsi :

$$S_{eff}^2 = S_0^2 \quad \text{d'où} \quad S_{eff} = S_0.$$



Le facteur $\sqrt{2}$, entre l'amplitude S_0 et la valeur efficace S_{eff} , est caractéristique du régime purement sinusoïdal, ne se généralise pas pour les autres formes d'ondes.

3.4 Comparaison avec un signal constant

Pourquoi définir une valeur efficace par la valeur moyenne du carré du signal, et pourquoi la qualifier d'efficace ?

Soit un circuit très simple, composé d'un générateur qui délivre une tension $e(t)$ et d'une résistance R qui modélise un chauffage ou un éclairage. L'intensité $i(t)$ du courant qui circule est donnée par la loi d'Ohm :

$$i(t) = \frac{e(t)}{R}.$$

Que vaut la puissance $\mathcal{P}$ délivrée en moyenne dans la résistance (puissance thermique dans le cas d'un chauffage, puissance lumineuse dans le cas d'un éclairage) ? Si $e(t)$ est une constante E_0 , alors l'intensité du courant est elle-aussi une constante $I_0 = E_0/R$; la puissance vaut dans ce cas :

$$\mathcal{P} = E_0 I_0 = \frac{E_0^2}{R}.$$

Si $e(t)$ est sinusoïdale, ce qui est le cas de la tension délivrée par EDF, elle s'écrit $e(t) = E_0 \cos(\omega t)$ et $i(t) = \frac{E_0}{R} \cos(\omega t)$. La puissance instantanée dans la résistance vaut :

$$p(t) = e(t) i(t) = \frac{e(t)^2}{R} = \frac{E_0^2}{R} \cos^2(\omega t),$$

et en moyenne :

$$\mathcal{P} = \langle p(t) \rangle = \frac{\langle e(t)^2 \rangle}{R} = \frac{E_{eff}^2}{R} = \frac{E_0^2}{2R}.$$

La puissance a chuté de moitié par rapport au cas constant !
Pour retrouver la même puissance, il faut alimenter le circuit avec une tension dont la valeur efficace vaut E_0 , c'est-à-dire une tension $e(t) = E_0\sqrt{2}\cos(\omega t)$. Ceci est illustré sur la figure 12.13.

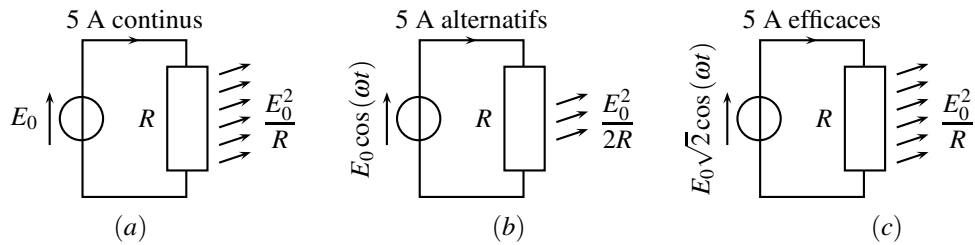


Figure 12.13 – Puissances dans un circuit résistif pour une intensité de 5 A (a) cas constant (b) cas sinusoïdal de même amplitude (c) cas sinusoïdal de même valeur efficace.

Un signal sinusoïdal de **valeur efficace** E_0 est aussi « efficace » qu'un signal constant E_0 pour l'alimentation en puissance. De plus, c'est bien $\langle e^2 \rangle$, c'est-à-dire E_{eff}^2 , qui intervient dans le calcul de la puissance moyenne.

4 Filtrage linéaire d'un signal non sinusoïdal

4.1 Position du problème

Le but de ce paragraphe est d'observer spectralement la réponse $s(t)$ d'un système, en régime permanent, lorsqu'il est soumis à un signal périodique non sinusoïdal, par exemple un signal créneau $e(t)$, de valeur moyenne différente de 0, de période T variable :

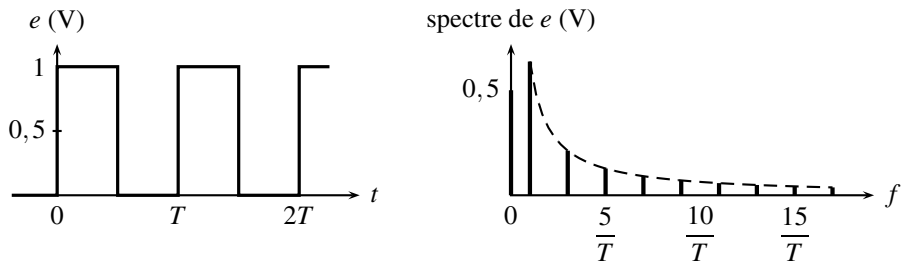


Figure 12.14 – Forme d'onde et spectre du signal $e(t)$.

Il sera important d'interpréter qualitativement la forme du signal de sortie avec l'éventuelle présence de la valeur moyenne et des hautes fréquences.

4.2 Filtrage passe-bas

On utilise un filtre passe-bas du premier ordre de transmittance : $\underline{H}(j\omega) = \frac{1}{1 + j\tau\omega}$.

CHAPITRE 12 – FILTRAGE LINÉAIRE

L'allure de la sortie dépend de la valeur de T par rapport à celle de τ , ce qui est très visible quand on trace, simultanément, le diagramme de Bode, limité au gain, et le spectre du signal d'entrée, accompagné de son enveloppe. La valeur moyenne de $e(t)$, de pulsation nulle, est rejetée à l'infini à gauche en échelle logarithmique. De plus, les hautes fréquences sont d'une amplitude trop faible pour être observées (l'amplitude des harmoniques décroît en $1/n$), mais restent bien présentes.

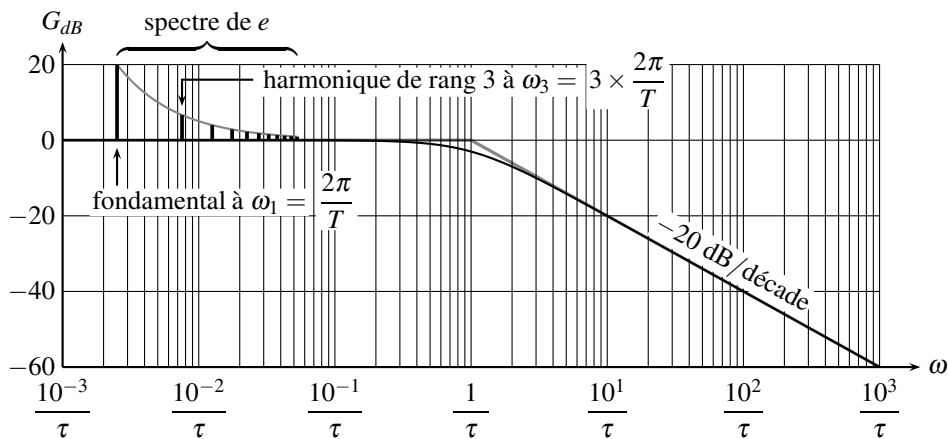


Figure 12.15 – Gain du système et spectre de l'entrée dans le cas où $\frac{2\pi}{T} \ll \frac{1}{\tau}$.

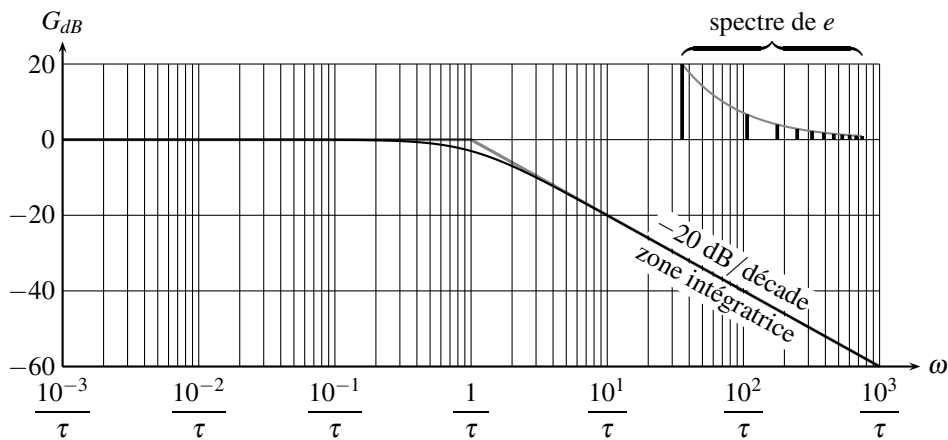


Figure 12.16 – Gain du système et spectre de l'entrée dans le cas où $\frac{2\pi}{T} \gg \frac{1}{\tau}$.

On observe trois cas distincts pour la sortie $s(t)$:

- $2\pi/T \ll 1/\tau$: le spectre est dans la zone passante du filtre. Le signal de sortie est peu différent de celui d'entrée, avec toutefois un arrondissement aux demi-périodes, car les hautes fréquences, indispensables pour reconstituer les brusques variations, sont filtrées.

Plus $1/T$ est faible devant $1/\tau$, plus les nombre d'harmoniques qui passent à travers le filtre est élevé, plus le signal de sortie tend vers le signal d'entrée.

- $2\pi/T \approx 1/\tau$: c'est une zone de transition, une partie seulement du signal d'entrée est filtrée.
- $2\pi/T \gg 1/\tau$: à part la valeur moyenne, qui passe intacte à travers le filtre, le spectre est totalement dans la zone à -20 dB/décade, où le signal est intégré. On observe donc un signal de sortie triangulaire, mais de très faible amplitude autour de la valeur moyenne car le système présente une forte atténuation. On a ainsi réalisé un **moyenueur**.

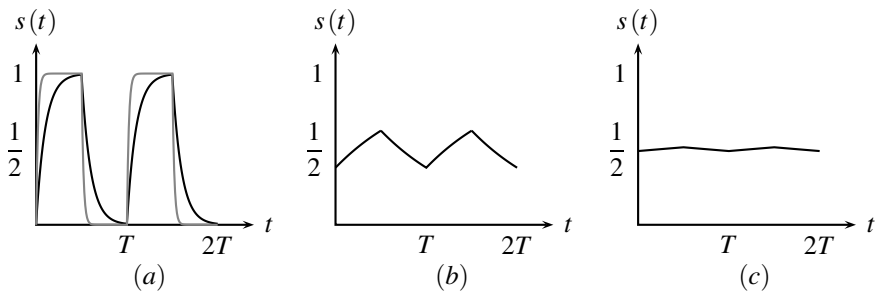


Figure 12.17 – Filtrage passe-bas d'un signal créneau (a) $T = 50\tau$ en gris et $T = 10\tau$ en noir, (b) $T = \tau$, (c) $T = 0, 1\tau$.

4.3 Réalisation d'un moyenueur

On voit nettement dans le cas (c) du paragraphe précédent, que le signal de sortie est quasiment constant, égal à sa valeur moyenne. En effet, seule la valeur moyenne du signal d'entrée passe à travers le filtre. On en déduit une méthode systématique de réalisation d'un moyenueur au moyen d'un passe-bas :

- identifier la période T du signal $e(t)$ dont on cherche la valeur moyenne,
- se souvenir que toutes les harmoniques contenues dans $e(t)$ auront une pulsation supérieure ou égale à $2\pi/T$, pulsation du fondamental, sauf la valeur moyenne qui a une pulsation nulle,
- construire un passe-bas qui laisse passer la pulsation nulle mais élimine toutes les pulsations à partir de $2\pi/T$, c'est-à-dire un filtre de pulsation caractéristique très inférieure à $2\pi/T$.

Un **moyenueur** est un filtre passe-bas de pulsation caractéristique très inférieure à celle du fondamental du signal à moyenner.

4.4 Filtrage passe-haut

On utilise maintenant un filtre passe-haut du premier ordre de transmittance :

$$H(j\omega) = \frac{j\tau\omega}{1 + j\tau\omega}.$$

CHAPITRE 12 – FILTRAGE LINÉAIRE

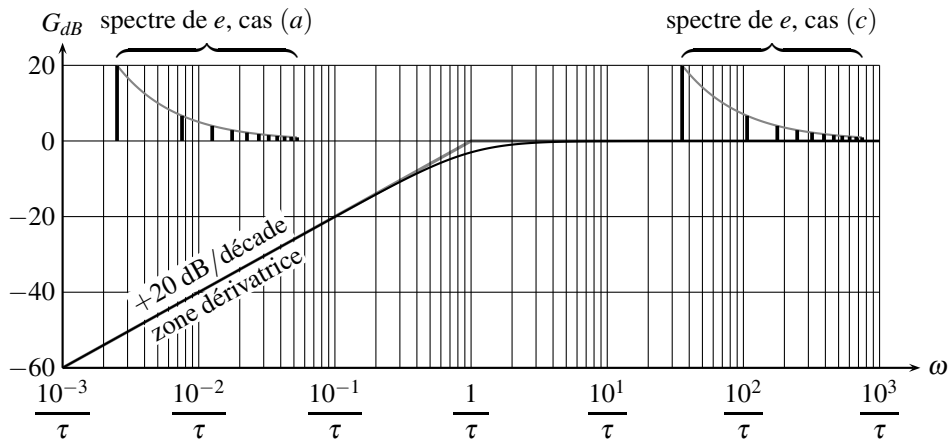


Figure 12.18 – Gain du système et spectre de l'entrée pour un filtrage passe-haut.

On observe dans chaque cas que la valeur moyenne est complètement éliminée par le filtre. La valeur moyenne du signal de sortie est donc nulle. De plus, l'allure de la sortie dépend encore de la valeur de T par rapport à celle de τ :

- $2\pi/T \ll \frac{1}{\tau}$: le spectre est dans la zone à $+20$ dB/décade, où le signal est dérivé. On observe donc un signal de sortie qui présente de brusques variations aux demi-périodes, dues aux hautes fréquences. L'amplitude du signal est due à ces hautes fréquences, d'amplitudes trop faibles pour être distinguées sur le spectre, qui passent à travers le filtre.
- $2\pi/T \approx \frac{1}{\tau}$: c'est une zone de transition, une partie seulement du signal d'entrée est filtrée.
- $2\pi/T \gg \frac{1}{\tau}$: le spectre est dans la zone passante du filtre. Le signal de sortie est peu différent de celui d'entrée, avec de brusques variations nettement marquées aux demi-périodes, à cause des hautes fréquences.

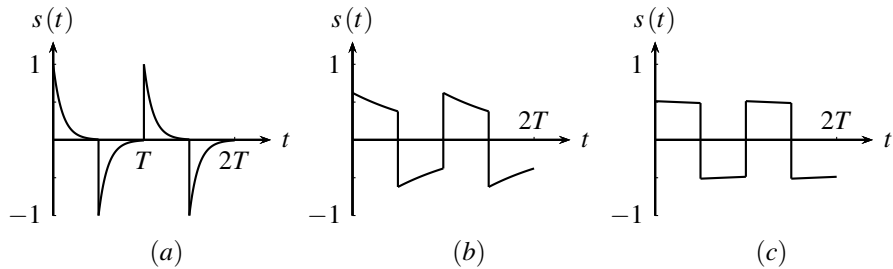


Figure 12.19 – Filtrage passe-haut d'un signal créneau (a) $T = 10\tau$, (b) $T = \tau$, (c) $T = 0,1\tau$.

4.5 Filtrage passe-bande

Ce filtrage est réalisé par un filtre du second ordre de transmittance :

$$\underline{H}(j\omega) = \frac{2\xi \frac{j\omega}{\omega_0}}{1 + 2\xi \frac{j\omega}{\omega_0} + \left(\frac{j\omega}{\omega_0}\right)^2}.$$

On traite ici le cas d'un filtre très sélectif, c'est-à-dire très peu amorti ou de fort coefficient de qualité :

$$Q = 50 \quad \text{soit} \quad \xi = \frac{1}{2m} = 10^{-2}.$$

Le filtre ne laisse passer qu'une étroite bande de pulsation autour de ω_0 , où son gain vaut 1. En dehors de cette fréquence, le gain est tellement négatif que le signal est complètement absorbé par le passe-bande. En effet, un gain -60 dB, par exemple, représente une division du signal par 1000.

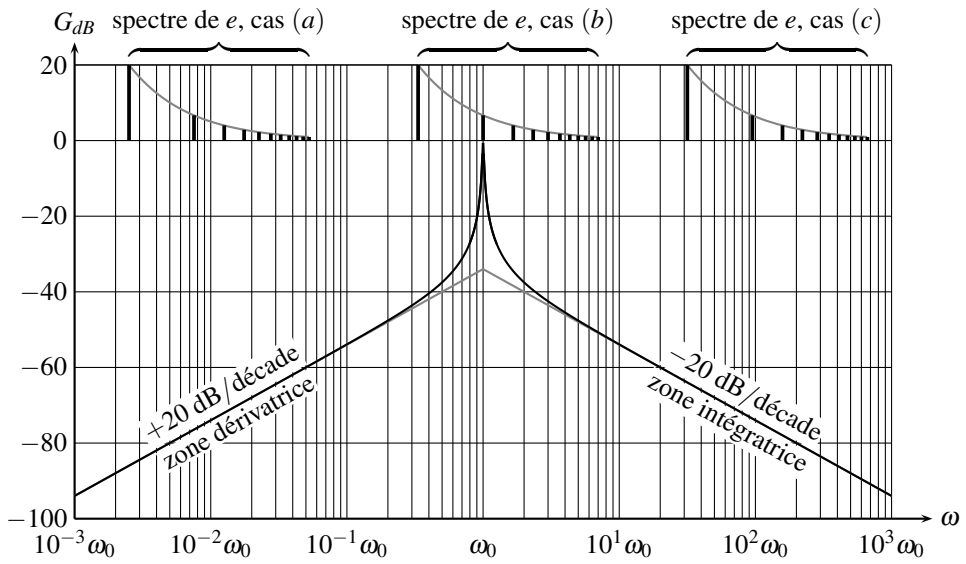


Figure 12.20 – Gain du système et spectre de l'entrée pour un filtrage passe-bande.

La valeur moyenne, composante de pulsation nulle, est complètement filtrée. La valeur moyenne du signal de sortie est donc nulle. De plus, l'allure de la sortie dépend encore de la valeur de $\frac{2\pi}{T}$ par rapport à celle de ω_0 :

- $\frac{2\pi}{T} \ll \omega_0$: le spectre est dans la zone à $+20$ dB/décade, où le signal est dérivé. Toutefois, l'atténuation est tellement forte que l'amplitude du signal est quasiment nulle.

CHAPITRE 12 – FILTRAGE LINÉAIRE

- $3 \times \frac{2\pi}{T} \approx \omega_0$: c'est l'exemple du bas (b) sur la figure. L'harmonique de rang 3, de pulsation $3 \times \frac{2\pi}{T}$, arrive exactement en ω_0 , seule composante que le filtre laisse passer. Cette harmonique est la seule à passer à travers le filtre, on la retrouve en sortie.
- $\frac{2\pi}{T} \gg \omega_0$: le spectre est dans la zone à -20 dB/décade, où le signal est intégré. Toutefois, l'atténuation est tellement forte que l'amplitude du signal triangulaire est très faible.

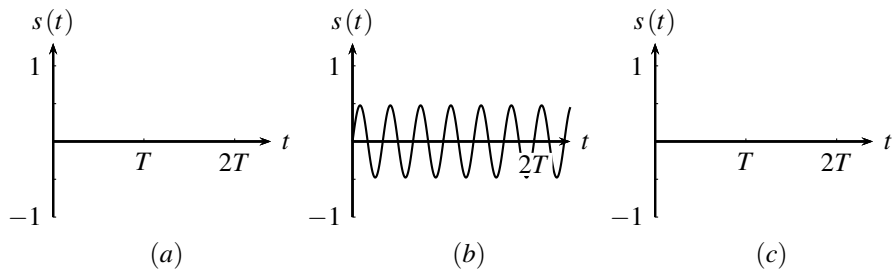


Figure 12.21 – Filtrage passe-bande d'un signal créneau : (a) $\frac{2\pi}{T} \ll \omega_0$,
(b) $3 \times \frac{2\pi}{T} = \omega_0$, (c) $\frac{2\pi}{T} \gg \omega_0$.

5 Réponse indicielle et contenu spectral

5.1 Possibilité de discontinuité, valeur moyenne

La réponse indicielle d'un système du premier ou du deuxième ordre est la réponse du filtre à un échelon. Elle a été étudiée dans les chapitres précédents. On rappelle, dans trois cas, les résultats alors obtenus, auxquels on ajoute la fonction de transfert du système.

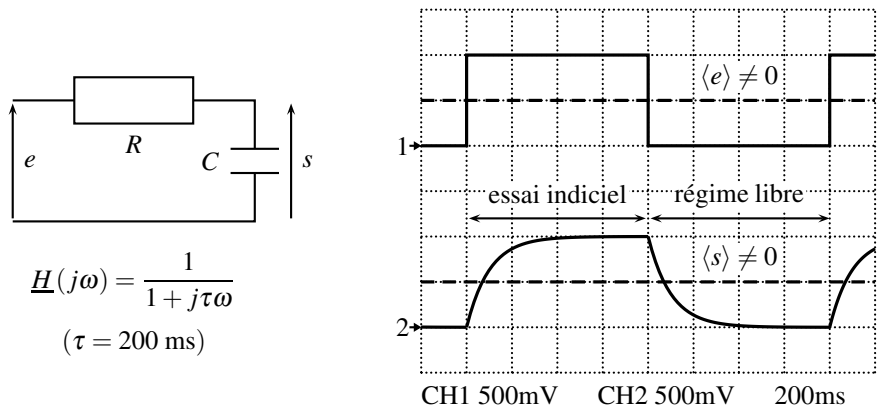


Figure 12.22 – Montage, transmittance, essai indiciel et régime libre d'un passe-bas du premier ordre.

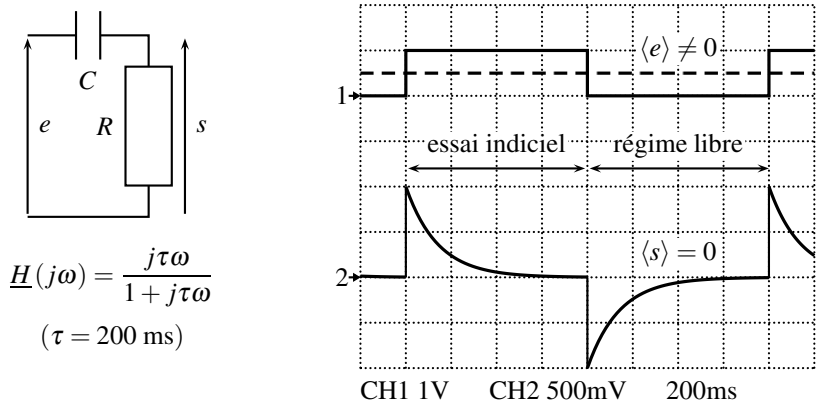


Figure 12.23 – Montage, transmittance, essai indiciel et régime libre d'un passe-haut du premier ordre.

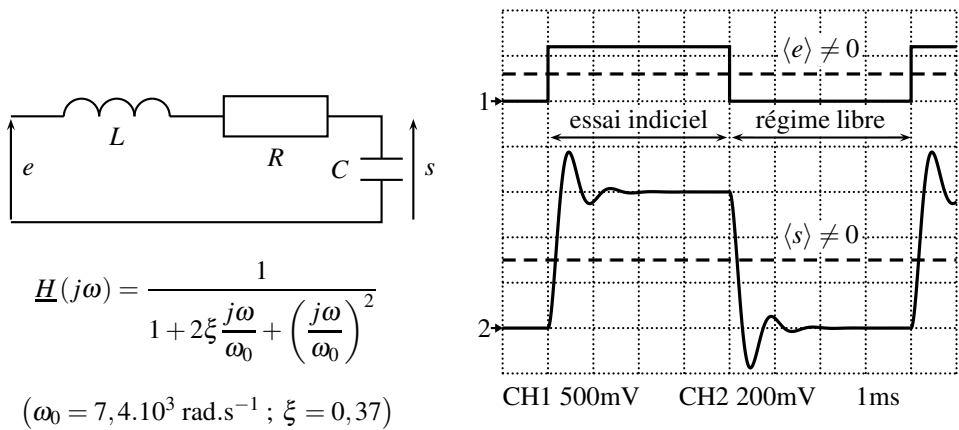


Figure 12.24 – Montage, transmittance, essai indiciel et régime libre d'un passe-bas du deuxième ordre.

On observe que pour les passe-bas, tant du premier que du deuxième ordre, le signal de sortie est continu, c'est-à-dire que la sortie admet la même valeur juste après l'échelon que juste avant. C'est la même chose pour le régime libre, la valeur de la sortie reste continue lors de la mise à la masse de l'entrée.

Pour le passe-haut, au contraire, on observe une sortie discontinue, c'est-à-dire que la valeur de la tension de sortie juste après l'échelon est différente de celle juste avant ; de même dans le cas du régime libre.

Comment interpréter ces résultats ?

La forme générale d'un signal est contenue dans ses basses fréquences, alors que les détails et les brusques variations le sont dans ses hautes fréquences. Ainsi, lors d'un filtrage passe-bas, les hautes fréquences sont supprimées ; il est donc impossible de reconstituer une brusque variation du signal en sortie. Lors du filtrage passe-haut, au contraire, les hautes fréquences,

CHAPITRE 12 – FILTRAGE LINÉAIRE

continues dans le signal créneau, sont transmises et expliquent la brusque variation du signal de sortie.

Il ne peut y avoir de discontinuité du signal de sortie que si le système présente un caractère passe-haut.

De plus, on observe que le signal de sortie des passe-bas a une valeur moyenne différente de zéro, alors que la sortie du passe-haut a une valeur moyenne nulle.

Le contenu spectral d'un signal explique ce phénomène. En effet, la valeur moyenne est la composante de pulsation nulle ($\omega = 0$) du signal. Elle passe à travers les passe-bas et se retrouve intacte en sortie (dans les deux cas, $\underline{H}(0) = 1$) : la valeur moyenne du créneau en entrée se retrouve en sortie. Le filtrage passe-haut, quant à lui, coupe les basses fréquences, donc la valeur moyenne. Le signal de sortie d'un passe-haut a donc nécessairement une valeur moyenne nulle.

Le signal de sortie ne peut avoir de valeur moyenne non nulle que si le système présente un caractère passe-bas.

Attendu que la valeur moyenne est la composante de pulsation nulle du signal, en notant $e(t)$ l'entrée et $s(t)$ la sortie d'un système de transmittance $\underline{H}(j\omega)$:

$$\langle s \rangle = \underline{H}(0) \langle e \rangle .$$

Et pour un passe-bande, qui ne laisse passer ni les basses ni les hautes fréquences ? *A priori*, le signal de sortie, lors d'un essai indiciel ou d'un régime libre :

- admet une valeur moyenne nulle, car la composante $\omega = 0$ est filtrée,
- reste continu à l'arrivée de l'échelon ou du régime libre, car les hautes fréquences, indispensables à une brusque variation du signal, sont filtrées.

L'observation expérimentale valide ces deux résultats :

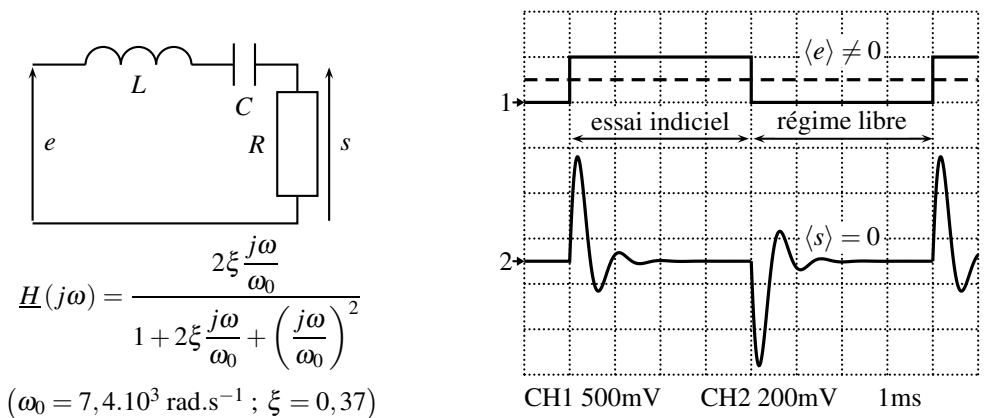


Figure 12.25 – Montage, transmittance, essai indiciel et régime libre d'un passe-bande du deuxième ordre.

5.2 Complément : lien avec le théorème de la valeur initiale

Le cours de SII introduit un outil très utile pour l'étude des essais indiciels : le théorème de la valeur initiale. Si sa démonstration ne relève pas du présent cours, ses enseignements sont totalement liés à ceux du paragraphe précédent.

La fonction de transfert de Laplace est la même que celle en régime permanent sinusoïdal, en remplaçant $j\omega$ par p . Le théorème de la valeur initiale affirme que, lorsque les limites existent :

$$\lim_{t \rightarrow 0} s(t) = \lim_{p \rightarrow \infty} pS(p) = \lim_{p \rightarrow \infty} pH(p)E(p).$$

La transformée de Laplace d'un échelon de hauteur E_0 est $E(p) = E_0/p$ et donc, pour un essai indiciel :

$$\lim_{t \rightarrow 0} s(t) = \lim_{p \rightarrow \infty} H(p)E_0 \quad \text{soit} \quad s(0^+) = H(\infty)E_0.$$

Ainsi, si le filtre présente un caractère passe-bas, $H(\infty) = 0$ et $s(0^+) = 0$; la sortie reste continue. Par contre, si le filtre présente un caractère passe-haut, $H(\infty) \neq 0$ et $s(0^+) \neq 0$; la sortie est discontinue lors d'un essai indiciel. De plus, c'est la valeur de la fonction de transfert en l'infini qui importe, c'est-à-dire pour les hautes fréquences. On retrouve qu'une brusque variation d'un signal est due aux hautes fréquences.

6 Approche documentaire : accéléromètre

La notion de filtrage n'est pas réservée aux circuits électriques. Elle peut être étendue à tout type de système, par exemple mécanique. Il existe de nombreux filtres mécaniques, comme les sismomètres, les amortisseurs, les accéléromètres...

Un extrait de la notice constructeur d'un accéléromètre est présentée sur la page suivante. Il s'agit d'une puce qui transforme une accélération mécanique en un signal électrique. Ce genre de puce est intégrée dans toutes les manettes de jeux de console, dans les téléphones portables...

L'accéléromètre présenté est le ADXL330. Les deux premières lettres d'une puce indiquent le fabricant, ici Analog Device ; les lettres et les chiffres qui suivent précisent de quel circuit il s'agit. Le lecteur retrouvera la fiche technique entière sur internet en cherchant « *ADXL330 datasheet* ».

Quelques questions :

1. le ADXL330 mesure-t-il l'accélération suivant un seul axe ou selon les 3 coordonnées d'espace ?
2. Les fréquences qui passent à travers un filtre constituent la **bande passante**, *bandwidth* en anglais. Quelles sont les bandes passantes maximales suivant les trois axes ?
3. L'utilisateur peut-il diminuer ces bandes passantes ? les augmenter ?
4. Qu'arrive-t-il aux fréquences qui sont en dehors de la bande passante ?
5. Les dimensions de la puce sont précisées dans l'extrait. Dessiner de profil la puce et indiquer à quoi correspondent les trois axes.

CHAPITRE 12 – FILTRAGE LINÉAIRE



Small, Low Power, 3-Axis $\pm 3\text{ g}$
iMEMS® Accelerometer

ADXL330

FEATURES

- 3-axis sensing
- Small, low-profile package
4 mm $\times$ 4 mm $\times$ 1.45 mm LFCSP
- Low power
180 μA at $V_s = 1.8\text{ V}$ (typical)
- Single-supply operation
1.8 V to 3.6 V
- 10,000 g shock survival
- Excellent temperature stability
- BW adjustment with a single capacitor per axis
- RoHS/WEEE lead-free compliant

APPLICATIONS

- Cost-sensitive, low power, motion- and tilt-sensing applications
- Mobile devices
- Gaming systems
- Disk drive protection
- Image stabilization
- Sports and health devices

GENERAL DESCRIPTION

The ADXL330 is a small, thin, low power, complete 3-axis accelerometer with signal conditioned voltage outputs, all on a single monolithic IC. The product measures acceleration with a minimum full-scale range of $\pm 3\text{ g}$. It can measure the static acceleration of gravity in tilt-sensing applications, as well as dynamic acceleration resulting from motion, shock, or vibration.

The user selects the bandwidth of the accelerometer using the C_x , C_y , and C_z capacitors at the X_{OUT} , Y_{OUT} , and Z_{OUT} pins. Bandwidths can be selected to suit the application, with a range of 0.5 Hz to 1600 Hz for X and Y axes, and a range of 0.5 Hz to 550 Hz for the Z axis.

The ADXL330 is available in a small, low profile, 4 mm $\times$ 4 mm $\times$ 1.45 mm, 16-lead, plastic lead frame chip scale package (LFCSP_LQ).

FUNCTIONAL BLOCK DIAGRAM

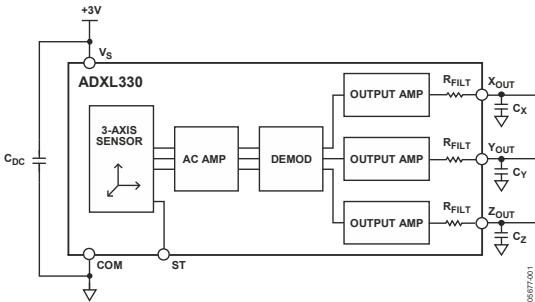


Figure 1.

Rev. A
Information furnished by Analog Devices is believed to be accurate and reliable. However, no responsibility is assumed by Analog Devices for its use, nor for any infringements of patents or other rights of third parties that may result from its use. Specifications subject to change without notice. No license is granted by implication or otherwise under any patent or patent rights of Analog Devices. Trademarks and registered trademarks are the property of their respective owners.

- 6. Quelle est l'ordre de grandeur de la fréquence maximale d'oscillation qu'un expérimentateur normal peut atteindre avec une manette de jeux ou un téléphone ? Le résultat est-il en accord avec les spécifications du ADXL330 ?
- 7. Pourquoi la bande passante doit-elle être supérieure à la fréquence maximale d'oscillation précédente ?

Des réponses sont proposées à la page suivante.

SYNTHÈSE

SAVOIRS

- développement en série de Fourier
- définir la valeur moyenne
- définir la valeur efficace

SAVOIR-FAIRE

- utiliser une fonction de transfert ou sa représentation graphique pour conduire l'étude de la réponse d'un système linéaire à une excitation sinusoïdale, à une somme finie d'excitations sinusoïdales, à un signal périodique
- expliciter les conditions d'utilisation d'un filtre en moyeneur

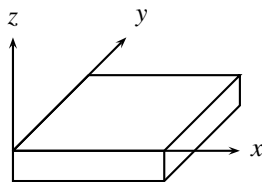
MOTS-CLÉS

- développement en série de Fourier
- valeur efficace
- réponse d'un système linéaire
- valeur moyenne
- moyeneur

CHAPITRE 12 – FILTRAGE LINÉAIRE

Réponses proposées aux questions de l'analyse documentaire :

1. le ADXL330 mesure l'accélération suivant les 3 coordonnées d'espace : *The ADXL is a [...] complete 3-axis accelerometer.*
2. La bande passante maximale est de 1600 Hz suivant (Ox) et (Oy) et de 550 Hz suivant (Oz).
3. L'utilisateur peut les diminuer jusqu'à 0,5 Hz, mais ne peut pas les augmenter.
4. Les fréquences en dehors de la bande passante sont éliminées. La documentation ne précise pas si l'atténuation est forte ou pas, l'utilisateur doit donc considérer que toute fréquence qui sort de la bande passante est filtrée.
5. Les dimensions sont 4 mm $\times$ 4 mm $\times$ 1,45 mm. Attendu que les caractéristiques sont identiques suivant (Ox) et (Oy), ces deux axes sont équivalents : la puce mesure 4 mm suivant ces axes, et donc 1,45 mm suivant (Oz).

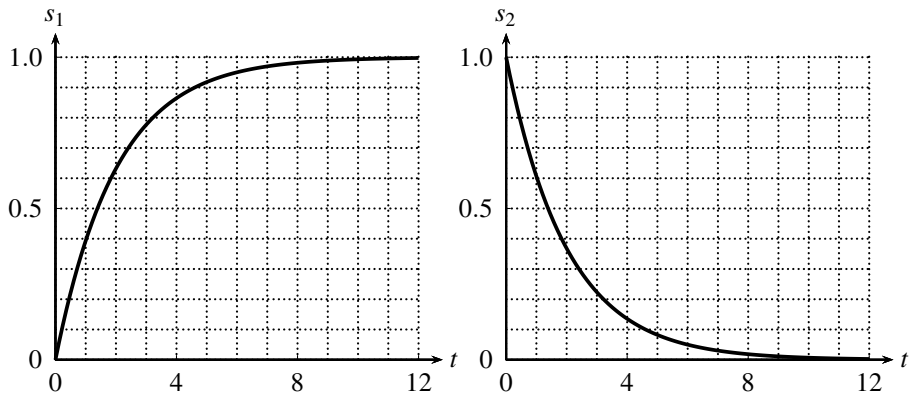


6. L'auteur arrive à mener environ 10 oscillations du poignet par seconde au maximum, ce qui correspond à une fréquence de 10 Hz. Les oscillations physiques sont donc comprises entre une fréquence de 0 Hz, pour un mouvement extrêmement lent, à 10 Hz, pour des oscillations rapides. Cette gamme de fréquence entre dans la bande passante du ADXL330.
7. La question précédente n'envisage qu'un mouvement harmonique. Un mouvement d'un point à un autre de l'espace, avec un départ très rapide, n'est pas sinusoïdal, mais contient de nombreuses fréquences, en particulier des hautes fréquences, qui permettent de décrire un signal qui varie très rapidement. Pour bien capter un mouvement quelconque, l'accéléromètre doit être capable de prendre en compte ces hautes fréquences. C'est pourquoi la bande passante doit être largement supérieure à 10 Hz.

S'ENTRAÎNER

12.1 Identification d'un système (★)

Identifier la fonction de transfert du système dont on présente la réponse indicielle unitaire, c'est-à-dire lorsque l'entrée est un échelon de hauteur $E_0 = 1\text{ V}$:

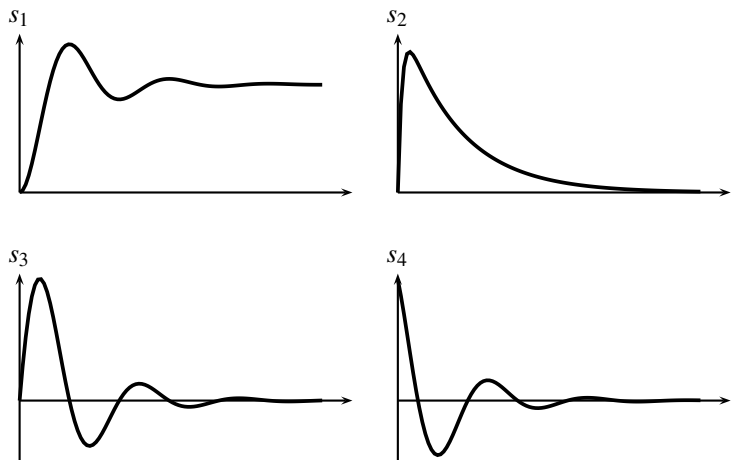


12.2 Filtrage par un premier ordre (★)

Rappeler la fonction de transfert canonique d'un passe-bas du premier ordre. Quelle est la sortie du montage si l'entrée vaut $e(t) = E \sin^3(\omega_0 t)$? On précise : $\sin^3(\theta) = \frac{3}{4} \sin(\theta) - \frac{1}{4} \sin(3\theta)$

12.3 Identification graphique d'un système (★)

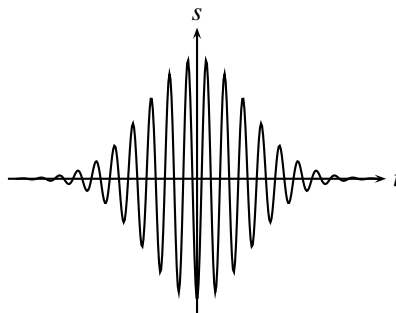
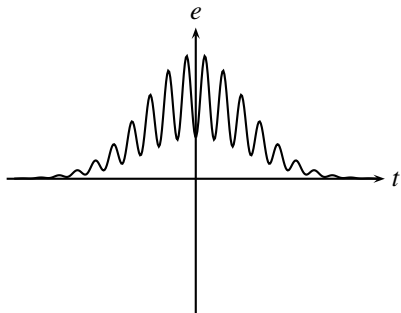
Préciser sans calcul le caractère passe-bas, passe-haut, passe-bande, premier ou second ordre des systèmes dont on présente la réponse indicielle (on ne demande pas d'identifier la fonction de transfert) :



CHAPITRE 12 – FILTRAGE LINÉAIRE

12.4 Filtrage graphique (★)

Quel type de filtre permet-il de passer du signal e au signal s ? Expliquer.



12.5 Action d'un filtre passe-haut sur un signal (★)

On considère un filtre passe-haut du premier ordre dont la fréquence de coupure est 100 Hz. Donner l'allure du signal recueilli en sortie du filtre si on envoie en entrée :

1. une sinusoïde d'amplitude 4 V centrée autour de 1 V et de fréquence 2 kHz,
2. une sinusoïde d'amplitude 4 V centrée autour de 0 V et de fréquence 2 kHz,
3. un créneau d'amplitude 4 V centré autour de 1 V et de fréquence 2 kHz,
4. un créneau d'amplitude 4 V centré autour de 0 V et de fréquence 2 kHz.

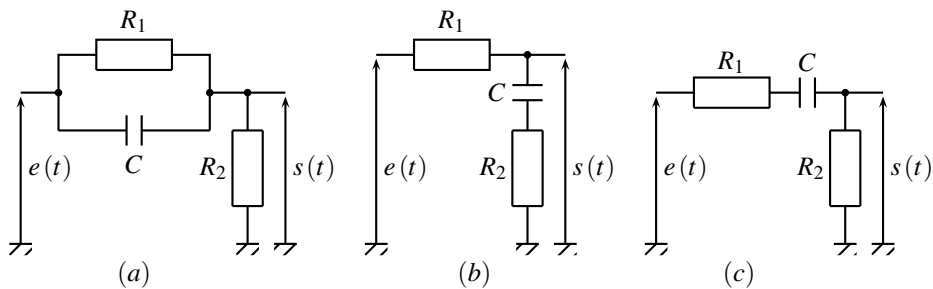
12.6 Action d'un filtre passe-bas sur un signal (★)

On considère un filtre passe-bas du premier ordre dont la fréquence de coupure est 100 Hz. Donner l'allure du signal recueilli en sortie du filtre si on envoie en entrée :

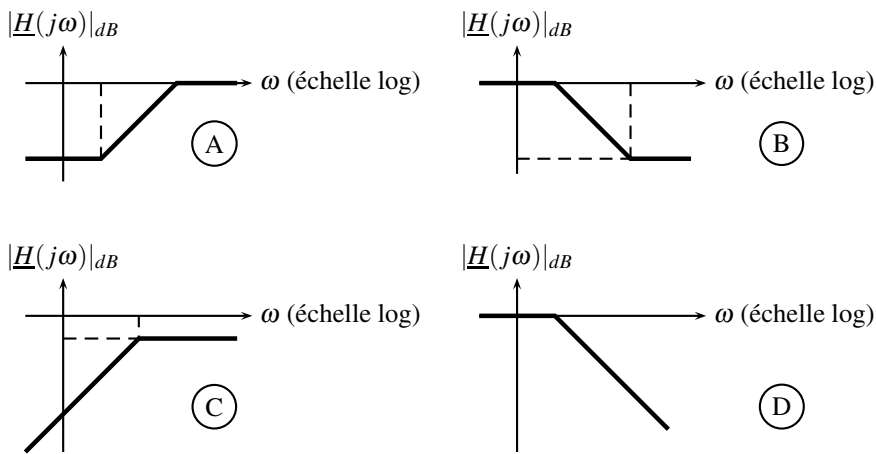
1. une sinusoïde d'amplitude 4 V centrée autour de 1 V et de fréquence 2 kHz,
2. une sinusoïde d'amplitude 4 V centrée autour de 0 V et de fréquence 2 kHz,
3. un créneau d'amplitude 4 V centré autour de 1 V et de fréquence 2 kHz,
4. un créneau d'amplitude 4 V centré autour de 0 V et de fréquence 2 kHz,
5. un créneau d'amplitude 4 V centré autour de 1 V et de fréquence 75 Hz,
6. un créneau d'amplitude 4 V centré autour de 0 V et de fréquence 75 Hz.

APPROFONDIR

12.7 Diagrammes de Bode, essais indiciels et filtrage (★)



- 1. Établir les fonctions de transfert de chaque circuit présenté page suivante, H_a , H_b et H_c .
- 2. Associer chaque fonction de transfert à un diagramme de Bode asymptotique proposé ci-après, en précisant les pentes et les pulsations remarquables. En déduire les phases dans chaque domaine.

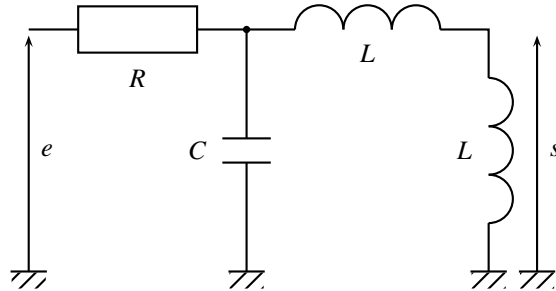


- 3. Calculer et représenter l'essai indiciel, c'est à dire la sortie associé à un échelon de tension de hauteur E_0 :
- $$\begin{cases} e(t) = 0 & \text{pour } t < 0 \\ e(t) = E & \text{pour } t \geq 0 \end{cases}$$
- 4. Expliquer sans calcul pourquoi toutes les réponses indicielles ont le même comportement en $t = 0$ (sortie continue ou discontinue).

CHAPITRE 12 – FILTRAGE LINÉAIRE

12.8 Filtre de Hartley (★★)

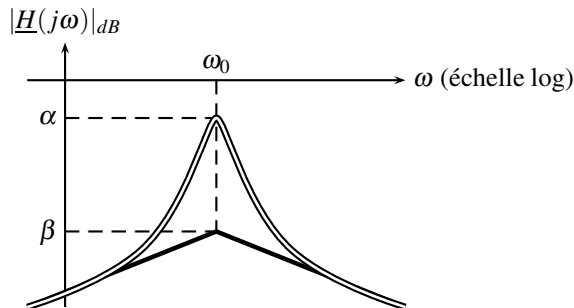
On réalise le montage ci-après.



1. Établir sa transmittance $H(p)$ sous la forme : $H(p) = \frac{H_0 2m \frac{p}{\omega_0}}{1 + 2m \frac{p}{\omega_0} + \frac{p^2}{\omega_0^2}}$ et exprimer H_0, m

et ω_0 en fonction de R, L et C .

2. Dans le cas où $L = 1,0 \text{ mH}$, $C = 100,0 \text{ nF}$ et $R = 10,0 \text{ k}\Omega$, le diagramme de Bode en amplitude a l'allure présentée page suivante. Identifier les pentes des asymptotes, les valeurs de α et β . En déduire l'allure du diagramme de Bode en phase.



3. Le montage peut-il servir d'intégrateur ou de dérivateur? Si oui, dans quelle bande de fréquence?

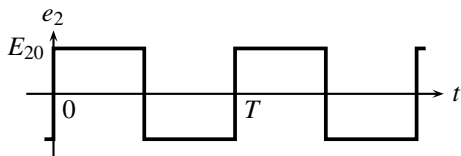
4. On étudie la sortie $s_1(t)$ associée à l'entrée $e_1(t) = E_0 + E_1 \cos(\omega_1 t)$ où $\omega_1 = \frac{1}{\sqrt{2LC}}$.

Comment réaliser expérimentalement ce signal au laboratoire d'électronique?

5. Calculer l'expression littérale de la sortie $s_1(t)$, observée sur l'oscilloscope en régime permanent.

6. On étudie maintenant la sortie $s_2(t)$ associée au signal créneaux $e_2(t)$, de période $T = 6\pi\sqrt{2LC}$, d'amplitude $E_{20} = 1 \text{ V}$. Il est décomposable en série de Fourier.

$$e_2(t) = \frac{4E_{20}}{\pi} \left[\sin(\omega_2 t) + \frac{\sin(3\omega_2 t)}{3} + \frac{\sin(5\omega_2 t)}{5} + \dots + \frac{\sin[(2n+1)\omega_2 t]}{2n+1} + \dots \right]$$



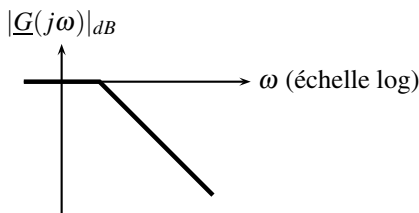
- Calculer la valeur efficace E_{2eff} de e_2 .
7. Tracer l'allure du spectre de e_2 . Préciser numériquement les pulsations correspondantes aux 3 premières harmoniques.
8. Calculer numériquement les amplitudes des 3 premières harmoniques du signal de sortie s_2 . Justifier alors le nom de « tripleur de fréquence » donné au montage. Préciser l'expression numérique de $s_2(t)$.
9. On alimente dorénavant le montage avec un échelon $e_3(t)$ de hauteur $E_{30} = 1\text{ V}$. La sortie est notée $s_3(t)$.
- Tracer l'allure de cet échelon. À quoi sert un essai indiciel ?
10. Expliquer sans calcul pourquoi $s_3(0^+) = 0$.
11. En admettant que $\frac{ds_3}{dt}(0^+) = \frac{E_{30}}{2RC}$, établir l'expression numérique de $s_3(t)$. Les valeurs numériques permettront de grandement simplifier la résolution de l'équation différentielle. Que dire de la pseudopulsation qui apparaît dans le résultat ?

12.9 Filtre passe-bas (★)

On étudie un montage dont la fonction de transfert est : $\underline{G} = -\frac{1}{1 + 3jRC_1\omega + R^2C_1C_2(j\omega)^2}$.

1. Montrer qu'on peut l'écrire sous la forme : $\underline{G} = \frac{G_0}{1 + 2jm\frac{\omega}{\omega_0} - \frac{\omega^2}{\omega_0^2}}$. On explicitera les expressions de m , G_0 et ω_0 .
2. En déduire l'équation différentielle vérifiée par $s(t)$.
3. Que décrit la solution générale de cette équation ? Préciser les différentes allures possibles de cette solution générale.
4. On applique une tension constante E_0 à partir de l'instant $t = 0$. Donner, en les justifiant, les valeurs de $s(t)$ quand $t \rightarrow 0^+$ (notée $s(0^+)$) et quand $t \rightarrow \infty$ (notée $s(\infty)$). À quoi correspond physiquement la valeur de $s(\infty)$?
5. On veut obtenir $m = \frac{\sqrt{2}}{2}$. Calculer, en fonction de C_1 , la valeur qu'il faut donner à C_2 pour y parvenir. Quelle est alors la nature du régime libre ?
6. On veut faire varier la fréquence $f_0 = \frac{\omega_0}{2\pi}$ entre 1 kHz et 4 kHz. Entre quelles limites (exprimées en fonction de C_1) doit-on choisir R ? Dans la suite, on supposera que $f_0 = 2\text{ kHz}$.
7. L'allure du diagramme de Bode en amplitude est représentée page suivante. Préciser la valeur de la pulsation de cassure, les pentes et les valeurs limites du déphasage.
8. Ce circuit peut-il être utilisé en intégrateur ? en dérivateur ?
9. On envoie un signal de fréquence $f = 1,5\text{ kHz}$.

CHAPITRE 12 – FILTRAGE LINÉAIRE



a. Le signal est sinusoïdal centré. Représenter, en les justifiant, les signaux d'entrée et de sortie de ce filtre.

b. Mêmes questions pour un signal créneau pair de valeur basse 0 et de valeur haute a . On rappelle que la décomposition en séries de Fourier d'un tel signal est :

$$\frac{a}{2} + \frac{a}{\pi} \sum_{p=0}^{+\infty} \frac{(-1)^p}{2p+1} \cos((2p+1)2\pi ft)$$

On calculera éventuellement des ordres de grandeur pour justifier l'allure obtenue.

c. Mêmes questions pour un signal triangulaire pair, centré, d'amplitude a .

10. Quel est l'avantage de ce filtre par rapport à un filtre passe-bas du premier ordre ?

12.10 Filtrage d'un signal électrique (★)

On étudie un filtre de fonction de transfert $H(j\omega) = -\frac{1}{1 + 3jRC_1\omega + R^2C_1C_2(j\omega)^2}$.

1. a. Quelle est la nature de ce filtre ?

b. Écrire la fonction de transfert sous forme canonique. On précisera les valeurs des paramètres en fonction de R , C_1 et C_2 .

c. Montrer alors qu'il est possible d'obtenir les valeurs des capacités connaissant la valeur de la résistance R et les caractéristiques du filtre.

d. Si ce filtre est réalisé à l'aide de résistances dont les valeurs sont connues avec une précision infinie et qu'on a une précision de 5 % sur la détermination des caractéristiques du filtre, quelle est la précision (qu'on supposera identique) sur les mesures des capacités ?

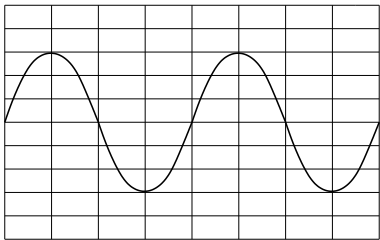
e. Application numérique : $R = 1,00 \text{ k}\Omega$, facteur de qualité $Q = 0,707$ et fréquence de résonance $f_0 = 20,0 \text{ kHz}$. Donner les valeurs numériques de C_1 et C_2 .

2. a. On dispose d'un générateur basses fréquences de fréquence maximale 2 MHz et d'un oscilloscope deux voies ainsi que de tous les cables nécessaires. Décrire la manière dont on pourra réaliser le tracé du diagramme de Bode.

b. On observe le chronogramme page suivante pour un signal d'entrée particulier. On a les calibres suivants : verticalement 1 V par division, horizontalement $12,5 \mu\text{s}$ par division. Quelle(s) est(sont) la(es) fréquence(s) de ce signal d'entrée et du signal de sortie correspondant ?

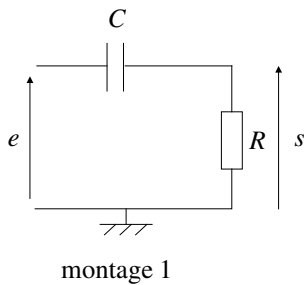
c. Calculer avec les valeurs numériques précédemment fournies le module de la fonction de transfert ainsi que son déphasage.

d. Représenter sur un même chronogramme les signaux d'entrée et de sortie.

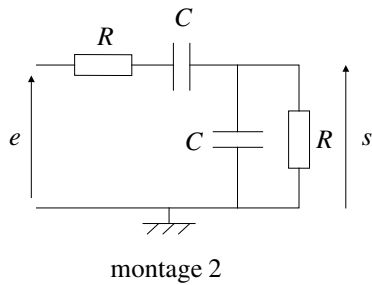


e. L'oscilloscope dispose d'un module permettant de tracer la décomposition en série de Fourier des signaux étudiés. Comment peut-on observer le diagramme de Bode en gain en utilisant un signal créneau ? On rappelle que l'analyse harmonique d'un créneau donne une décroissance en $\frac{1}{n}$ des amplitudes et qu'il est donc possible de considérer que les harmoniques de rangs élevés ont des amplitudes approximativement constantes.

12.11 Étude et utilisation de filtres, d'après ENSAIT 2002 (★)

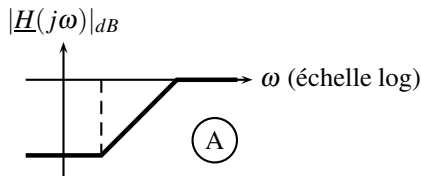


montage 1

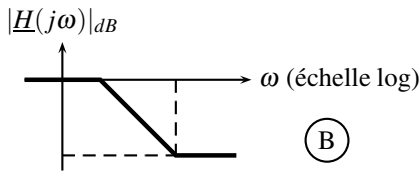


montage 2

1. Déterminer la fonction de transfert $\underline{H}$ du montage 1.
2. Préciser la nature du filtre.
3. Définir le sens de variation du gain.
4. Quel est le diagramme de Bode en amplitude qui correspond au filtre ? A, B ? C, D (représentés page suivante) ? Préciser la(les) pulsation(s) de cassure et la(les) pente(s).



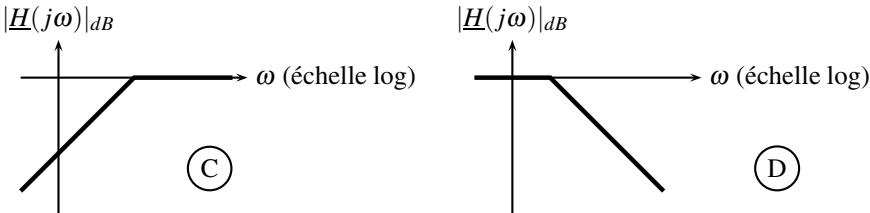
(A)



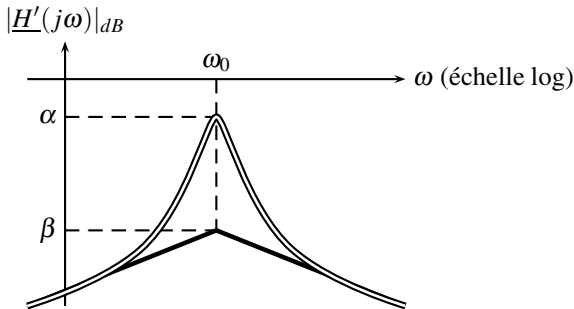
(B)

5. Déterminer la valeur à donner à la capacité C pour que la tension de sortie soit affaiblie de 6 dB par rapport à la tension d'entrée lorsque celle-ci est un signal sinusoïdal de fréquence 1,0 kHz et que $R = 10,0 \text{ k}\Omega$.
6. Déterminer la fonction de transfert $\underline{H}'$ du montage 2.
7. Préciser la nature du filtre.

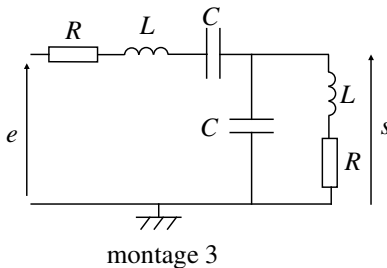
CHAPITRE 12 – FILTRAGE LINÉAIRE



8. L'écrire sous la forme canonique.
9. L'allure du diagramme de Bode en amplitude est la suivante :



- Préciser la valeur des pentes et la pulsation de cassure.
10. En prenant les mêmes valeurs que précédemment pour R et C , donner la valeur de la fréquence de résonance.
11. Pour une fréquence $f = 1,0$ kHz, donner la valeur du gain pour le fondamental et les harmoniques d'ordre inférieure ou égal à 10.
12. Que peut-on en conclure sur l'influence du filtre sur un signal de fréquence f ?
13. On considère le filtre du montage 3. Sur quelle fréquence est-il accordé sachant que la fréquence d'accord est la fréquence de résonance en intensité du circuit R, L, C série ?



14. Calculer les impédances du circuit R, L, C série et du circuit R, L série en parallèle avec C à la fréquence d'accord.
15. On applique une tension sinusoïdale de valeur efficace 0,5 V et de fréquence la fréquence d'accord. Déterminer la tension de sortie du filtre à la fréquence d'accord.

CORRIGÉS

12.1 Identification d'un système

1. Pour le système 1, la sortie est celle d'un passe-bas du premier ordre (variation exponentielle) de transmittance canonique : $H_1(j\omega) = \frac{H_0}{1 + j\tau\omega}$.

Alimenté par une échelon de hauteur E_0 , $s_1(\infty) = H_0E_0$. On retrouve facilement ce point avec la résolution de l'équation différentielle qui modélise le comportement du système :

$(1 + j\tau\omega) \underline{s_1} = H_0 \underline{e_1}$ donc $s_1 + \tau \frac{ds_1}{dt} = H_0 e_1 = H_0 E_0$, qui s'intègre en : $s_1(t) = H_0 E_0 + \mu \exp\left(-\frac{t}{\tau}\right) \xrightarrow{t \rightarrow \infty} H_0 E_0$.

On peut aussi utiliser le théorème de la valeur finale, vue en SI, sachant que la transformée d'un échelon est E_0/p :

$$\lim_{t \rightarrow \infty} s_1(t) = \lim_{p \rightarrow 0} p S_1(p) = \lim_{p \rightarrow 0} p H_1(p) E_1(p) = \lim_{p \rightarrow 0} p H_1(p) \frac{E_0}{p} = H_1(0) E_0 = H_0 E_0.$$

L'échelon est unitaire donc $E_0 = 1$; par lecture graphique : $H_0 = 1$.

La réponse indicielle vaut 63% de la valeur finale en $t = \tau$. On lit graphiquement : $\tau = 2s$.

2. Pour le système 2, la sortie est celle d'un passe-haut du premier ordre (variation exponentielle) de transmittance canonique : $H_2(p) = H_0 \frac{j\tau\omega}{1 + j\tau\omega}$. En effet, attendu que le signal est discontinu en $t = 0$, les hautes fréquences doivent passer dans le système.

Pour un échelon de hauteur E_0 , $s_2(0) = H_0 E_0$. La méthode la plus simple pour retrouver facilement ce point est d'utiliser le théorème de la valeur initiale, vue en SI, sachant que la transformée d'un échelon est E_0/p : $\lim_{t \rightarrow 0} s_2(t) = \lim_{p \rightarrow \infty} p S_2(p) = \lim_{p \rightarrow \infty} p H_2(p) \frac{E_0}{p} = H_2(\infty) E_0 = H_0 E_0$.

La réponse indicielle a perdu 63% de la valeur initiale en $t = \tau$. On lit graphiquement : $\tau = 2s$. L'échelon est unitaire donc $E_0 = 1$; par lecture graphique : $H_0 = 1$.

12.2 Filtrage par un premier ordre

La transmittance canonique d'un passe-bas du premier ordre est $H(j\omega) = \frac{H_0}{1 + j\tau\omega}$.

Le signal d'entrée n'est pas harmonique mais décomposable en une somme de fonctions harmoniques : $e(t) = E \sin^3(\omega_0 t) = \frac{E}{4} (3 \sin(\omega_0 t) - \sin(3\omega_0 t))$.

En **régime sinusoïdal permanent** (indispensable pour utiliser les notations complexes), la sortie est la somme des sorties associées à $e_1(t) = \frac{3E}{4} \sin(\omega_0 t)$ (le fondamental) et $e_3(t) = -\frac{E}{4} \sin(3\omega_0 t)$ (l'harmonique de rang 3). Calculons $s_1(t)$ et $s_3(t)$, pour $\omega = \omega_0$:

CHAPITRE 12 – FILTRAGE LINÉAIRE

$$\begin{aligned}\underline{s}_1(t) &= \underline{H}(j\omega = j\omega_0) \underline{e}_1(t) \\ \underline{s}_1(t) &= \frac{H_0}{1 + j\tau\omega_0} \frac{3E}{4} \exp\left(j\left(\omega_0 t - \frac{\pi}{2}\right)\right) \\ \underline{s}_1(t) &= \frac{H_0}{\sqrt{1 + (\tau\omega_0)^2}} \exp(-j \arctan \tau\omega_0) \frac{3E}{4} \exp\left(j\left(\omega_0 t - \frac{\pi}{2}\right)\right) \\ \underline{s}_1(t) &= \frac{3EH_0}{4\sqrt{1 + (\tau\omega_0)^2}} \exp\left(j\left(\omega_0 t - \arctan \tau\omega_0 - \frac{\pi}{2}\right)\right) \\ s_1(t) &= \frac{3EH_0}{4\sqrt{1 + (\tau\omega_0)^2}} \sin(\omega_0 t - \arctan \tau\omega_0).\end{aligned}$$

et pour $\omega = 3\omega_0$:

$$\begin{aligned}\underline{s}_3(t) &= \underline{H}(j\omega = 3j\omega_0) \underline{e}_3(t) \\ \underline{s}_3(t) &= -\frac{H_0}{1 + 3j\tau\omega_0} \frac{E}{4} \exp\left(j\left(3\omega_0 t - \frac{\pi}{2}\right)\right) \\ \underline{s}_3(t) &= -\frac{H_0}{\sqrt{1 + (3\tau\omega_0)^2}} \exp(-j \arctan 3\tau\omega_0) \frac{E}{4} \exp\left[j\left(3\omega_0 t - \frac{\pi}{2}\right)\right] \\ \underline{s}_3(t) &= -\frac{EH_0}{4\sqrt{1 + (3\tau\omega_0)^2}} \exp\left[j\left(3\omega_0 t - \arctan 3\tau\omega_0 - \frac{\pi}{2}\right)\right] \\ s_3(t) &= -\frac{EH_0}{4\sqrt{1 + (3\tau\omega_0)^2}} \sin(3\omega_0 t - \arctan 3\tau\omega_0).\end{aligned}$$

La sortie totale $s(t)$ est donc en régime sinusoïdal permanent :

$$s(t) \frac{EH_0}{4} \left(\frac{3}{\sqrt{1 + (\tau\omega_0)^2}} \sin(\omega_0 t - \arctan \tau\omega_0) - \frac{1}{\sqrt{1 + (3\tau\omega_0)^2}} \sin(3\omega_0 t - \arctan 3\tau\omega_0) \right).$$

12.3 Identification d'un système

1. Système 1 : pas de discontinuité en $t = 0$ donc pas de haute fréquence dans le signal. Valeur finale $\neq 0$ donc composante continue présente. Le système est donc un passe-bas du second ordre (oscillations).
2. Système 2 : pas de discontinuité en $t = 0$ donc pas de haute fréquence dans le signal. Valeur finale nulle donc pas de composante continue. Le système est donc un passe-bande du second ordre (car il n'existe pas de passe-bande du premier ordre).
3. Système 3 : pas de discontinuité en $t = 0$ donc pas de haute fréquence dans le signal. Valeur finale nulle donc pas de composante continue. Le système est donc un passe-bande du second ordre (oscillations).
4. Système 4 : discontinuité en $t = 0$ donc hautes fréquences dans le signal. Valeur finale nulle donc pas de composante continue. Le système est donc un passe-haut du second ordre (oscillations).

12.4 Filtrage graphique

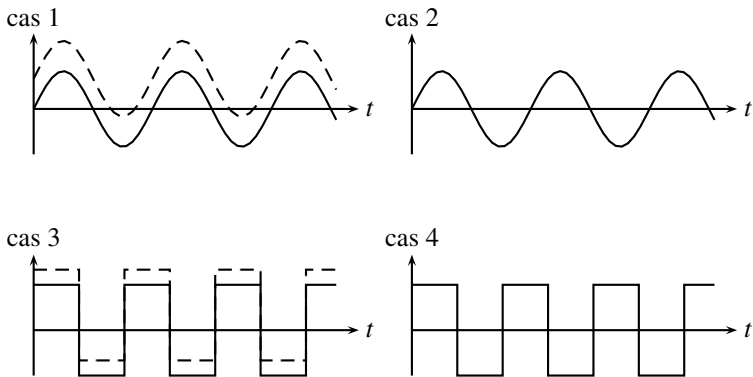
Les oscillations de période de T de e se retrouvent dans s , elles doivent passer à travers le filtre.

e a une valeur moyenne différente de 0, mais s a une valeur moyenne nulle. La composante continue, c'est à dire la pulsation nulle, ne passe pas à travers le filtre.

C'est donc un filtre passe-haut qui permet de passer de e à s . Sa pulsation caractéristique ω_0 soit être inférieure à $2\pi/T$ pour laisser passer les oscillations de e .

12.5 Action d'un filtre passe-haut sur un signal

En supposant que le filtre est idéal, il ne laisse passer que les fréquences supérieures à 100 Hz et coupe toutes les autres. Pour des signaux de 2 kHz, cela revient à conserver le fondamental et toutes les harmoniques et à supprimer la composante continue. On obtient donc les signaux d'entrée centrés autour de la valeur 0 V dont les allures sont les suivantes (s est en traits pleins, e en pointillés) :



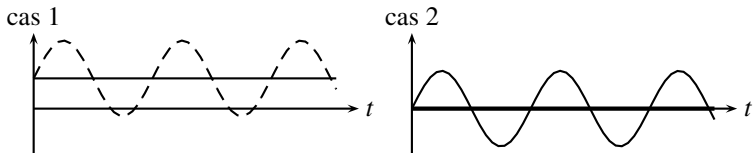
12.6 Action d'un filtre passe-bas sur un signal

En supposant que le filtre est idéal, il ne laisse passer que les fréquences inférieures à 100 Hz et coupe toutes les autres.

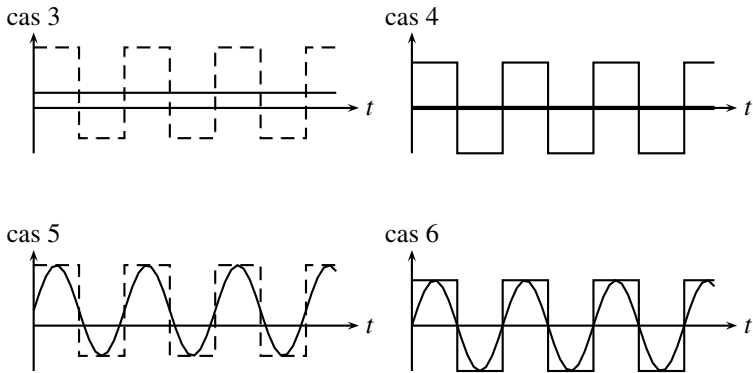
Pour des signaux de 2 kHz, cela revient à ne conserver que la composante continue et à supprimer le fondamental et toutes les harmoniques. On obtient un signal continu.

Pour des signaux de 75 Hz, cela revient à ne conserver que la composante continue et le fondamental et à supprimer toutes les autres harmoniques : on obtient donc une sinusoïde, pas nécessairement centrée sur 0.

Les allures sont donc les suivantes (s est en traits pleins, e en pointillés) :



CHAPITRE 12 – FILTRAGE LINÉAIRE



12.7 Diagrammes de Bode, essais indiciels et filtrage

1. Pour le circuit (a), on reconnaît un diviseur de tension :

$$\frac{s_a}{e_a} = \frac{R_2}{R_2 + \frac{1}{\frac{1}{R_1} + jC\omega}} = \frac{R_2 \left(\frac{1}{R_1} + jC\omega \right)}{R_2 \left(\frac{1}{R_1} + jC\omega \right) + 1}.$$

Mettons numérateur et dénominateur sous forme canonique :

$$\frac{s_a}{e_a} = \frac{\frac{R_2}{R_1} (1 + jR_1C\omega)}{\left(1 + \frac{R_2}{R_1} \right) \left(1 + j \frac{R_1R_2}{R_1 + R_2} C\omega \right)} = \frac{R_2}{R_1 + R_2} \frac{1 + j\tau\omega}{1 + ja\tau\omega}.$$

Finalement : $H_a(j\omega) = a \frac{1 + j\tau\omega}{1 + ja\tau\omega}$ avec $a = \frac{R_2}{R_1 + R_2} < 1$ et $\tau = R_1C$.

Le circuit (b) est aussi un diviseur de tension : $\frac{s_b}{e_b} = \frac{\frac{1}{jC\omega} + R_2}{\frac{1}{jC\omega} + R_2 + R_1} = \frac{1 + jR_2C\omega}{1 + j(R_1 + R_2)C\omega} =$

$H_b(j\omega)$.

Encore un diviseur de tension pour le circuit (c) : $\frac{s_c}{e_c} = \frac{R_2}{R_1 + \frac{1}{jC\omega} + R_2} = \frac{jR_2C\omega}{1 + j(R_1 + R_2)C\omega} =$

$H_c(j\omega)$.

2. Dressons un tableau asymptotique pour simplifier H_a , sachant que deux pulsations caractéristiques apparaissent, $1/\tau$ au numérateur et $1/(a\tau)$ au dénominateur (page suivante). On reconnaît le diagramme (A).

Un tableau asymptotique permet, de même, de simplifier H_b . Les pulsations caractéristiques sont $\frac{1}{R_2C}$ au numérateur et $\frac{1}{(R_1 + R_2)C}$ au dénominateur. Ainsi, on reconnaît le diagramme (B).

	$\frac{1}{\tau}$ $\frac{1}{(a\tau)}$	
a	a	a
$1 + j\tau\omega$	1	$j\tau\omega$
$1 + ja\tau\omega$	1	$ja\tau\omega$
$H_{a,asympt}(j\omega)$	a	1
pente	nulle	$+20\text{ dB/dec}$
phase	0	$+\pi/2$

	$\frac{1}{(R_1 + R_2)C}$ $\frac{1}{R_2C}$	
$1 + jR_2C\omega$	1	$jR_2C\omega$
$1 + j(R_1 + R_2)C\omega$	1	$j(R_1 + R_2)C\omega$
$H_{b,asympt}(p)$	1	$\frac{R_2}{R_1 + R_2} < 1$
pente	nulle	-20 dB/dec
phase	0	$-\pi/2$

H_c est simplifié avec un tableau asymptotique dans lequel intervient l'unique pulsation, au dénominateur. Le numérateur n'est pas simplifiable car composé d'un seul terme ; il n'est pas négligeable devant un autre. Ainsi :

	$\frac{1}{(R_1 + R_2)C}$	
$jR_2C\omega$	$jR_2C\omega$	$jR_2C\omega$
$1 + j(R_1 + R_2)C\omega$	1	$j(R_1 + R_2)C\omega$
$H_{c,asympt}(j\omega)$	$jR_2C\omega$	$\frac{R_2}{R_1 + R_2} < 1$
pente	$+20\text{ dB/dec}$	nulle
phase	$\pi/2$	0

On reconnaît le diagramme (C).

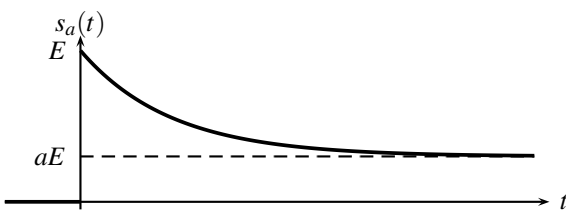
3. Pour la réponse indicielle du circuit (a), l'équation différentielle régissant l'évolution du système est trouvée avec l'équation différentielle : $(1 + ja\tau\omega)\underline{s_a} = a(1 + j\tau\omega)\underline{e_a}$ donc $a\tau\frac{ds_a}{dt} + s_a(t) = a\left[\tau\frac{de_a}{dt} + e_a(t)\right]$.

Comme $e_a(t) = E$ pour $t > 0$, l'équation se simplifie en : $a\tau\frac{ds_a}{dt} + s(t) = aE$ et s'intègre en : $s_a(t) = aE + \lambda \exp\left(-\frac{t}{a\tau}\right)$.

On trouve la constante d'intégration λ avec la condition initiale :

CHAPITRE 12 – FILTRAGE LINÉAIRE

- le condensateur est initialement déchargé, la tension à ses bornes est nulle en $t = 0$. Il n'y a donc pas de différence de potentiel entre $e_a(t)$ et $s_a(t)$ en $t = 0$ (continuité de la tension aux bornes d'un condensateur). Ainsi $s_a(0^+) = E$.
 - on retrouve ce résultat avec le théorème de la valeur initiale, vu en cours de SI. Sachant que $p = j\omega$ et que la transformée de Laplace d'un échelon de hauteur E est E/p : $\lim_{t \rightarrow 0} s_a(t) = \lim_{p \rightarrow \infty} pS_a(p) = \lim_{p \rightarrow \infty} pH_a(p) \frac{E}{p} = H_a(\infty)E = E$.
- Ainsi $\lambda = -E(a - 1)$. Finalement : $s_a(t > 0) = aE - (a - 1)E \exp\left(-\frac{t}{a\tau}\right)$.



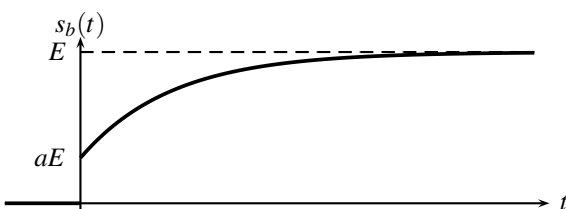
Avec, $(1 + j(R_1 + R_2)C\omega) \underline{s_b} = (1 + jR_2C\omega) \underline{e_b}$, l'équation différentielle régissant l'évolution du circuit (b) est : $s_b(t) + (R_1 + R_2)C \frac{ds_b}{dt} = e_b(t) + R_2C \frac{de_b}{dt}$.

Comme $e_b(t > 0) = E$, l'équation se simplifie en $s_b(t) + (R_1 + R_2)C \frac{ds_b}{dt} = E$ et s'intègre en : $s_b(t) = E + \mu \left(-\frac{t}{(R_1 + R_2)C} \right)$.

On trouve μ avec la condition initiale :

- le condensateur est initialement déchargé, donc la tension à ses bornes est nulle en $t = 0$. On a alors un pont diviseur instantané formé de R_1 et R_2 (il est qualifié d'instantané car il n'existe qu'en $t = 0$; après C a un effet) : $s_b(0^+) = \frac{R_2}{R_1 + R_2} e_b(0^+)$.
- Avec le théorème de la valeur initiale, vu en SI : $\lim_{t \rightarrow 0} s_b(t) = H_b(\infty)E = \frac{R_2}{R_1 + R_2} E$.

Ainsi $\mu = -E \frac{R_1}{R_1 + R_2}$ et : $s_b(t > 0) = E \left[1 - \frac{R_1}{R_1 + R_2} \exp\left(-\frac{t}{(R_1 + R_2)C}\right) \right]$.



L'équation différentielle régissant l'évolution du circuit (c) est trouvée avec la transmittance :

$(1 + j(R_1 + R_2)C\omega) \underline{s_c} = jR_2C\omega \underline{e_c}$ donc $(R_1 + R_2)C \frac{ds_c}{dt} + s_c(t) = R_2C \frac{de_c}{dt}$.

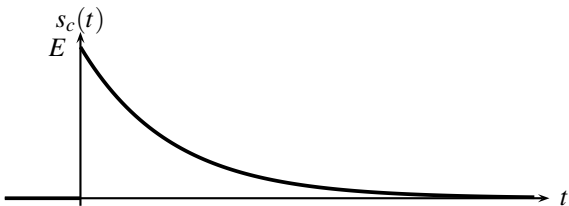
Elle devient sur $t > 0$ ($e_c(t > 0) = E$) : $(R_1 + R_2) \frac{ds_c}{dt} + s_c(t) = 0$ et s'intègre en : $s_c(t) =$

$$\mu \exp \left(-\frac{t}{(R_1 + R_2)C} \right).$$

On trouve la constante d'intégration μ avec la condition initiale :

- la tension aux bornes du condensateur est initialement nulle ; on a donc un pont diviseur résistif instantané : $s_c(0^+) = \frac{R_2}{R_1 + R_2} e_c(0^+) = \mu$
- Avec le théorème de la valeur initiale, vu en SI : $\lim_{t \rightarrow 0} s_c(t) = H_c(\infty) E = \frac{R_2}{R_1 + R_2} E$.

Finalement : $s_c(t) = E \frac{R_2}{R_1 + R_2} \exp \left(-\frac{t}{(R_1 + R_2)C} \right).$

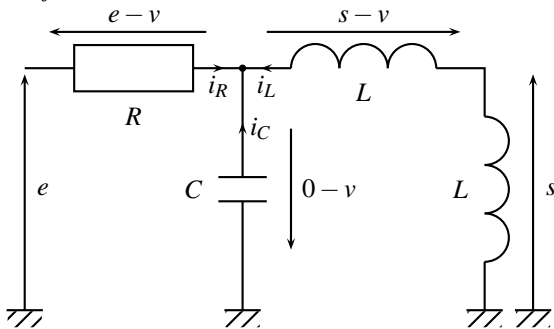


4. Dans chaque cas, le système laisse passer les hautes fréquences, car $H_{a,b \text{ ou } c}(\infty) \neq 0$. Les hautes fréquences présentes dans le signal d'entrée se retrouvent intégralement en sortie, qui est à son tour discontinue.

12.8 Filtre de Hartley

1. Soit v la tension aux bornes du condensateur. Les deux bobines forment un diviseur de tension entre s et v : $\frac{s}{v} = \frac{jL\omega}{jL\omega + jL\omega} = \frac{1}{2}$.

La somme des courants qui arrivent au nœud central (entre R , C et L) est nul. Le potentiel de ce nœud est v (car la tension v aux bornes de C est égale à la différence de potentiel $v - 0$). Alors, en notant bien les tensions en convention récepteur : $i_R + i_C + i_L = 0$ donc $\frac{e-v}{R} + jC\omega(0-v) + \frac{s-v}{jL\omega} = 0$.



Avec $v = 2s$, on aboutit à : $\frac{s}{e} = H(j\omega) = \frac{j\frac{L}{R}\omega}{1 + 2j\frac{L}{R}\omega + 2LC(j\omega)^2}.$

CHAPITRE 12 – FILTRAGE LINÉAIRE

On identifie les paramètres canoniques : $\frac{1}{\omega_0^2} = 2LC$ donc $\omega_0 = \frac{1}{\sqrt{2LC}}$, $\frac{2m}{\omega_0} = 2\frac{L}{R}$ donc $m =$

$\frac{1}{R}\sqrt{\frac{L}{2C}}$ et $H_0\frac{2m}{\omega_0} = \frac{L}{R}$ donc $H_0 = \frac{1}{2}$.

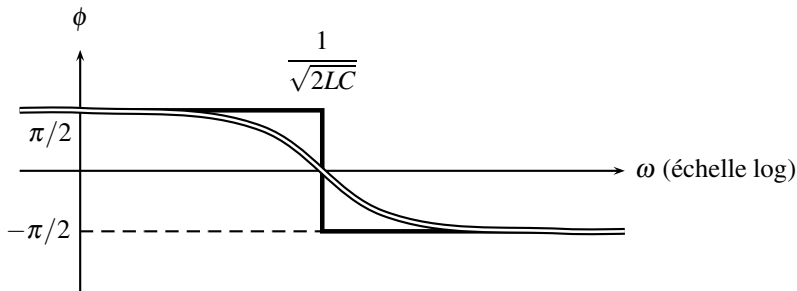
2. Le filtre est un passe-bande du deuxième ordre dont le dénominateur est simplifiable. Dressons un tableau asymptotique, où les pentes apparaissent :

	ω_0	
$1 + 2j\frac{L}{R}\omega + 2LC(j\omega)^2$	1	$2LC(j\omega)^2$
$H_{asympt}(j\omega)$	$j\frac{L}{R}\omega$	$\frac{1}{j2RC\omega}$
pente	+20 dB/dec	-20 dB/dec
phase	$+\pi/2$	$-\pi/2$

Le **gain asymptotique** β vaut en ω_0 : $|H_{asympt}(j\omega_0)| = \frac{1}{R}\sqrt{\frac{L}{2C}}$ donc $|H_{asympt}(j\omega_0)|_{dB} = -43\text{ dB}$

Quant au **gain réel** α en ω_0 : $H(j\omega_0) = \frac{1}{2}$ donc $|H(j\omega_0)|_{dB} = -6\text{ dB}$ et $\varphi(j\omega_0) = 0$.

Puis, avec les phases obtenues dans le tableau asymptotique :



3. Intégrateur pour $\omega \gg \omega_0$ et dérivateur pour $\omega \ll \omega_0$.

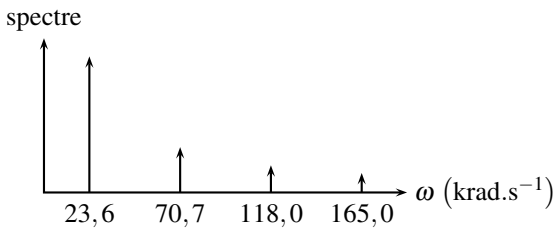
4. On utilise un GBF qui délivre un signal sinusoïdal accompagné d'un offset.

5. En régime permanent harmonique : $s_1(t) = E_0 H(j0) + E_1 e^{j\omega_1 t} H(j\omega_1)$.

Or $\omega_1 = \omega_0$ donc $H(j0) = 0$ et $H(j\omega_1) = \frac{1}{2}$: $s_1(t) = \frac{E_1}{2} \cos(\omega_1 t)$.

6. $E_{2eff}^2 = \langle e_2^2 \rangle = \frac{1}{T} \int_0^T e_2^2(t) dt = \frac{1}{T} \int_0^{T/2} E_{20}^2 dt + \frac{1}{T} \int_{T/2}^T (-E_{20})^2 dt = E_{20}^2$, donc $E_{2eff} = E_{20}$.

7. Le fondamental a une pulsation $\omega_2 = \frac{2\pi}{T} = \frac{1}{3\sqrt{2LC}} = \frac{\omega_0}{3} = 23,6\text{ krad.s}^{-1}$:



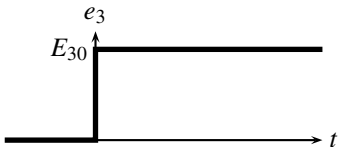
8. Le gain du filtre à la pulsation ω est : $|H(p = j\omega)| = \frac{\frac{L}{R}\omega}{\sqrt{(1 - 2LC\omega^2)^2 + (2\frac{L}{R}\omega)^2}}$.

Ainsi :

	amplitude de l'entrée	× gain du filtre	= amplitude de la sortie
fondamental ($\omega = \omega_0/3$)	$\frac{4E_{20}}{\pi} = 1,27$	× $2,65 \cdot 10^{-3}$	$= 3,38 \cdot 10^{-3}$
harmonique de rang 3 ($\omega = 3 \times \omega_0/3$)	$\frac{4E_{20}}{3\pi} = 0,42$	× 0,5	$= 2,12 \cdot 10^{-1}$
harmonique de rang 5 ($\omega = 5 \times \omega_0/3$)	$\frac{4E_{20}}{5\pi} = 0,25$	× $6,63 \cdot 10^{-3}$	$= 1,69 \cdot 10^{-3}$
harmonique de rang 7 ($\omega = 7 \times \omega_0/3$)	$\frac{4E_{20}}{7\pi} = 1,70 \cdot 10^{-2}$	× $3,71 \cdot 10^{-3}$	$= 6,75 \cdot 10^{-4}$

L'harmonique de rang 3 a une amplitude largement supérieure aux autres composantes (63 fois celle du fondamental, 125 fois celle de l'harmonique 5). On ne verra donc qu'elle dans le signal de sortie. Avec $\arg[H(3\omega_2 = \omega_0)] = 0$ car $H(3\omega_2 = \omega_0) = 1/2$: $s_2(t) = 2,12 \cdot 10^{-1} \cos(70,7 \cdot 10^3 t)$ (V).

9. Un essai indiciel sert à identifier un système inconnu. L'observation de la sortie permet de proposer une fonction de transfert pour le système.



10. Le système ne laisse pas passer les hautes fréquences, car $H(\infty) = 0$. Il n'y a donc en sortie aucune haute fréquence pour synthétiser une discontinuité. Ainsi, s_3 est continu en $t = 0$, c'est à dire $s_3(0^+)$. On retrouve ce résultat avec le théorème de la valeur initiale :

$$s_3(0^+) = \lim_{p \rightarrow \infty} pS(p) = \lim_{p \rightarrow \infty} pH(p) \frac{E_{30}}{p} = H(\infty)E_{30} = 0.$$

CHAPITRE 12 – FILTRAGE LINÉAIRE

La dérivée en $t = 0^+$ est calculable facilement avec le théorème de la valeur initiale : $\frac{ds_3}{dt}(0^+) =$

$$\lim_{p \rightarrow \infty} p [pS(p) - s_3(0^+)] = \lim_{p \rightarrow \infty} p^2 H(p) \frac{E_{30}}{p} = \frac{E_{30}}{2RC}.$$

11. L'équation différentielle à laquelle obéit le système provient de la transmittance :

$$(1 + 2j\frac{L}{R}\omega + 2LC(j\omega)^2)\underline{s} = j\frac{L}{R}\omega\underline{e}.$$

Et donc pour $t > 0$, où $e_3(t) =$ constante : $2LC\frac{d^2s_3}{dt^2} + 2\frac{L}{R}\frac{ds_3}{dt} + s_3 = \frac{L}{R}\frac{de_3}{dt} = 0$. Le discriminant de l'équation réduite associée est : $\Delta = 4\frac{L^2}{R^2} - 4 \times 2LC = 4.10^{-14} - 8.10^{-10} (s^2) = -8.10^{-10} s^2 < 0$.

On prendra donc dans la suite $\Delta = -8LC$. Les solutions de l'équation réduite associée sont

$$\text{donc : } r_{1,2} = \frac{-2\frac{L}{R} \pm \sqrt{\Delta}}{4LC} = -\frac{1}{2RC} \pm \frac{1}{\sqrt{2LC}} = -500 \pm 70,7.10^3 \text{ rad.s}^{-1}.$$

Ainsi : $s_3(t) = \exp\left(-\frac{t}{2RC}\right) \left[\lambda \cos\left(\frac{t}{\sqrt{2LC}}\right) + \mu \sin\left(\frac{t}{\sqrt{2LC}}\right) \right]$, soit numériquement $s_3(t) = e^{-500t} [\lambda \cos(70,7.10^3 t) + \mu \sin(70,7.10^3 t)]$ (V).

Les conditions initiales imposent : $s_3(0^+) = 0 = \lambda$ et $\frac{ds_3}{dt}(0^+) = \frac{E_{30}}{2RC} = \frac{\mu}{\sqrt{2LC}}$, donc $\mu =$

$$\frac{E_{30}}{R} \sqrt{\frac{L}{2C}} = 7,1.10^{-3} \text{ V}.$$

Finalement : $s_3(t) = \frac{E_{30}}{R} \sqrt{\frac{L}{2C}} \exp\left(-\frac{t}{2RC}\right) \sin\left(\frac{t}{\sqrt{2LC}}\right)$, numériquement :

$$s_3(t) = 7,1.10^{-3} e^{-500t} \sin(70,7.10^3 t) \text{ (V)}.$$

La pseudopulsation $\omega_0 = 70,7.10^3 \text{ rad.s}^{-1}$ est la seule fréquence qui passe à travers le filtre, d'après l'étude de la question 8. Il est donc normal de la retrouver en sortie.

12.9 Filtre passe-bas

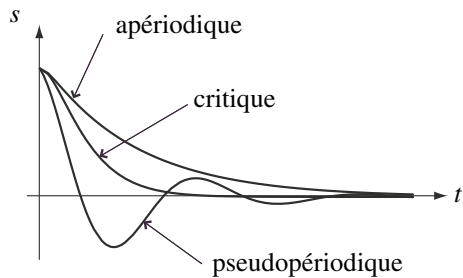
1. Par identification avec l'expression proposée dans l'énoncé, on trouve : $\omega_0 = \frac{1}{R\sqrt{C_1C_2}}$,

$$m = \frac{3}{2} \sqrt{\frac{C_1}{C_2}} \text{ et } G_0 = -1.$$

2. On utilise le passage de la fonction de transfert à l'équation différentielle grâce à l'équivalence entre la multiplication par $j\omega$ en notation complexe et la dérivation par rapport au temps en notation réelle. On en déduit : $\frac{1}{\omega_0^2} \frac{d^2s}{dt^2} + \frac{2m}{\omega_0} \frac{ds}{dt} + s = G_0 e$.

3. La solution générale de l'équation différentielle homogène associée décrit le régime libre. L'équation caractéristique associée $\frac{1}{\omega_0^2} r^2 + \frac{2m}{\omega_0} r + 1 = 0$ admet pour discriminant réduit $\Delta' =$

$\frac{m^2 - 1}{\omega_0^2}$. On a donc trois régimes possibles pour le régime libre : apériodique si $\Delta' > 0$ ou $m > 1$, critique si $\Delta' = 0$ ou $m = 1$ et pseudopériodique si $\Delta' < 0$ ou $m < 1$. Les allures sont les suivantes :



4. Le système est un passe-bas donc ne laisse pas passer les hautes fréquences, qui ne peuvent donc pas reconstituer de discontinuité en sortie : $s(0^+) = 0$. On retrouve ce résultat avec le théorème de la valeur initiale, vu en SI, avec une netrée en échelon.

$t \rightarrow \infty$ correspond au régime permanent qui sera décrit par une solution particulière constante de l'équation différentielle. On obtient : $s(+\infty) = G_0 E_0$.

5. On rappelle que $m = \frac{3}{2} \sqrt{\frac{C_1}{C_2}}$. Si on veut $m = \frac{\sqrt{2}}{2} < 1$, il faut donc choisir C_2 tel que $C_2 = \frac{9}{2} C_1$. Le régime libre est alors apériodique d'après ce qui a été fait auparavant.

6. On souhaite $10^3 \text{ Hz} \leq f_0 = \frac{\sqrt{2}}{6\pi RC_1} \leq 4 \cdot 10^3 \text{ Hz}$, ce qui impose $\frac{\sqrt{2}}{24\pi 10^3 C_1} \leq R \leq \frac{\sqrt{2}}{6\pi 10^3 C_1}$.

7. Il n'y a qu'une pulsation caractéristique ω_0 dans ce système. Simplifions la fonction de transfert dans un tableau asymptotique :

	ω_0	
$1 + 3jRC_1\omega + R^2C_1C_2(j\omega)^2$	1	$R^2C_1C_2(j\omega)^2$
$H_{asymp}(j\omega)$	-1	$-\frac{1}{R^2C_1C_2(j\omega)^2}$
pente	nulle	-40 dB/dec
phase	$-\pi$	-2π

On n'oublie pas que le -1 dans la fonction de transfert, ajoute une phase de $-\pi$.
le tableau asymptotique répond à la question.

8. On a un comportement intégrateur lorsque la fonction de transfert peut s'identifier à la forme : $\frac{1}{j\frac{\omega}{\omega_0}}$, ce qui n'est pas le cas.

Pour avoir un comportement dérivateur, il faudrait que la fonction de transfert puisse s'identifier à la forme $j\frac{\omega}{\omega_0}$, ce qui n'est jamais le cas ici.

9. a. Comme $1,5 < 2$, la fréquence du signal est dans la bande passante donc $s = -e$ c'est-à-dire un déphasage de π et une amplitude identique.

b. Pour le fondamental, on a les mêmes résultats qu'à la question précédente. Pour la première harmonique obtenue pour $n = 3$, son amplitude vaut 0,11a et le gain à cette fréquence est également de 0,11 : l'amplitude en sortie du filtre est donc de 1% par rapport à celle du fondamental. On peut la négliger et le signal en sortie est une sinusoïde en opposition de phase par rapport au créneau.

CHAPITRE 12 – FILTRAGE LINÉAIRE

c. On effectue le même raisonnement et on obtient le même résultat qu’à la question précédente.

10. À l’ordre 2, la pente de l’asymptote obtenue pour $\omega \gg \omega_0$ est de -40 dB/déc : on a donc une diminution plus rapide des composantes de hautes fréquences qu’on souhaite supprimer.

12.10 Filtrage d’un signal électrique

1. a. Il s’agit d’un filtre passe-bas du second ordre.

b. Sous forme canonique $\frac{H_0}{1 + j2\xi \frac{\omega}{\omega_0} - \left(\frac{\omega}{\omega_0}\right)^2}$, avec $H_0 = -1$, $\omega_0 = \frac{1}{R\sqrt{C_1C_2}}$ et $\xi = \frac{3}{2}\sqrt{\frac{C_1}{C_2}}$

c. On en déduit les valeurs de C_1 et C_2 à partir des expressions de ω_0 et $Q = \frac{1}{2\xi} = \frac{1}{3}\sqrt{\frac{C_2}{C_1}}$ soit $C_2 = \frac{3Q}{R\omega_0}$ et $C_1 = \frac{1}{3RQ\omega_0}$.

d. À partir de l’expression de ω_0 , on en déduit l’incertitude relative $\frac{\Delta\omega_0}{\omega_0} = \frac{\Delta R}{R} + \frac{\Delta C}{C} = \frac{\Delta C}{C}$ du fait de la précision infinie sur R et de la précision sur C à 5 %.

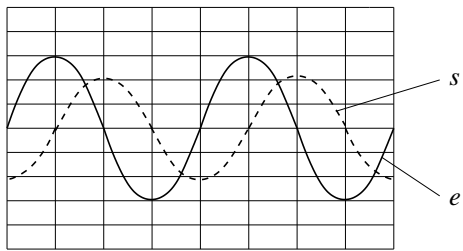
e. L’application numérique donne $C_2 = 16,9 \text{ nF}$ et $C_1 = 3,75 \text{ nF}$.

2. a. Pour réaliser le tracé du diagramme de Bode, on visualise les deux voies en fonction du temps, on mesure l’amplitude (en tenant compte du calibre et en mesurant le nombre de carreaux ou en utilisant les outils intégrés) et du déphasage.

b. Une période correspond à 4 carreaux, chaque carreau représentant $12,5 \mu\text{s}$. On en déduit que la période est $T = 50 \mu\text{s}$ et la fréquence $f = \frac{1}{T} = 20 \text{ kHz}$. Quant à l’amplitude d’entrée, elle est de trois carreaux correspondant chacun à 1,0 V donc l’amplitude est de 3,0 V.

c. Pour une fréquence $f = f_0$, la fonction de transfert s’écrit $\underline{H} = \frac{H_0}{2j\xi} = jQ$ car $H_0 = -1$. En faisant l’application numérique, on en déduit la valeur du gain 0,707 et du déphasage $\frac{\pi}{2}$.

d. Par conséquent, l’amplitude de sortie vaut 2,12 V et on obtient l’allure suivante :



e. Pour observer le diagramme de Bode avec ce dispositif à l'aide d'un créneau, il faut prendre un signal de très faible fréquence afin que les harmoniques intéressantes pour le filtre (au voisinage de 20 kHz) soient à peu près constantes d'après le rappel de l'énoncé. On envoie ce signal à l'entrée du filtre et on fait l'analyse de Fourier de la sortie : le spectre de la sortie pour les harmoniques élevées donne l'allure du gain.

12.11 Étude et utilisation de filtres, d'après ENSAIT 2002

1. En appliquant la relation du pont diviseur de tension, on obtient $\underline{H} = \frac{jRC\omega}{1 + jRC\omega}$.
2. Il s'agit d'un filtre passe-haut du premier ordre.
3. Le gain s'écrit $G = |\underline{H}| = \frac{1}{\sqrt{1 + \frac{1}{(RC\omega)^2}}}$. La dérivée de $f(\omega) = 1 + \frac{1}{R^2C^2\omega^2}$ est $f'(\omega) = -\frac{2}{R^2C^2\omega^3} < 0$. On en déduit que f est décroissante et que le gain G est une fonction croissante.
4. Une seule pulsation caractéristique intervient dans ce système, $\omega_0 = \frac{1}{RC}$. Simplifions la transmittance avec un tableau asymptotique :

	ω_0	
$1 + jRC\omega$	1	$jRC\omega$
$H_{asympt}(j\omega)$	$jRC\omega$	1
pente	+20 dB/dec	nulle
phase	$+\pi/2$	0

On en conclut que le diagramme est le (C). La pulsation de cassure est $\omega_0 = \frac{1}{RC}$ et les pentes sont indiquées dans le tableau.

5. Pour obtenir la valeur de C cherchée, on doit effectuer la résolution de :
 $-20\log \sqrt{1 + \frac{1}{R^2C^2\omega^2}} = -6$. On obtient $C = \frac{1}{2\pi Rf\sqrt{10^{0,6} - 1}} = 9,2 \cdot 10^{-9}$ F.
6. Pour le montage 2, on applique à nouveau la relation du pont diviseur de tension et on obtient $\underline{H} = \frac{\underline{Z}_{\parallel}}{\underline{Z}_s + \underline{Z}_{\parallel}}$ avec $\underline{Z}_s = \frac{1 + jRC\omega}{jC\omega}$ et $\underline{Z}_{\parallel} = \frac{R}{1 + jRC\omega}$. En explicitant les valeurs des impédances dans la fonction de transfert, on en déduit $\underline{H} = \frac{jRC\omega}{1 + 3jRC\omega - R^2C^2\omega^2}$.
7. Il s'agit d'un filtre passe-bande du second ordre.

8. La fonction de transfert peut s'écrire sous la forme $\underline{H} = H_0 \frac{2j\xi \frac{\omega}{\omega_0}}{1 + 2j\xi \frac{\omega}{\omega_0} + \left(j \frac{\omega}{\omega_0}\right)^2}$ avec
 $\omega_0 = \frac{1}{RC}$, $\xi = \frac{1}{2}$ et $H_0 = \frac{1}{3}$.

CHAPITRE 12 – FILTRAGE LINÉAIRE

9. La pulsation caractéristique qui apparaît au dénominateur est $\omega_0 = \frac{1}{RC}$. Simplifions la transmittance dans un tableau asymptotique :

	ω_0	
	1	$-R^2C^2\omega^2$
$H_{asymp}(j\omega)$	$jRC\omega$	$\frac{1}{jRC\omega}$
pente	$+20\text{ dB/dec}$	-20 dB/dec
phase	$+\pi/2$	$-\pi/2$

La lecture du tableau répond à la question.

10. L'application numérique donne $f_0 = \frac{1}{2\pi RC} = 1,7\text{ kHz}$.

11. Pour une fréquence $f = 1,0\text{ kHz}$, on a un gain de $G = 0,30$ pour le fondamental (fréquence f), $G = 0,33$ pour l'harmonique d'ordre 2, $G = 0,30$ pour celle d'ordre 3, $G = 0,28$ pour celle d'ordre 4, $G = 0,25$ pour celle d'ordre 5, $G = 0,22$ pour celle d'ordre 6, $G = 0,20$ pour celle d'ordre 7, $G = 0,18$ pour celle d'ordre 8, $G = 0,17$ pour celle d'ordre 9 et $G = 0,15$ pour celle d'ordre 10.

12. Le signal de sortie est quasi-identique à celui d'entrée avec une amplitude divisée approximativement par 3 puisque le gain est sensiblement égal pour toutes les premières harmoniques.

13. Pour un circuit R, L, C série, on a $\underline{i} = \frac{\underline{e}}{R + j\left(L\omega - \frac{1}{C\omega}\right)}$. La résonance en intensité se

produit pour la fréquence $f_0 = \frac{1}{2\pi\sqrt{LC}} = 1,0\text{ MHz}$.

14. Pour le circuit R, L, C série, l'impédance est $\underline{Z}_s = R + j\left(L\omega - \frac{1}{C\omega}\right)$. Pour $\omega = \omega_0$, on a $\underline{Z}_s = R = 1,0 \cdot 10^4\ \Omega$.

Pour le circuit proposé, on a l'impédance $\underline{Z}_{||} = \frac{1}{jC\omega + \frac{1}{R + jL\omega}} = \frac{1 + jL\omega}{1 - LC\omega^2 + jRC\omega} =$

$\frac{L}{RC} - j\sqrt{\frac{L}{C}} = 250 \cdot 10^6 - 1,58 \cdot 10^6\ \Omega$ soit $\underline{Z}_{||} \simeq \frac{L}{RC} = 250 \cdot 10^6\ \Omega$ pour $\omega = \omega_0$.

15. On applique la relation du pont diviseur de tension et $\underline{s} = \frac{\underline{Z}_{||}}{\underline{Z}_s + \underline{Z}_{||}} \underline{e}$ soit compte tenu des valeurs numériques un gain de quasi 1 : on a une amplitude de 0,5 V et un déphasage quas nul.

Deuxième partie

Mécanique 1

Trouver la bibliothèque des livres ici : <https://1024terabox.com/s/1mg0wxl22G8NjM8MA2RzgFw>

Cinématique du point

13

La cinématique est une partie de la mécanique qui s'intéresse à l'étude des mouvements, indépendamment des causes qui les ont produits. Dans ce chapitre, on fixe le cadre de l'étude mécanique, on définit la notion de point en physique et on décrit les outils qui permettent de décrire géométriquement sa position, sa vitesse et son accélération au cours du temps.

1 Notion de point en physique

Aucun objet physique n'est rigoureusement ponctuel au sens mathématique du terme. Pour étudier la mécanique du point, il faut donc préciser dans quelles conditions un objet mécanique peut être modélisé par un point.

1.1 Définition d'un solide

Un solide est système matériel dont les points restent à distance constante les uns des autres. Pour repérer un solide dans l'espace, il faut six paramètres :

- les trois coordonnées d'un point du solide (on considère souvent la position de son centre de gravité G) ;
- trois angles qui définissent l'orientation d'un repère lié au solide par rapport au référentiel d'étude (voir figure 13.1).

Dans le cadre de la cinématique, les solides sont supposés indéformables. On exclut donc toute déformation et toute rupture du solide de cette étude.

1.2 Définition d'un point

Un point matériel est un solide dont on peut négliger l'extension spatiale et la rotation sur lui-même. Pour repérer complètement un point dans l'espace, il suffit de donner trois paramètres : ses trois coordonnées.

En dynamique, on a besoin de définir sa masse. Cela mène à la notion de point matériel. Ce n'est pas nécessaire en cinématique.

CHAPITRE 13 – CINÉMATIQUE DU POINT

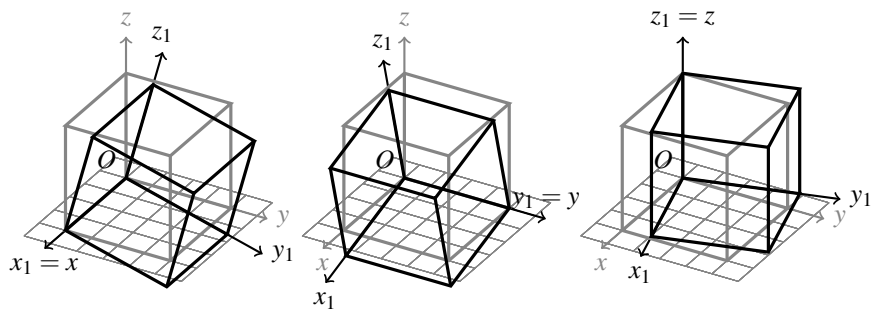


Figure 13.1 – Cube de sommet O fixe repéré dans l'espace. À gauche, le cube a tourné autour l'axe (Ox) ; au centre, autour de (Oy) et à droite, autour de (Oz) . Le cas général est composé de ces trois rotations. Il faut donc trois angles pour déterminer l'orientation d'un cube dans l'espace.

1.3 Quand peut-on assimiler un système à un point ?

La question n'a pas de réponse simple et on doit se référer à la définition : tout dépend si on peut négliger son extension spatiale et sa rotation sur lui-même.

Exemple

Le mouvement orbital de la Terre autour du Soleil dans le référentiel de Copernic s'inscrit sur une trajectoire quasi-circulaire de centre confondu avec le centre du Soleil et de rayon orbital $d_{ST} = 150.10^6$ km. L'extension spatiale de la Terre (son rayon $R_T = 6400$ km) est négligeable devant le rayon orbital : $R_T \ll d_{ST}$. On peut négliger l'extension de la Terre et sa rotation sur elle-même et l'assimiler à un point matériel lors de l'étude de son mouvement orbital.

Par contre, si l'on étudie le mouvement de rotation propre de la Terre autour de son axe Nord-Sud dans le référentiel géocentrique pour connaître la durée du jour, on ne peut négliger ni l'extension de la Terre, ni sa rotation sur elle-même. On ne peut pas assimiler la Terre à un point matériel lors de l'étude de sa rotation propre.

2 Repérage d'un point du plan

2.1 Intérêt d'avoir plusieurs systèmes de coordonnées

Lors d'expériences de mécanique, on repère le mouvement d'un point au cours du temps. La trajectoire suivie par le mobile peut prendre des formes très différentes. Si l'on considère trois mouvements plans vus au lycée, différents paramétrages géométriques sont utilisés, qui correspondent à différents systèmes de coordonnées.

Exemple

Lors d'une chute libre, la trajectoire suivie par le mobile est une portion de droite. On repère son mouvement sur un axe à l'aide d'un repérage cartésien réduit à un axe.

On lâche un petit objet sans vitesse initiale depuis l'origine du repère à $t_0 = 0$ s dans le champ de pesanteur terrestre. Il chute verticalement sur une trajectoire rectiligne dirigée vers le bas et à l'instant t , on le repère par son abscisse $x(t)$ sur un axe gradué. Le mouvement est rectiligne accéléré.

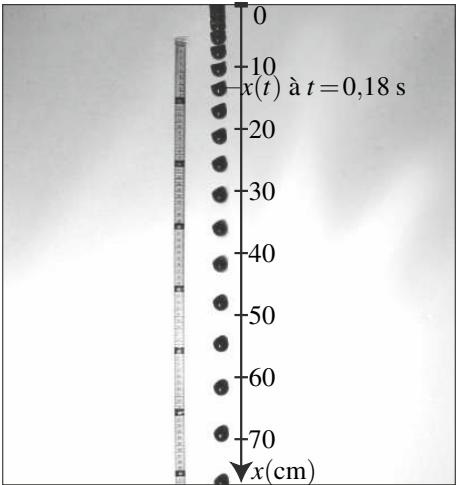


Figure 13.2 – Chronophotographie d'un petit objet en chute libre. Le mètre permet d'obtenir l'échelle. Fréquence de prise de vue : 50 images par seconde. Vitesse d'obturation : (1/1000)s.

Exemple

Lors d'un vol balistique, la trajectoire suivie par le mobile est une parabole. On repère son mouvement dans un plan à l'aide d'un repérage cartésien à deux dimensions.

On lance un petit objet avec une vitesse initiale depuis l'origine du repère à $t_0 = 0$ s dans le champ de pesanteur terrestre. La vitesse initiale de norme $\simeq 2,5 \text{ m}\cdot\text{s}^{-1}$ fait un angle $\simeq 60^\circ$ sur l'horizontale. L'objet effectue un vol balistique suivant une trajectoire plane. À l'instant t , on repère sa position par ses coordonnées $x(t)$ et $y(t)$. C'est un repérage cartésien de la position.

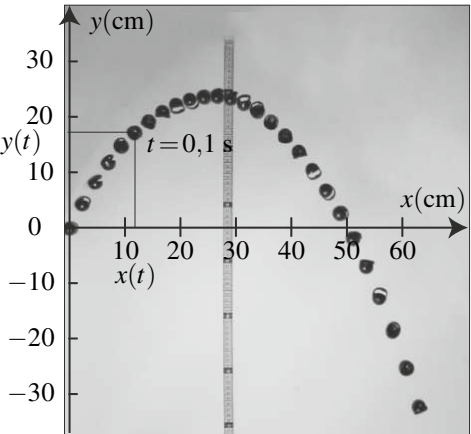


Figure 13.3 – Chronophotographie d'un petit objet en vol balistique. Le mètre permet d'obtenir l'échelle. Fréquence de prise de vue : 50 images par seconde. Vitesse d'obturation : (1/1000)s.

CHAPITRE 13 – CINÉMATIQUE DU POINT

Exemple

Lors d'un mouvement circulaire et uniforme, la trajectoire suivie par le mobile est un cercle. On repère son mouvement sur un cercle à l'aide d'un repérage polaire.

Un petit objet est posé sur un plateau à une distance $r = 30\text{ cm}$ de l'axe de rotation du plateau qui tourne à vitesse angulaire constante dans le sens trigonométrique. La position à l'instant t est repérée par la donnée de l'angle $\theta(t)$. C'est un repérage polaire de la position.

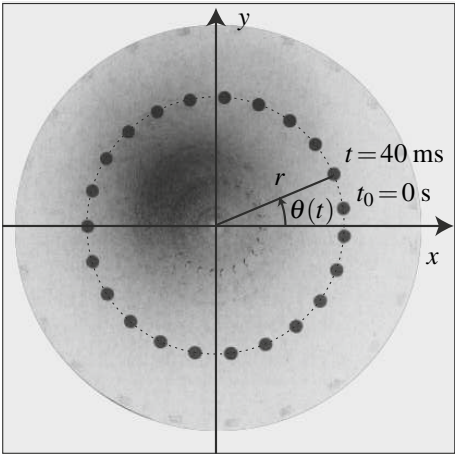


Figure 13.4 – Chronophotographie d'un petit objet en mouvement circulaire et uniforme. Fréquence de prise de vue : 25 images par seconde. Vitesse d'obturation : $(1/500)\text{s}$.

2.2 Repérage d'un point sur une droite.

Ce système de coordonnées est utilisé pour repérer un point M sur une droite fixe $\mathcal{D}$. Il est adapté à l'étude des mouvements rectilignes.

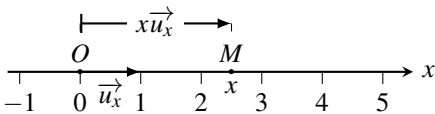


Figure 13.5 – Repérage d'un point sur une droite

On place un point O et un vecteur unitaire $\vec{u}_x$ sur la droite $\mathcal{D}$. Le point M est repéré sur la droite orientée $(O, \vec{u}_x)$ par la valeur algébrique $\overline{OM} = x$, appelée coordonnée de M .

- $x > 0$ si $\overline{OM}$ est dans le sens de $\vec{u}_x$ donc si M est à droite de O sur le schéma ci-dessus.
- $x < 0$ si $\overline{OM}$ est dans le sens opposé à $\vec{u}_x$ donc si M est à gauche de O .

Sur la droite (Ox) , un point M de coordonnée x est noté $M(x)$. Son vecteur position s'écrit :

$$\overrightarrow{OM} = x\vec{u}_x.$$

2.3 Repérage d'un point dans le plan

Pour repérer un point M dans un plan fixe $\mathcal{P}$, on utilise les systèmes de coordonnées cartésien ou polaire. Ces systèmes de coordonnées sont adaptés à l'étude des mouvements plans. Suivant les cas, on choisit un repérage cartésien ou un repérage polaire. Le repérage polaire est particulièrement indiqué pour les mouvements circulaires.

a) Coordonnées cartésiennes

Description Le repérage cartésien consiste à quadriller le plan à l'aide d'une grille carrée. Les points du plan sont repérés par leurs coordonnées sur deux axes orthogonaux comme sur la figure 13.6.

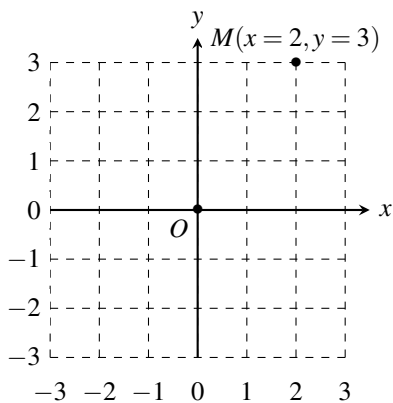


Figure 13.6 – Quadrillage cartésien du plan. M a pour coordonnées cartésiennes $x = 2$ et $y = 3$.

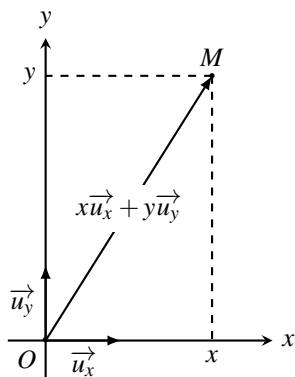


Figure 13.7 – Repérage d'un point M du plan (xOy) en coordonnées cartésiennes.

Système de coordonnées : repérage d'un point du plan On place un point O et deux vecteurs fixes $\vec{u}_x$ et $\vec{u}_y$ dans le plan $\mathcal{P}$ tels que le doublet $(\vec{u}_x, \vec{u}_y)$ forme une base orthonormée directe de ce plan. Le point M est repéré par deux paramètres géométriques x et y où :

- x est la coordonnée du projeté orthogonal de M sur l'axe orienté $(O, \vec{u}_x)$;
- y est la coordonnée du projeté orthogonal de M sur l'axe orienté $(O, \vec{u}_y)$.

Dans le plan (Oxy) , un point M de **coordonnées cartésiennes** x et y est noté $M(x, y)$. Son vecteur position s'écrit :

$$\overrightarrow{OM} = x\vec{u}_x + y\vec{u}_y,$$

où x et y sont les projections orthogonales de $\overrightarrow{OM}$ sur les vecteurs $\vec{u}_x$ et $\vec{u}_y$.

CHAPITRE 13 – CINÉMATIQUE DU POINT

Base de projection : repérage d'un vecteur du plan Le doublet $(\vec{u}_x, \vec{u}_y)$ forme une **base de projection fixe** du plan $\mathcal{P} = (Oxy)$. Tout vecteur $\vec{w}$ du plan s'écrit de manière unique sous la forme :

$$\vec{w} = w_x \vec{u}_x + w_y \vec{u}_y \quad \text{noté} \quad \vec{w} = \begin{pmatrix} w_x \\ w_y \end{pmatrix}.$$

w_x et w_y sont les composantes du vecteur $\vec{w}$ sur $\vec{u}_x$ et $\vec{u}_y$. Ce sont les projections orthogonales de $\vec{w}$ sur les vecteurs $\vec{u}_x$ et $\vec{u}_y$. En particulier, le vecteur position $\vec{OM} = x\vec{u}_x + y\vec{u}_y$ est noté :

$$\vec{OM} = \begin{pmatrix} x \\ y \end{pmatrix}.$$

Le repérage cartésien est le plus naturel dans le cas des mouvements plans, mais il n'est pas toujours le plus pertinent. Il faut avoir à l'esprit le système de coordonnées polaires, notamment dans le cas des mouvements circulaires.

b) Coordonnées polaires

Description Le repérage polaire consiste à quadriller le même plan $\mathcal{P}$ à l'aide d'une grille centrée sur un point O du plan. Les points du plan sont alors repérés par leur distance à O et un angle mesuré par rapport à un axe fixe comme sur la figure ci-dessous.

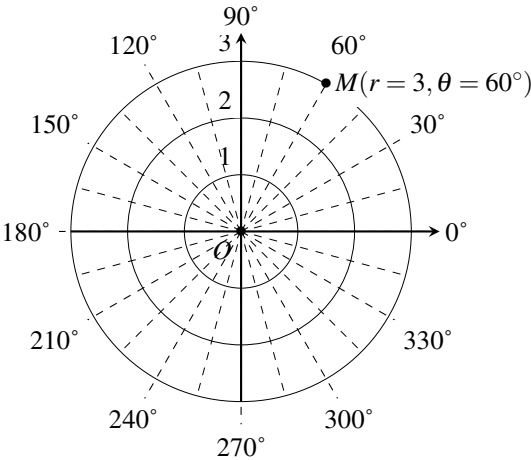


Figure 13.8 – Quadrillage du plan en coordonnées polaires. Le point M a pour coordonnées polaires $r = 3$ et $\theta = 60^\circ$.

Système de coordonnées : repérage d'un point du plan En général, on utilise le point O et le vecteur $\vec{u}_x$ des coordonnées cartésiennes pour définir la droite orientée (Ox) prise comme référence des angles. Le point M est repéré par ses coordonnées polaires r et θ où :

- θ est l'angle orienté que fait le vecteur $\vec{OM}$ avec le vecteur $\vec{u}_x$;
- r est la distance OM . C'est la norme du vecteur $\vec{OM}$: $r = ||\vec{OM}||$.

Un point M de coordonnées polaires r et θ est noté $M(r, \theta)$ et il faut avoir en tête que :

- par définition, r étant une distance, elle est à valeur positive ;
- pour décrire l'intégralité du plan, θ doit varier sur un intervalle de largeur 2π radian, le plus souvent $\theta \in [0; 2\pi[$.

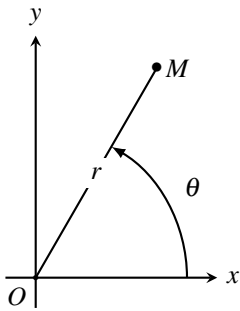


Figure 13.9 – Repérage d'un point M du plan en polaires.

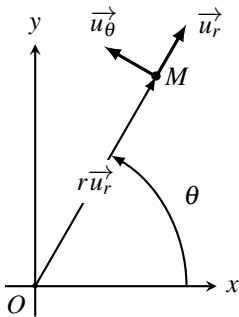


Figure 13.10 – Base polaire locale pour un point M du plan.

Base de projection : repérage d'un vecteur du plan On définit la **base de projection mobile** $(\vec{u}_r, \vec{u}_\theta)$ orthonormée directe adaptée aux coordonnées polaires de la manière suivante :

- $\vec{u}_r$ est un vecteur unitaire dont la direction et le sens sont ceux de $\overrightarrow{OM}$:

$$\vec{u}_r = \frac{\overrightarrow{OM}}{\|\overrightarrow{OM}\|};$$

- $\vec{u}_\theta$ est obtenu en faisant tourner $\vec{u}_r$ d'un angle de $\frac{\pi}{2}$ rad.

Dans le plan (Oxy) , un point M de **coordonnées polaires** r et θ est noté $M(r, \theta)$. Le vecteur position s'écrit :

$$\overrightarrow{OM} = r\vec{u}_r,$$

où r est la distance OM .

La base de projection $(\vec{u}_r, \vec{u}_\theta)$ est définie pour une position donnée du point M . Elle est donc définie localement en tout point M et tourne pour que le vecteur $\vec{u}_r$ pointe toujours de O vers M . Pour ces raisons, la base $(\vec{u}_r, \vec{u}_\theta)$ est appelée **base locale** ou **base mobile**.

Remarque

Le lecteur attentif aura noté que les vecteurs de base $(\vec{u}_r, \vec{u}_\theta)$ sont définis en un point M donné et que M doit être différent de O . Cependant, étant donné que les vecteurs ne sont pas attachés à un point, ils sont parfois reportés de M à O pour améliorer la clarté des figures.

CHAPITRE 13 – CINÉMATIQUE DU POINT

Pour un point M donné, le doublet $(\vec{u}_r, \vec{u}_\theta)$ forme une base mobile du plan $\mathcal{P}$. Tout vecteur $\vec{w}$ du plan peut s’écrire sous la forme :

$$\vec{w} = w_r \vec{u}_r + w_\theta \vec{u}_\theta \quad \text{noté} \quad \vec{w} = \begin{pmatrix} w_r \\ w_\theta \end{pmatrix}.$$

- w_r est la composante du vecteur $\vec{w}$ sur $\vec{u}_r$ appelée **composante radiale** ;
- w_θ est la composante du vecteur $\vec{w}$ sur $\vec{u}_\theta$ appelée **composante orthoradiale**.

w_r et w_θ sont les projections orthogonales de $\vec{w}$ sur les vecteurs $\vec{u}_r$ et $\vec{u}_\theta$. En particulier, le vecteur position $\vec{OM} = r\vec{u}_r$ s’écrit :

$$\vec{OM} = \begin{pmatrix} r \\ 0 \end{pmatrix}.$$

! Il ne faut pas confondre les coordonnées du point M et les composantes du vecteur position. Ces deux notions mathématiques, qui coïncident en coordonnées cartésiennes, ne sont plus du tout les mêmes en coordonnées cylindriques.

3 Repérage d’un point dans l’espace

3.1 Repérage cartésien

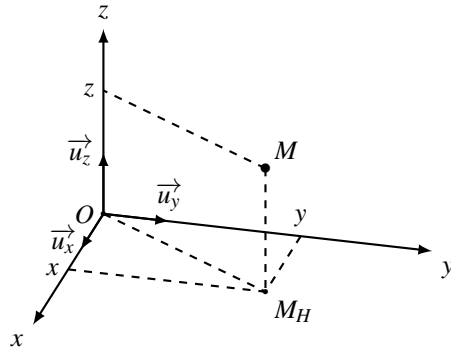


Figure 13.11 – Repérage d’un point M dans l’espace en coordonnées cartésiennes. Le projeté de M sur le plan (Oxy) noté M_H est repéré en coordonnées cartésiennes planes. M est alors à une « altitude » z de M_H selon l’axe (Oz) .

La définition des coordonnées cartésiennes utilisée dans le plan peut être étendue à un repérage dans l’espace à trois dimensions en rajoutant un axe (Oz) orienté orthogonal au plan (Oxy) . L’orientation de (Oz) est celle du vecteur $\vec{u}_z$ choisi pour que le triplet $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ soit une base orthonormée directe (voir annexe mathématique). Le vecteur position $\vec{OM}$ vaut alors :

$$\vec{OM} = x\vec{u}_x + y\vec{u}_y + z\vec{u}_z \quad \text{noté} \quad \vec{OM} = \begin{pmatrix} x \\ y \\ z \end{pmatrix}.$$

3.2 Repérage cylindrique

La définition des coordonnées polaires utilisée dans le plan peut être étendue à un repérage dans l'espace à trois dimensions en rajoutant un axe (Oz) orthogonal au plan (Oxy) . On obtient alors le **système de coordonnées cylindriques**. L'orientation de (Oz) est celle du vecteur $\vec{u}_z$ choisi pour que le triplet $(\vec{u}_r, \vec{u}_\theta, \vec{u}_z)$ soit une base orthonormée directe. Le vecteur position $\vec{OM}$ vaut alors :

$$\vec{OM} = r\vec{u}_r + z\vec{u}_z \quad \text{noté} \quad \vec{OM} = \begin{pmatrix} r \\ 0 \\ z \end{pmatrix}.$$

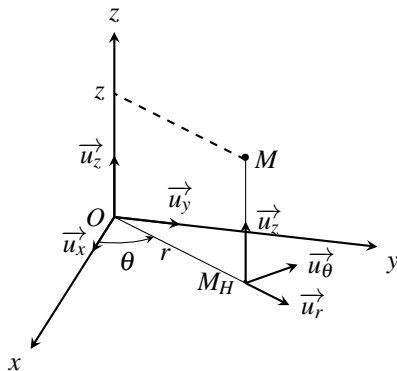


Figure 13.12 – Repérage d'un point M dans l'espace en coordonnées cylindriques.
Le projeté de M sur le plan (Oxy) , noté M_H , est repéré en coordonnées polaires planes. M est alors à une « altitude » z de M_H selon le vecteur $\vec{u}_z$.

3.3 Repérage sphérique

Description Le repérage sphérique d'un point M consiste à repérer le point par la distance OM , puis sa position sur la sphère de rayon OM par deux angles (voir figure 13.13). Comme il s'agit d'un repérage courant en géographie, le vocabulaire qui suit est emprunté aux géographes.

On appelle équateur le cercle marquant l'intersection de la sphère avec le plan (Oxy) et plan équatorial le plan de l'équateur. Un parallèle est un cercle de la sphère parallèle à l'équateur. On repère un parallèle par sa latitude.

L'axe (Oz) est l'axe Nord/Sud. Un méridien est un cercle qui contient l'axe Nord/Sud de la sphère. On repère un méridien par sa longitude. Un plan qui contient un méridien est un plan méridien. Il contient également l'axe Nord/Sud.

Le point M est situé à l'intersection d'un parallèle et d'un méridien. On le repère par sa longitude et sa latitude.

CHAPITRE 13 – CINÉMATIQUE DU POINT

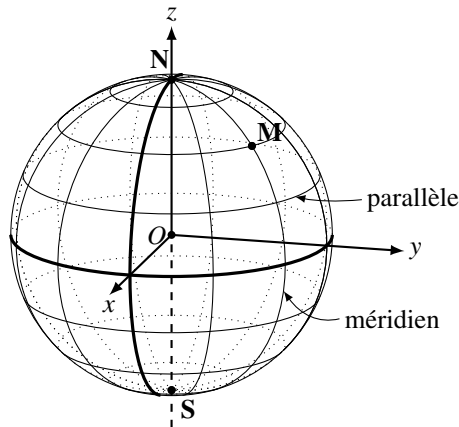


Figure 13.13 – Quadrillage d’une sphère de rayon OM en coordonnées sphériques. Les méridiens sont tracés tous les 30° de longitude. Les parallèles tracés ont pour des latitudes $\pm 15^\circ$, $\pm 45^\circ$, $\pm 75^\circ$. Le point M a pour coordonnées 60° de longitude et 45° de latitude.

Système de coordonnées : repérage d’un point de l’espace En général, on utilise le méridien du plan (Oxz) des coordonnées cartésiennes comme référence des longitudes φ . On délaisse la latitude λ au profit de la co-latitude $\theta = \frac{\pi}{2} - \lambda$.

Le point M est repéré par ses coordonnées sphériques r , θ et φ où :

- r est la distance OM . C’est la norme du vecteur $\overrightarrow{OM}$: $r = ||\overrightarrow{OM}||$;
- la longitude φ permet de repérer le plan méridien. C’est l’angle orienté que fait le plan méridien passant par M avec le plan méridien (Oxz) pris comme référence des longitudes ;
- la co-latitude θ permet de repérer le point M dans le plan méridien. C’est l’angle orienté que fait le vecteur $\overrightarrow{OM}$ avec le vecteur $\vec{u}_z$.

Un point M de coordonnées sphériques r , θ et φ est noté $M(r, \theta, \varphi)$ et il faut noter que :

- par définition, r étant une distance, elle est à valeur positive
- pour décrire l’intégralité de l’espace, θ doit varier sur l’intervalle $[0; \pi]$ et φ sur l’intervalle $[0; 2\pi[$.



Il faut veiller à toujours prendre $\theta \in [0; \pi]$ et $\varphi \in [0; 2\pi[$. On justifiera ce choix plus tard mais il faut dès à présent en avoir conscience pour éviter des problèmes de signe.

Base de projection : repérage d’un vecteur On définit la base mobile $(\vec{u}_r, \vec{u}_\theta, \vec{u}_\varphi)$ orthonormée directe adaptée aux coordonnées sphériques de la manière suivante :

- $\vec{u}_r$ est un vecteur unitaire dont la direction et le sens sont ceux de $\overrightarrow{OM}$: $\vec{u}_r = \frac{\overrightarrow{OM}}{||\overrightarrow{OM}||}$;
- $\vec{u}_\theta$ un vecteur unitaire du plan méridien orthogonal à $\vec{u}_r$ donc tangent au méridien passant par M et dirigé dans le sens où θ augmente ;
- $\vec{u}_\varphi$ un vecteur unitaire tel que la base $(\vec{u}_r, \vec{u}_\theta, \vec{u}_\varphi)$ soit orthonormée directe. $\vec{u}_\varphi$ est perpendiculaire au plan méridien. Il est tangent au parallèle passant par M .

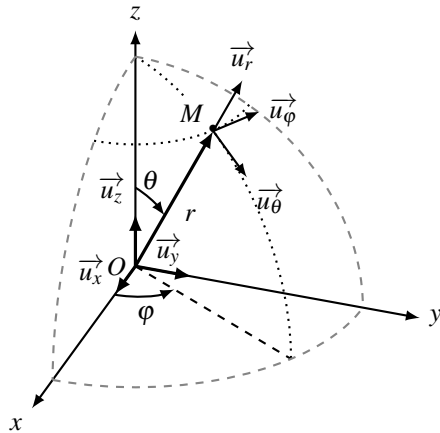


Figure 13.14 – Repérage d'un point M en coordonnées sphériques.

Un point M de l'espace de **coordonnées sphériques** r , θ et φ est noté $M(r, \theta, \varphi)$. Le vecteur position s'écrit :

$$\overrightarrow{OM} = r\overrightarrow{u_r},$$

où r est la distance OM .

La base de projection $(\overrightarrow{u_r}, \overrightarrow{u_\theta}, \overrightarrow{u_\varphi})$ est définie pour une position donnée du point M . Elle est définie localement en tout point M et tourne pour que le vecteur $\overrightarrow{u_r}$ pointe toujours de O vers M . Pour un point M donné, le triplet $(\overrightarrow{u_r}, \overrightarrow{u_\theta}, \overrightarrow{u_\varphi})$ forme une base mobile de l'espace. Tout vecteur $\overrightarrow{w}$ de l'espace peut s'écrire sous la forme :

$$\overrightarrow{w} = w_r\overrightarrow{u_r} + w_\theta\overrightarrow{u_\theta} + w_\varphi\overrightarrow{u_\varphi} \quad \text{noté} \quad \overrightarrow{w} = \begin{pmatrix} w_r \\ w_\theta \\ w_\varphi \end{pmatrix}.$$

w_r est la composante du vecteur $\overrightarrow{w}$ sur $\overrightarrow{u_r}$ appelée composante radiale. Les autres composantes n'ont pas de dénominations particulières.

w_r , w_θ et w_φ sont les projections orthogonales de $\overrightarrow{w}$ sur les vecteurs $\overrightarrow{u_r}$, $\overrightarrow{u_\theta}$ et $\overrightarrow{u_\varphi}$. En particulier, le vecteur position $\overrightarrow{OM} = r\overrightarrow{u_r}$ s'écrit :

$$\overrightarrow{OM} = \begin{pmatrix} r \\ 0 \\ 0 \end{pmatrix}.$$



Encore une fois en coordonnées sphériques, il ne faut pas confondre les coordonnées du point M et les composantes du vecteur position.

CHAPITRE 13 – CINÉMATIQUE DU POINT

4 Cinématique du point

4.1 Notion de référentiel

a) Exemple introductif

On considère un observateur O_1 situé dans un train qui avance à vitesse uniforme sur des rails rectilignes. Un observateur O_2 est posté au bord des rails et voit le train passer de gauche à droite. L'observateur O_1 situé dans le train lâche une petite balle, qui rebondit à ses pieds et remonte dans sa main. Le mouvement est rectiligne (Figure 13.15). L'observateur O_2 observe le même mouvement. Il voit la balle suivre une trajectoire parabolique vers le bas puis parabolique vers le haut pour remonter dans la main de l'observateur O_1 (Figure 13.15). Le mouvement de la balle observé par les deux observateurs est différent. La description d'un même phénomène dépend de l'observateur. Pour pouvoir comparer les observations des deux observateurs, il faut impérativement qu'ils précisent le cadre spatio-temporel de leurs observations : ils doivent préciser le référentiel d'étude.

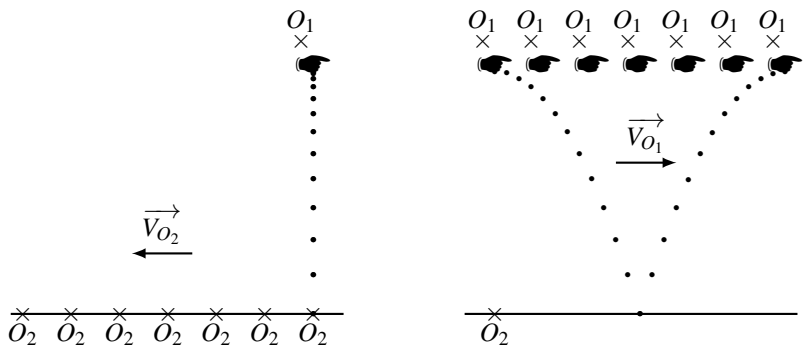


Figure 13.15 – À gauche, mouvement de la balle tel qu'observé par l'observateur O_1 .
À droite, mouvement de la balle tel qu'observé par l'observateur O_2 .

b) Définition

Le **référentiel** d'une étude mécanique est le cadre spatio-temporel de l'étude. Pour le définir, on a besoin de deux repères : un repère de temps et un repère d'espace.

Relativité du mouvement La description du mouvement du système dépend du référentiel dans lequel on l'étudie, on parle de relativité du mouvement. Pour une étude mécanique, il faut systématiquement préciser le référentiel d'étude.

Exemple

Pour fixer les idées, on va décrire les référentiels de l'exemple introductif.
Le référentiel $\mathcal{R}_1$ de l'observateur O_1 lié au train est caractérisé par un repère de temps et un repère d'espace fixe par rapport au train d'origine O_1 et d'axes dirigés vers l'avant

du wagon, vers sa gauche et vers le haut. L'observateur O_1 mesure le mouvement par rapport à ce référentiel $\mathcal{R}_1$. Dans ce référentiel, la vitesse initiale de la balle est nulle et son mouvement est rectiligne selon la verticale.

Le référentiel $\mathcal{R}_2$ de l'observateur O_2 lié au sol est caractérisé par un repère de temps et un repère d'espace fixe par rapport au sol d'origine O_2 et d'axes dirigés vers la droite de l'observateur, vers l'avant de l'observateur et vers le haut. L'observateur O_2 mesure le mouvement par rapport à ce référentiel $\mathcal{R}_2$. Dans ce référentiel, la balle a une vitesse initiale dirigée vers la droite et son mouvement est un mouvement parabolique.

c) Repère de temps

Pour dater un événement, on a besoin d'un repère de temps constitué :

- d'une mesure de temps définie par l'unité de durée : la seconde de symbole s ;
- d'une origine des temps.

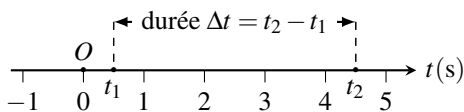


Figure 13.16 – Repérage d'un instant sur la flèche du temps.

On définit alors un vocabulaire scientifique précis sur cette flèche du temps :

- un instant est un point de la flèche du temps ;
- l'origine des dates est un instant particulier choisi comme instant zéro ;
- une date t repère la position d'un instant sur la flèche du temps. C'est la durée séparant l'origine des temps et cet instant. Une date est une grandeur algébrique : elle est positive si l'instant t est postérieur à l'instant initial et négative s'il est antérieur à l'instant initial.
- une durée $\Delta t = |t_2 - t_1|$ est la largeur d'un intervalle de temps séparant deux instants de dates t_1 et t_2 . Une durée est positive.

Remarque

L'origine des temps n'a pas de sens physique (seules les durées en ont) et peut être modifiée si nécessaire. La seconde est définie au niveau international par un traité diplomatique appelé *Convention du Mètre*. La définition de la seconde a varié dans l'histoire des sciences. Elle est actuellement définie à l'aide d'horloges atomiques.

d) Repère d'espace

Pour repérer un point dans l'espace, on a besoin d'un repère d'espace constitué :

- d'une mesure d'espace définie par l'unité de longueur : le mètre, symbole m ;
- d'une origine d'espace : un point fixe dans le référentiel d'étude ;
- de trois directions orthogonales fixes repérées par 3 axes vu précédemment.

CHAPITRE 13 – CINÉMATIQUE DU POINT

Remarque

L'origine et l'unité de longueur sont arbitraires. L'origine n'a pas de sens physique et peut être modifiée si nécessaire. Le mètre est défini au niveau international par la *Convention du Mètre*. La définition du mètre a varié dans l'histoire des sciences. Il est actuellement défini à partir de la seconde et de la vitesse de la lumière. En effet, la vitesse de la lumière est une constante universelle fixée par convention à la valeur de $c = 299\,792\,458 \text{ m}\cdot\text{s}^{-1}$. Fixer l'unité de temps fixe automatiquement l'unité de longueur.

4.2 Vecteurs position, déplacement, vitesse et accélération

Soit O un point fixe du référentiel d'étude $\mathcal{R}$ et M un point mobile.

a) Vecteur position

Le vecteur $\overrightarrow{OM}$ est le vecteur position. Son évolution temporelle $\overrightarrow{OM}(t)$ est donnée par les **équations horaires** du mouvement.

La courbe décrite par le point M au cours du mouvement est la **trajectoire** du point M .

Remarque

La notion de temps n'est plus présente dans la notion de trajectoire.

b) Vecteur déplacement

On considère que le mobile M est situé au point $M(t)$ à l'instant t et au point $M(t + \Delta t)$ à l'instant $t + \Delta t$.

Entre t et $t + \Delta t$, M se déplace de $\Delta\overrightarrow{OM} = \overrightarrow{M(t)M(t + \Delta t)}$ appelé **vecteur déplacement**.



En physique, la notation Δf est utilisée pour parler de la différence entre la valeur finale f_{finale} et la valeur initiale f_{initiale} de f : $\Delta f = f_{\text{finale}} - f_{\text{initiale}}$.

Ici, le vecteur position $\overrightarrow{OM}$ évolue entre $\overrightarrow{OM}(t)$ (valeur initiale) et $\overrightarrow{OM}(t + \Delta t)$ (valeur finale) donc $\Delta\overrightarrow{OM} = \overrightarrow{OM}(t + \Delta t) - \overrightarrow{OM}(t) = \overrightarrow{M(t)M(t + \Delta t)}$.

c) Vecteur déplacement élémentaire

Lorsque la durée Δt devient très faible, on la note dt et on parle de durée élémentaire ou infinitésimale. Pendant cette durée dt , le mobile M passe du point $M(t)$ au point $M(t + dt)$.

Entre t et $t + dt$, M se déplace de $d\overrightarrow{OM} = \overrightarrow{M(t)M(t + dt)}$. C'est le **déplacement élémentaire** ou infinitésimal.

De part sa définition, le déplacement élémentaire est tangent à la trajectoire.

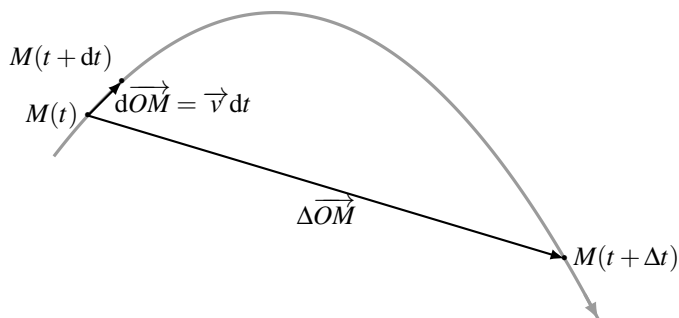


Figure 13.17 – Vecteurs déplacement et déplacement infinitésimal sur une trajectoire tracée en gris et décrite dans le sens de la flèche. Lorsque la durée Δt du déplacement devient infinitésimale et égale à dt , le déplacement $\overrightarrow{\Delta OM}$ devient infinitésimal et égal à $d\overrightarrow{OM}$, tangent en $M(t)$ à la trajectoire.

d) Vecteur vitesse

Le vecteur **vitesse moyenne** entre l’instant t et l’instant $t + \Delta t$ est défini par le rapport du vecteur déplacement à la durée de ce déplacement :

$$\langle \vec{v} \rangle = \frac{\overrightarrow{\Delta OM}}{\Delta t}.$$

Lorsque l’intervalle de temps devient infinitésimal, ce rapport devient le vecteur **vitesse instantanée** :

$$\vec{v}(t) = \frac{d\overrightarrow{OM}}{dt}.$$

C’est la dérivée du vecteur position $\overrightarrow{OM}$. La norme de ces vitesses se mesure en $\text{m}\cdot\text{s}^{-1}$.

e) Lien entre le déplacement élémentaire et la vitesse

La dérivée temporelle, en physique comme en mathématique, est la limite du taux de variation lorsque l’intervalle de temps entre deux mesures devient très petit. L’intervalle de temps dt et le déplacement élémentaire $d\overrightarrow{OM}$ ont un sens et une existence propre en physique. Ils correspondent à la notion de différentielle en mathématique.

Le vecteur déplacement élémentaire est la différentielle du vecteur position et s’exprime en fonction de la vitesse de la manière suivante :

$$d\overrightarrow{OM} = \overrightarrow{OM}(t + dt) - \overrightarrow{OM}(t) = \overrightarrow{M(t)M(t + dt)} = \vec{v} dt.$$

La vitesse est colinéaire au déplacement élémentaire. Comme lui, elle est parallèle à la trajectoire.

CHAPITRE 13 – CINÉMATIQUE DU POINT

Remarque

La durée infinitésimale dt peut devenir aussi courte que l'on veut ; le déplacement infinitésimal également. On écrit alors :

$$\begin{cases} d\vec{OM} &= \lim_{\Delta t \rightarrow 0} \Delta\vec{OM} \\ \vec{v} = \frac{d\vec{OM}}{dt} &= \lim_{\Delta t \rightarrow 0} \frac{\Delta\vec{OM}}{\Delta t}. \end{cases}$$

Le vecteur vitesse est bien la dérivée temporelle du vecteur position.

f) Vecteur accélération

Le vecteur accélération moyenne entre l'instant t et un instant $t + \Delta t$ est défini par le rapport de la variation du vecteur vitesse à la durée Δt :

$$\langle \vec{a} \rangle = \frac{\vec{v}(t + \Delta t) - \vec{v}(t)}{\Delta t} = \frac{\Delta \vec{v}}{\Delta t}.$$

Lorsque l'intervalle de temps devient infinitésimal, ce rapport devient le vecteur accélération instantanée :

$$\vec{a} = \frac{d\vec{v}}{dt} = \frac{d^2\vec{OM}}{dt^2}.$$

La norme de ces accélérations se mesure en $\text{m} \cdot \text{s}^{-2}$.

5 Utilisation des différents systèmes de coordonnées

5.1 Coordonnées cartésiennes

a) Cas d'un mouvement plan

Dans un premier temps, on fait l'hypothèse que le mouvement est réalisé dans le plan (Oxy) . On se place dans le référentiel $\mathcal{R}$ lié au repère d'espace fixe (Oxy) . On utilise la base de projection $(\vec{u}_x, \vec{u}_y)$ fixe dans $\mathcal{R}$ pour exprimer les vecteurs position, déplacement élémentaire, vitesse et accélération.

b) Vecteur position

L'expression du vecteur position a été déterminée précédemment pour un point $M(x, y)$. Le point M se déplace au cours du temps donc ses coordonnées x et y sont maintenant des fonctions du temps et :

$$\vec{OM} = x(t)\vec{u}_x + y(t)\vec{u}_y.$$

c) Déplacement élémentaire

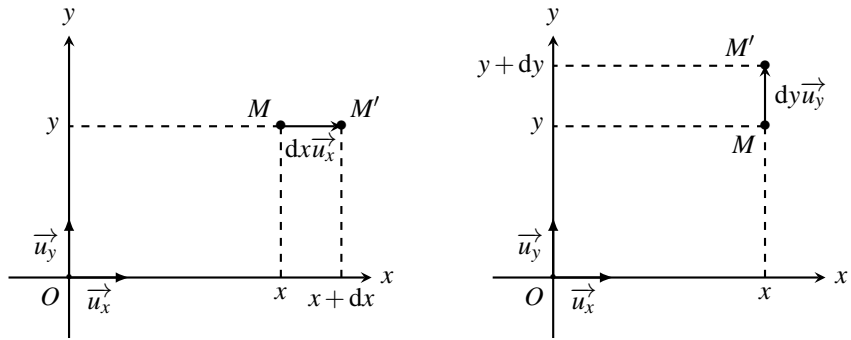


Figure 13.18 – Déplacement élémentaire en coordonnées cartésiennes. M' est tel que $\overrightarrow{MM'} = d\overrightarrow{OM}$. À gauche, x varie à y fixé ; à droite, y varie à x fixé.

Lorsque x varie de dx à y fixé, le déplacement a lieu selon le vecteur $\vec{u}_x$ et vaut $d\overrightarrow{OM} = dx\vec{u}_x$. Lorsque y varie de dy à x fixé, le déplacement a lieu selon le vecteur $\vec{u}_y$ et vaut $d\overrightarrow{OM} = dy\vec{u}_y$. Les coordonnées x et y étant indépendantes et la base $(\vec{u}_x, \vec{u}_y)$ orthonormée, on trouve le déplacement élémentaire dans le cas général où x et y varient simultanément en sommant ces deux expressions :

$$d\overrightarrow{OM} = dx\vec{u}_x + dy\vec{u}_y.$$

d) Vecteur vitesse

Le vecteur vitesse est égal au rapport du vecteur déplacement élémentaire à la durée élémentaire dt de ce déplacement :

$$\vec{v}(t) = \frac{d\overrightarrow{OM}}{dt} = \frac{dx}{dt}\vec{u}_x + \frac{dy}{dt}\vec{u}_y.$$

En notant $\dot{x}$ le rapport $\frac{dx}{dt}$ qui est également la dérivée temporelle de x , et $\dot{y} = \frac{dy}{dt}$, on obtient l'expression de la vitesse :

$$\vec{v}(t) = \dot{x}\vec{u}_x + \dot{y}\vec{u}_y.$$

Remarque

Cela revient à dériver le vecteur position terme à terme, en utilisant le fait que les vecteurs de la base $(\vec{u}_x, \vec{u}_y)$ sont fixes donc indépendant du temps.

e) Vecteur accélération

On continue alors la dérivation terme à terme de la vitesse pour trouver le vecteur accélération :

$$\vec{a}(t) = \frac{d\vec{v}}{dt} = \ddot{x}\vec{u}_x + \ddot{y}\vec{u}_y.$$

CHAPITRE 13 – CINÉMATIQUE DU POINT

f) Généralisation : mouvement dans l'espace repéré en cartésien

On peut généraliser l'analyse du mouvement pour un point M se déplaçant dans l'espace à trois dimensions en introduisant le point M_H , projeté orthogonal de M sur le plan (Oxy) . On obtient alors :

$$\overrightarrow{OM} = \overrightarrow{OM_H} + \overrightarrow{M_HM}.$$

Le point M_H est repéré en coordonnées cartésiennes dans le plan (Oxy) . Le point M est repéré sur l'axe orienté $(M_H, \vec{u}_z)$ par sa coordonnée z . On obtient alors les caractéristiques des mouvements de M par sommation des mouvement de M_H dans le plan (Oxy) et de M sur la droite $(M_H, \vec{u}_z)$.

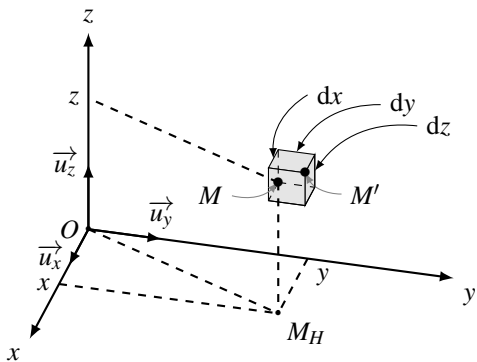


Figure 13.19 – Déplacement élémentaire en coordonnées cartésiennes. M' est tel que $\overrightarrow{MM'} = d\overrightarrow{OM}$.

g) Bilan en repérage cartésien

Les vecteurs position, déplacement élémentaire, vitesse et accélération en coordonnées cartésiennes exprimés sur la base $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ sont respectivement :

$$\overrightarrow{OM} = \begin{pmatrix} x \\ y \\ z \end{pmatrix} ; \quad d\overrightarrow{OM} = \begin{pmatrix} dx \\ dy \\ dz \end{pmatrix} ; \quad \vec{v} = \begin{pmatrix} \dot{x} \\ \dot{y} \\ \dot{z} \end{pmatrix} ; \quad \vec{a} = \begin{pmatrix} \ddot{x} \\ \ddot{y} \\ \ddot{z} \end{pmatrix}.$$

5.2 Coordonnées cylindro-polaire

a) Coordonnées polaires

Dans un premier temps, on fait l'hypothèse que le mouvement s'inscrit sur le plan (Oxy) . On se place dans le référentiel $\mathcal{R}$ lié au repère d'espace fixe (Oxy) . On utilise la base de projection $(\vec{u}_r, \vec{u}_\theta)$ mobile dans $\mathcal{R}$ pour exprimer les vecteurs position, déplacement élémentaire, vitesse et accélération du point M .

b) Vecteur position

L'expression du vecteur position a été déterminée précédemment pour un point $M(r, \theta)$. Le point M se déplace au cours du temps donc ses coordonnées r et θ sont maintenant des fonctions du temps. Le vecteur position s'écrit :

$$\overrightarrow{OM} = r(t)\overrightarrow{u_r}(t).$$



Une nouveauté intervient en coordonnées polaires : le vecteur $\overrightarrow{u_r}$ suit le point M au cours de son mouvement et évolue au cours du temps. $\overrightarrow{u_r}$ est donc une fonction du temps.

c) Déplacement élémentaire

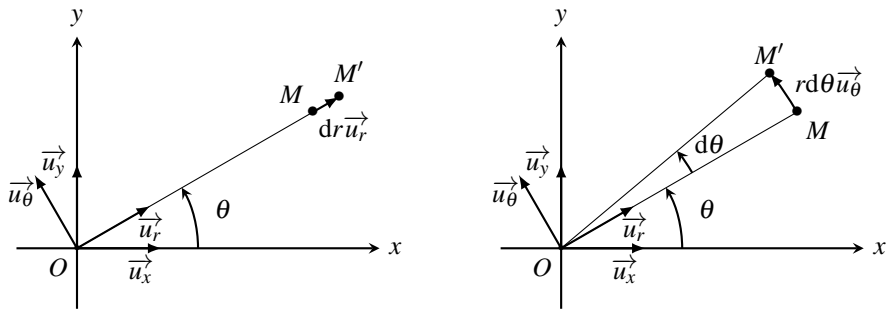


Figure 13.20 – Déplacement élémentaire en coordonnées polaires. M' est tel que $\overrightarrow{MM'} = d\overrightarrow{OM}$. À gauche, r varie à θ fixé ; à droite, θ varie à r fixé.

Lorsque r varie de dr à θ fixé, le déplacement a lieu selon le vecteur $\overrightarrow{u_r}$ et vaut : $d\overrightarrow{OM} = dr\overrightarrow{u_r}$.
Lorsque θ varie de $d\theta$ à r fixé, le déplacement a lieu selon $\overrightarrow{u_\theta}$ et vaut : $d\overrightarrow{OM} = r d\theta \overrightarrow{u_\theta}$.

Les coordonnées r et θ étant indépendantes et la base $(\overrightarrow{u_r}, \overrightarrow{u_\theta})$ orthonormée, on trouve le déplacement élémentaire dans le cas général où r et θ varient simultanément en sommant ces deux expressions :

$$d\overrightarrow{OM} = dr\overrightarrow{u_r} + r d\theta \overrightarrow{u_\theta}.$$



L'orientation de l'angle θ et le choix du vecteur de base $\overrightarrow{u_\theta}$ sont liés. $\overrightarrow{u_\theta}$ est dirigé dans le sens où θ augmente. Il faut toujours le vérifier sur son schéma. Une erreur d'orientation entraîne automatiquement des erreurs de signe.

d) Vecteur vitesse

Le vecteur vitesse est égal au rapport du vecteur déplacement élémentaire à la durée élémentaire dt de ce déplacement :

CHAPITRE 13 – CINÉMATIQUE DU POINT

$$\vec{v}(t) = \frac{d\vec{OM}}{dt} = \frac{dr}{dt}\vec{u}_r + r\frac{d\theta}{dt}\vec{u}_\theta.$$

On obtient alors l’expression du vecteur vitesse :

$$\vec{v}(t) = \dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta.$$

Remarque

On peut également trouver le vecteur vitesse en dérivant le vecteur position par rapport au temps :

$$\vec{v}(t) = \frac{d\vec{OM}}{dt} = \frac{d(r(t)\vec{u}_r(t))}{dt} = \frac{dr}{dt}\vec{u}_r + r\frac{d\vec{u}_r}{dt}.$$

On peut alors conjecturer que la dérivée de $\vec{u}_r$ par rapport au temps vaut $\frac{d\vec{u}_r}{dt} = \dot{\theta}\vec{u}_\theta$. C’est ce que nous allons démontrer dans le paragraphe suivant.

e) Vecteur accélération

Pour trouver le vecteur accélération, on ne peut plus utiliser de méthode géométrique simple. On va donc dériver le vecteur vitesse par rapport au temps. Pour cela, on doit établir les expression des dérivées temporelles de $\vec{u}_r$ et de $\vec{u}_\theta$.

Expression des vecteurs de la base mobile dans la base fixe On peut passer de la base de projection mobile à la base de projection fixe grâce aux figures de projection suivantes :

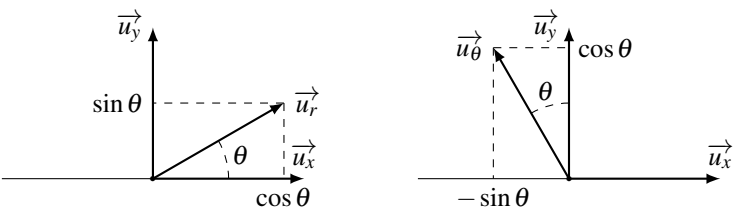


Figure 13.21 – Figures de projection des vecteurs de la base mobile dans la base cartésienne fixe.

Cela permet de trouver l’expression des vecteurs $\vec{u}_r$ et $\vec{u}_\theta$ dans la base fixe $(\vec{u}_x, \vec{u}_y)$:

$$\begin{cases} \vec{u}_r &= \cos \theta \vec{u}_x + \sin \theta \vec{u}_y \\ \vec{u}_\theta &= -\sin \theta \vec{u}_x + \cos \theta \vec{u}_y. \end{cases}$$

Remarque

Il est conseillé de réaliser ces figures avec $0 < \theta < 90^\circ$ et θ nettement inférieur à 45° afin de distinguer nettement que $\cos \theta > \sin \theta > 0$. On peut alors vérifier que l'on ne s'est pas trompé ni sur les signes, ni sur le choix entre $\sin \theta$ et $\cos \theta$.

Dérivation des vecteurs de la base mobile Lorsque le point M se déplace au cours du temps, r et θ varient. On voit sur la figure 13.20 qu'une variation de r ne fait pas bouger la base mobile, tandis qu'une variation de θ fait tourner la base mobile d'un angle $d\theta$. Cette dépendance en θ de $\vec{u}_r$ et $\vec{u}_\theta$ se voit également dans leurs expressions dans la base cartésienne exprimées ci-dessus.

Dans un premier temps, on dérive $\vec{u}_r$ et $\vec{u}_\theta$ par rapport à θ :

$$\begin{cases} \frac{d\vec{u}_r}{d\theta} = -\sin \theta \vec{u}_x + \cos \theta \vec{u}_y = \vec{u}_\theta \\ \frac{d\vec{u}_\theta}{d\theta} = -\cos \theta \vec{u}_x - \sin \theta \vec{u}_y = -\vec{u}_r \end{cases} \iff \begin{cases} d\vec{u}_r = d\theta \vec{u}_\theta \\ d\vec{u}_\theta = -d\theta \vec{u}_r. \end{cases}$$

On en déduit leurs dérivées par rapport au temps, qu'il faut absolument connaître :

$$\begin{cases} \frac{d\vec{u}_r}{dt} = \dot{\theta} \vec{u}_\theta \\ \frac{d\vec{u}_\theta}{dt} = -\dot{\theta} \vec{u}_r. \end{cases}$$

Remarque

On retrouve bien $\frac{d\vec{u}_r}{dt} = \dot{\theta} \vec{u}_\theta$ que l'on avait conjecturé précédemment.

Calcul du vecteur accélération Il ne reste plus qu'à dériver le vecteur vitesse terme à terme pour trouver le vecteur accélération :

$$\begin{aligned} \vec{a}(t) &= \frac{d\vec{v}}{dt} = \frac{d(\dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta)}{dt} \\ &= \frac{d\dot{r}}{dt}\vec{u}_r + \dot{r}\frac{d\vec{u}_r}{dt} + \frac{dr}{dt}\dot{\theta}\vec{u}_\theta + r\frac{d\dot{\theta}}{dt}\vec{u}_\theta + r\dot{\theta}\frac{d\vec{u}_\theta}{dt} \\ &= \ddot{r}\vec{u}_r + \dot{r}(\dot{\theta}\vec{u}_\theta) + \dot{r}\dot{\theta}\vec{u}_\theta + r\ddot{\theta}\vec{u}_\theta + r\dot{\theta}(-\dot{\theta}\vec{u}_r). \end{aligned}$$

Finalement :

$$\vec{a}(t) = (\ddot{r} - r\dot{\theta}^2)\vec{u}_r + (2\dot{r}\dot{\theta} + r\ddot{\theta})\vec{u}_\theta.$$

f) Généralisation : mouvement à trois dimensions en coordonnées cylindriques

On peut généraliser l'analyse du mouvement pour un point M se déplaçant dans l'espace à trois dimensions en introduisant le point M_H , projeté orthogonal de M sur le plan (Oxy) . On

CHAPITRE 13 – CINÉMATIQUE DU POINT

obtient alors :

$$\overrightarrow{OM} = \overrightarrow{OM_H} + \overrightarrow{M_H M}.$$

Le point M_H est repéré en coordonnées polaires dans le plan (Oxy) . Le point M est repéré sur l'axe orienté $(M_H, \vec{u}_z)$ par sa coordonnée z . On obtient alors les caractéristiques du mouvement de M par sommation des mouvements de M_H dans le plan (Oxy) et de M sur la droite $(M_H, \vec{u}_z)$.

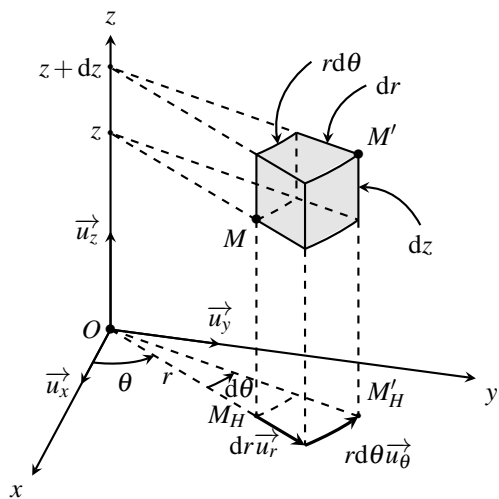


Figure 13.22 – Déplacement élémentaire en coordonnées cylindriques. M' est tel que $\overrightarrow{MM'} = d\overrightarrow{OM}$. Le déplacement de M est la somme du déplacement de M_H dans le plan (Oxy) repéré en polaire et du déplacement de M par rapport à M_H selon l'axe $(M_H, \vec{u}_z)$.

g) Bilan en repérage cylindro-polaire

Les vecteurs position, déplacement élémentaire, vitesse et accélération en coordonnées polaires exprimés sur la base mobile $(\vec{u}_r, \vec{u}_\theta, \vec{u}_z)$ sont respectivement

$$\overrightarrow{OM} = \begin{pmatrix} r \\ 0 \\ z \end{pmatrix} ; \quad d\overrightarrow{OM} = \begin{pmatrix} dr \\ r d\theta \\ dz \end{pmatrix} ; \quad \vec{v} = \begin{pmatrix} \dot{r} \\ r\dot{\theta} \\ \dot{z} \end{pmatrix} ; \quad \vec{a} = \begin{pmatrix} \ddot{r} - r\dot{\theta}^2 \\ 2\dot{r}\dot{\theta} + r\ddot{\theta} \\ \ddot{z} \end{pmatrix}.$$

Par ailleurs, il faut retenir que

$$\begin{cases} \frac{d\vec{u}_r}{dt} = \dot{\theta}\vec{u}_\theta \\ \frac{d\vec{u}_\theta}{dt} = -\dot{\theta}\vec{u}_r. \end{cases}$$

5.3 Coordonnées sphériques

a) Coordonnées sphériques

Le point M est repéré dans son mouvement à trois dimensions par ses coordonnées sphériques. On se place dans le référentiel $\mathcal{R}$ lié au repère d'espace fixe $(Oxyz)$. On utilise la base de projection $(\vec{u}_r, \vec{u}_\theta, \vec{u}_\varphi)$ mobile dans $\mathcal{R}$. L'expression de l'accélération est difficile à obtenir dans cette base et on se contentera ici d'exprimer les vecteurs position, déplacement élémentaire et vitesse du point M .

b) Vecteur position

L'expression du vecteur position a été déterminée précédemment pour un point $M(r, \theta, \varphi)$. Le point M se déplace au cours du temps donc ses coordonnées r , θ et φ sont des fonctions du temps. Le vecteur position s'écrit :

$$\overrightarrow{OM} = r(t)\vec{u}_r(t).$$

Le fait que θ et φ dépendent de temps implique que le vecteur $\vec{u}_r$ dépend du temps.

c) Déplacement élémentaire

On considère un point situé en M de coordonnées sphériques (r, θ, φ) à l'instant t et qui arrive en un point M' tel que $\overrightarrow{MM'} = d\overrightarrow{OM}$ de coordonnées $(r + dr, \theta + d\theta, \varphi + d\varphi)$ à l'instant $t + dt$. On note M_V la projection de M sur l'axe « vertical » (Oz).

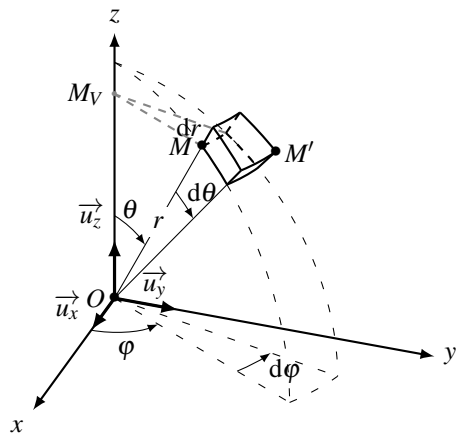


Figure 13.23 – Déplacement élémentaire en coordonnées sphériques.

Un déplacement à φ fixé correspond à un déplacement dans le plan méridien tandis qu'un déplacement à r et θ fixés correspond à un déplacement sur le parallèle passant par M .

Lorsque r varie de dr à θ et φ fixés, le déplacement a lieu selon $\vec{u}_r$ et vaut :

$$d\overrightarrow{OM} = dr\vec{u}_r.$$

CHAPITRE 13 – CINÉMATIQUE DU POINT

Lorsque θ varie de $d\theta$ à r et φ fixés, le déplacement a lieu selon $\vec{u}_\theta$ et vaut :

$$d\vec{OM} = rd\theta \vec{u}_\theta.$$

Lorsque φ varie de $d\varphi$, à r et θ fixés, le déplacement a lieu sur le parallèle passant par M de rayon $r \sin \theta$ et de centre M_V . Il se fait selon le vecteur $\vec{u}_\varphi$ et vaut :

$$d\vec{OM} = r \sin \theta d\varphi \vec{u}_\varphi.$$

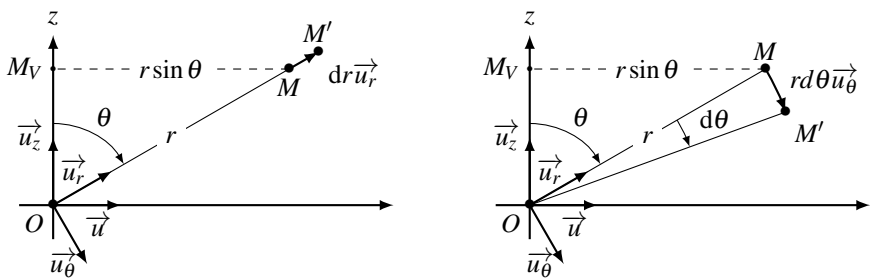


Figure 13.24 – Déplacement de M sur le plan méridien contenant M à φ fixé. Dans ce plan, M est repéré en coordonnées polaires. À gauche, r varie à θ et φ fixés ; à droite, θ varie à r et φ fixés.

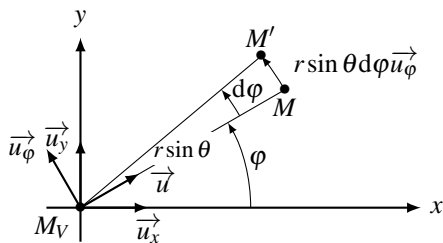


Figure 13.25 – Déplacement de M à r et θ fixés. Le déplacement a lieu sur le parallèle passant par M de rayon $M_V M = r \sin \theta$, où M_V est la projection orthogonale de M sur l'axe (Ox) .

Les coordonnées r , θ et φ étant indépendantes et la base $(\vec{u}_r, \vec{u}_\theta, \vec{u}_\varphi)$ orthonormée, on trouve le déplacement élémentaire dans le cas général où r , θ et φ varient simultanément en sommant ces trois expressions :

$$d\vec{OM} = dr \vec{u}_r + rd\theta \vec{u}_\theta + r \sin \theta d\varphi \vec{u}_\varphi.$$

Remarque

Le rayon d'un cercle est une distance donc une grandeur positive. $\sin \theta$ doit donc garder des valeurs positives. Pour cela, en coordonnées sphériques, θ varie dans l'intervalle $[0, \pi]$. Par suite, pour décrire l'ensemble de l'espace, φ varie sur $[0, 2\pi[$.



L'orientation de l'angle θ et le choix du vecteur $\vec{u}_\theta$ sont liés. $\vec{u}_\theta$ est dirigé dans le sens où θ augmente. De même, l'orientation de φ et le choix du vecteur $\vec{u}_\varphi$ sont liés. $\vec{u}_\varphi$ est dirigé dans le sens où φ augmente.

d) Vecteur vitesse

Le vecteur vitesse est égal au rapport du vecteur déplacement élémentaire à la durée élémentaire dt de ce déplacement :

$$\vec{v}(t) = \frac{d\vec{OM}}{dt} = \frac{dr}{dt} \vec{u}_r + r \frac{d\theta}{dt} \vec{u}_\theta + r \sin \theta \frac{d\varphi}{dt} \vec{u}_\varphi.$$

On obtient alors l'expression du vecteur vitesse :

$$\vec{v}(t) = \dot{r} \vec{u}_r + r \dot{\theta} \vec{u}_\theta + r \sin \theta \dot{\varphi} \vec{u}_\varphi.$$

6 Exemples de mouvements étudiés en coordonnées cartésiennes

6.1 Mouvements rectilignes

a) Mouvement rectiligne et uniforme

Exemple

Un train se déplaçant à vitesse constante sur des rails rectilignes suit un mouvement rectiligne et uniforme par rapport au sol.

On considère un point M en mouvement rectiligne et uniforme de vecteur vitesse $\vec{v}_0$ constant dans le référentiel $\mathcal{R}$. Le mouvement s'inscrit sur une droite fixe $\mathcal{D}$ qui contient forcément $\vec{v}_0$.

Choix du repère Le mouvement de M est rectiligne. On repère M sur une droite comme décrit au paragraphe 2.2. On choisit l'origine O du repère sur la droite $\mathcal{D}$ et le vecteur unitaire $\vec{u}_x$ sur cette droite. En général, on s'arrange pour que $\vec{u}_x$ soit dans le sens de $\vec{v}_0$ et on projette tous les vecteurs du problème sur $\vec{u}_x$. Le mouvement de M est entièrement paramétré par son abscisse $x(t)$ que l'on doit rechercher. C'est l'équation horaire du mouvement.

Équation horaire du mouvement Les vecteurs position, vitesse et accélération sont *a priori* de la forme :

$$\vec{OM} = x \vec{u}_x \quad ; \quad \vec{v} = \dot{x} \vec{u}_x \quad ; \quad \vec{a} = \ddot{x} \vec{u}_x.$$

CHAPITRE 13 – CINÉMATIQUE DU POINT

Le mouvement de M est uniforme. Cela signifie qu'il se fait à norme de la vitesse constante. Cela implique que $\dot{x} = v_0$ à tout instant puis $\ddot{x} = 0$.

L'intégration de l'équation $\dot{x} = v_0$ donne l'équation horaire du mouvement :

$$x = v_0 t + x_0.$$

Dans cette équation, x_0 est une constante d'intégration qui s'interprète comme la coordonnée du point M à l'instant $t = 0$ puisque : $x(t = 0) = x_0$.

Bilan pour un mouvement rectiligne et uniforme d'axe (Ox)

$$\vec{a} = \vec{0} \quad ; \quad \vec{v} = v_0 \vec{u}_x \quad ; \quad \vec{OM} = (v_0 t + x_0) \vec{u}_x.$$

Un mouvement rectiligne et uniforme est un mouvement à vecteur accélération nul : $\vec{a} = \vec{0}$ et à vecteur vitesse constant $\vec{v} = \vec{v}_0 = v_0 \vec{u}_x$. L'abscisse x de M évolue linéairement avec t sous la forme $x = v_0 t + x_0$ avec x_0 l'abscisse initiale de M .

b) Mouvement rectiligne uniformément varié

Exemple

L'exemple type de mouvement rectiligne uniformément varié dans le référentiel terrestre est le mouvement d'un point en chute libre dans le vide. Ce mouvement est illustré sur la figure 13.2.

On considère un point M en mouvement rectiligne dont le mouvement s'inscrit sur la droite fixe $\mathcal{D}$ dans le référentiel $\mathcal{R}$. Le mouvement est uniformément varié, cela signifie que la norme de sa vitesse est une fonction affine du temps.

Choix du repère Le repère choisi est identique au précédent ; la base de projection également. Le mouvement de M est entièrement paramétré par son abscisse $x(t)$.

Équation horaire du mouvement Les vecteurs position, vitesse et accélération sont *a priori* de la forme :

$$\vec{OM} = x \vec{u}_x \quad ; \quad \vec{v} = \dot{x} \vec{u}_x \quad ; \quad \vec{a} = \ddot{x} \vec{u}_x.$$

La vitesse est une fonction affine du temps. Elle s'écrit donc :

$$\dot{x} = a_0 t + v_0,$$

où v_0 est la composante de la vitesse du point M selon $\vec{u}_x$ à l'instant $t = 0$.

En dérivant cette expression, on déduit que le vecteur accélération est nécessairement constant et égal à :

$$\vec{a} = a_0 \vec{u}_x$$

à tout instant. a_0 est donc la valeur algébrique de l'accélération.

L'intégration de l'équation $\dot{x} = a_0 t + v_0$ donne l'équation horaire du mouvement :

$$x = a_0 \frac{t^2}{2} + v_0 t + x_0.$$

Dans cette équation, x_0 est une constante d'intégration qui s'interprète comme la position sur l'axe (Ox) du point M à l'instant $t = 0$ puisque $x(t = 0) = x_0$.

Bilan pour un mouvement rectiligne uniformément varié d'axe (Ox)

$$\vec{a} = a_0 \vec{u}_x \quad ; \quad \vec{v} = (a_0 t + v_0) \vec{u}_x \quad ; \quad \vec{OM} = \left(a_0 \frac{t^2}{2} + v_0 t + x_0 \right) \vec{u}_x.$$

Selon les signes de v_0 et a_0 , donc selon les orientations de $\vec{v}_0$ et $\vec{a}_0$, on a deux types de mouvements possibles :

- si $v_0 > 0$ et $a_0 > 0$, le vecteur vitesse initiale et le vecteur accélération sont de même sens. La norme de la vitesse est une fonction affine croissante du temps. Les équations horaires décrivent un **mouvement rectiligne et uniformément accéléré**. Cette situation correspond à un mobile que l'on lance verticalement vers le bas dans le référentiel terrestre, avec un axe (Ox) orienté vers le bas.
- si $v_0 < 0$ et $a_0 > 0$, le vecteur vitesse initiale et le vecteur accélération sont de sens opposés. Le mouvement se déroule en deux phases : la norme de la vitesse commence par diminuer (le mouvement est alors uniformément décéléré), s'annule (le mobile atteint un point extrême), puis recommence à augmenter (le mobile a rebroussé chemin et suit un mouvement rectiligne et uniformément accéléré). Cette situation correspond à un mobile que l'on lance verticalement vers le haut dans le référentiel terrestre, avec un axe (Ox) orienté vers le bas.

c) Généralisation : mouvement rectiligne quelconque

On considère un point M en mouvement rectiligne dans le référentiel $\mathcal{R}$. Le mouvement s'inscrit sur la droite fixe $\mathcal{D}$.

Choix du repère Le repère choisi est identique au précédent ; la base de projection également. Le mouvement de M est entièrement paramétré par son abscisse $x(t)$.

Équation horaire du mouvement Les vecteurs position, vitesse et accélération sont :

$$\vec{OM} = x \vec{u}_x \quad ; \quad \vec{v} = \dot{x} \vec{u}_x \quad ; \quad \vec{a} = \ddot{x} \vec{u}_x.$$

L'accélération $\vec{a} = \ddot{x} \vec{u}_x$ est une fonction du temps que l'on intègre pour obtenir $\dot{x}$:

$$\int_0^t \ddot{x} dt = [\dot{x}]_0^t = \dot{x}(t) - \dot{x}(t = 0) \quad \implies \quad \dot{x}(t) = \dot{x}(t = 0) + \int_0^t \ddot{x} dt.$$

On déduit alors x par intégration de $\dot{x}$ par rapport au temps :

$$\int_0^t \dot{x} dt = [x]_0^t = x(t) - x(t=0) \quad \implies \quad x(t) = x(t = 0) + \int_0^t \dot{x} dt.$$

On voit apparaître deux constantes d'intégration :

CHAPITRE 13 – CINÉMATIQUE DU POINT

- la position de M à l'instant $t = 0 : x(t = 0)$;
- la vitesse de M à l'instant $t = 0 : \dot{x}(t = 0)$.

Ce sont les conditions initiales du mouvement.

6.2 Mouvements à vecteur accélération constante

Exemple

Sur Terre, les corps lancés dans le vide suivent un mouvement à vecteur accélération constante $\vec{a} = \vec{g}$ où $\vec{g}$ est l'accélération de la pesanteur. Ce type de mouvement a été présenté sur les chronophotographies du paragraphe 2.1. Selon la vitesse initiale $\vec{v}_0$, deux cas sont à envisager :

- un mouvement rectiligne dont un exemple est donné sur la figure 13.2 ;
- un mouvement parabolique dont un exemple est donné sur la figure 13.3.

On considère un point M dont le vecteur accélération est constant et égal à $\vec{a} = \vec{g}$ avec $\vec{g}$ indépendant du temps. On obtient alors $\vec{v}$ par intégration par rapport au temps :

$$\vec{v} = \vec{v}_0 + \vec{g}t,$$

où $\vec{v}_0$ est la vitesse de M à l'instant $t = 0$. On obtient la position en intégrant une nouvelle fois :

$$\vec{OM} = \vec{OM}_0 + \vec{v}_0t + \frac{1}{2}\vec{g}t^2.$$

Choix de l'origine du repère On choisit pour origine du repère la position du point M à l'instant $t = 0$ de sorte que $\vec{OM}_0$ est égal à $\vec{0}$ et on obtient :

$$\vec{OM} = \vec{v}_0t + \frac{1}{2}\vec{g}t^2.$$

On peut alors tracer la trajectoire de M comme sur la figure 13.26.

Le mouvement est toujours plan et on le peut décomposer en un mouvement rectiligne et uniforme à la vitesse $\vec{v}_0$ et un mouvement rectiligne uniformément accéléré parallèle à $\vec{g}$.

Le mouvement est manifestement plan et on a deux possibilités :

- $\vec{v}_0 \parallel \vec{g}$ ou $\vec{v}_0 = \vec{0}$ et le mouvement est rectiligne selon la droite $(O, \vec{g})$;
- $\vec{v}_0$ n'est pas parallèle à $\vec{g}$ et le mouvement est parabolique.

Par la suite, on n'envisage que ce deuxième cas.

Choix du repère Le mouvement est plan. On choisit le repérage cartésien décrit au paragraphe 5.1 avec (Oy) parallèle à $\vec{g}$. On peut alors projeter tous les vecteurs sur la base de projection $(\vec{u}_x, \vec{u}_y)$. Or $\vec{g} = -g\vec{u}_y$ et $\vec{v}_0 = v_0 \cos \alpha \vec{u}_x + v_0 \sin \alpha \vec{u}_y$, on trouve alors :

$$\vec{a} = \begin{pmatrix} \ddot{x} = 0 \\ \ddot{y} = -g \end{pmatrix} ; \vec{v} = \begin{pmatrix} \dot{x} = v_0 \cos \alpha \\ \dot{y} = -gt + v_0 \sin \alpha \end{pmatrix} ; \vec{OM} = \begin{pmatrix} x = v_0 \cos \alpha t \\ y = -\frac{1}{2}gt^2 + v_0 \sin \alpha t \end{pmatrix}$$

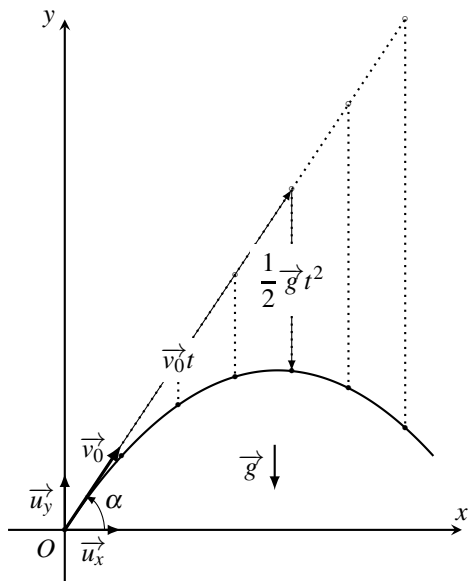


Figure 13.26 – Mouvement à accélération constante.

Les équations horaires du mouvements sont donc :

$$\begin{cases} x = v_0 \cos \alpha \, t \\ y = -\frac{1}{2} g t^2 + v_0 \sin \alpha \, t. \end{cases} \tag{13.1}$$

On obtient l'équation de la trajectoire de M en éliminant t entre les deux équations. Par hypothèse, $\vec{v}_0$ n'est pas parallèle à $\vec{g}$, de sorte que $\cos \alpha \neq 0$. On tire $t = \frac{x}{v_0 \cos \alpha}$ de la première équation et on l'injecte dans la deuxième :

$$y = -\frac{1}{2} \frac{g}{v_0^2 \cos^2 \alpha} x^2 + \tan \alpha \, x.$$

L'équation obtenue est celle d'une parabole : la trajectoire est parabolique.

Remarque

- Le mouvement rectiligne uniforme est un cas particulier du mouvement à accélération constante tel que $\vec{a} = \vec{0}$. On retrouve les équations horaires correspondante à partir des équations (13.1) en choisissant $\vec{u}_x$ parallèle à $\vec{v}_0$ puis en posant $g = 0$ et $\alpha = 0$.
- Le mouvement rectiligne et uniformément varié est un cas particulier du mouvement d'accélération constante pour lequel $\vec{v}_0$ et $\vec{g}$ ont même direction. On retrouve les équations horaires correspondantes à partir des équations (13.1) en posant $\alpha = \pm \frac{\pi}{2}$ et $g = -a_0$. Le mouvement s'inscrit alors sur la droite (Oy) .

CHAPITRE 13 – CINÉMATIQUE DU POINT

6.3 Mouvement rectiligne sinusoïdal : mouvement harmonique

Un mouvement rectiligne d'axe (Ox) tel que la coordonnée x du point M vérifie :

$$x(t) = X_m \cos(\omega_0 t + \varphi_0),$$

est appelé mouvement rectiligne sinusoïdal. Ce mouvement a été étudié en détail lors de l'étude de l'oscillateur harmonique.

Du point de vue cinématique, on remarque que la vitesse $\vec{v} = \dot{x}\vec{u}_x$ est telle que :

$$\dot{x}(t) = -\omega_0 X_m \sin(\omega_0 t + \varphi_0) = \omega_0 X_m \cos\left(\omega_0 t + \varphi_0 + \frac{\pi}{2}\right).$$

De même, l'accélération $\vec{a} = \ddot{x}\vec{u}_x$ vaut :

$$\ddot{x}(t) = -\omega_0^2 X_m \cos(\omega_0 t + \varphi_0) = \omega_0^2 X_m \cos(\omega_0 t + \varphi_0 + \pi).$$

La vitesse est donc en quadrature de phase positive avec la position (avance de phase de $\frac{\pi}{2}$) et l'accélération en opposition de phase (avance de phase de π).

À l'instant initial, la position vaut $x(t = 0) = X_m \cos \varphi_0$ et la vitesse $\dot{x}(t = 0) = -\omega_0 \sin \varphi_0$. L'abscisse $x(t)$ est maximale à l'instant t_0 tel que $\omega_0 t_0 + \varphi_0 = 0$ soit $t_0 = -\frac{\varphi_0}{\omega_0}$.

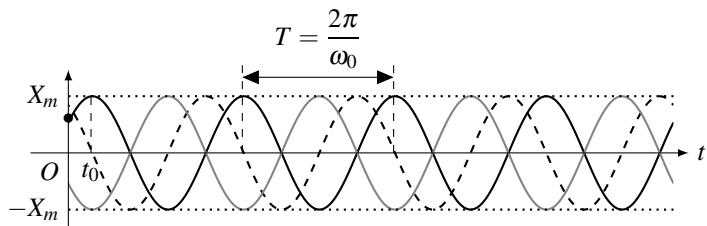


Figure 13.27 – Position, vitesse et accélération en fonction du temps lors d'un mouvement harmonique. $x(t)$ est en trait plein noir. $\frac{\dot{x}(t)}{\omega_0}$ en pointillés noirs et $\frac{\ddot{x}(t)}{\omega_0^2}$ en trait plein gris. On peut voir que $\dot{x}$ est en avance et en quadrature de phase par rapport à x et que $\ddot{x}$ est en opposition de phase avec x .

La caractéristique cinématique principale des mouvements harmoniques est que l'accélération $\ddot{x}$ et la position x sont reliées par une équation différentielle linéaire du deuxième ordre du type :

$$\ddot{x} + \omega_0^2 x = 0.$$

7 Mouvements circulaires

7.1 Mouvement circulaire et uniforme

Exemple

Le mouvement d'un corps posé sur un plateau tournant à vitesse angulaire constante est un mouvement circulaire et uniforme. Un exemple en est montré sur la figure 13.4 du paragraphe 2.1.

On considère un point M en mouvement circulaire et uniforme dans le référentiel $\mathcal{R}$. Le mouvement est réalisé sur un cercle $\mathcal{C}$ de centre C et de rayon R parcouru à vitesse constante (en norme).

Choix du repère Le mouvement de M est circulaire, il s'effectue donc dans un plan $\mathcal{P}$. On utilise le repérage polaire décrit au paragraphe 5.2 en choisissant le centre C du cercle $\mathcal{C}$ comme origine O du repère du plan et (Ox) un diamètre du cercle comme origine des angles. Tous les vecteurs du problème vont être projetés sur la base polaire mobile $(\vec{u}_r, \vec{u}_\theta)$.

Le mouvement de M étant inscrit sur un cercle de rayon R , la distance de M à l'origine O du repère est fixée à $r = R$. Le mouvement de M est entièrement paramétré par l'angle polaire $\theta(t)$ que l'on doit rechercher. C'est l'équation horaire du mouvement.

Équation horaire du mouvement Le vecteur position vaut :

$$\overrightarrow{OM} = R\vec{u}_r.$$

En utilisant le fait que R est constant et que $\frac{d\vec{u}_r}{dt} = \dot{\theta}\vec{u}_\theta$, on dérive $\overrightarrow{OM}$ pour trouver :

$$\vec{v} = R\dot{\theta}\vec{u}_\theta.$$

$\dot{\theta}$ est appelée la **vitesse angulaire** de M . De même $\ddot{\theta}$ est appelée **accélération angulaire**.

Le mouvement de M est uniforme. Cela signifie qu'il se fait à norme de la vitesse constante. La vitesse est donc de la forme :

$$\vec{v} = v_0\vec{u}_\theta.$$

Cela implique que $R\dot{\theta} = v_0$ est constante. La vitesse angulaire est constante et égale à $\dot{\theta} = \frac{v_0}{R}$.

On la note souvent ω_0 : $\dot{\theta} = \omega_0 = \frac{v_0}{R}$. La dérivation de cette relation montre que l'accélération angulaire est nulle : $\ddot{\theta} = 0$.

Son intégration permet de trouver l'équation horaire du mouvement :

$$\theta = \omega_0 t + \theta_0 = \frac{v_0}{R}t + \theta_0.$$

Dans cette équation, θ_0 est une constante d'intégration qui s'interprète comme la position angulaire de M à l'instant initial. Il faut remarquer que la norme de la vitesse vaut $|v_0|$ et que le signe de v_0 ou de ω_0 donne le sens de parcourt de la trajectoire circulaire :

- si $v_0 > 0$, la trajectoire est décrite dans le sens direct ;
- si $v_0 < 0$, la trajectoire est décrite dans le sens indirect.

CHAPITRE 13 – CINÉMATIQUE DU POINT

Bilan pour un mouvement circulaire et uniforme

$$\begin{cases} \overrightarrow{OM} = R\vec{u}_r \\ \vec{v} = R\omega_0\vec{u}_\theta = v_0\vec{u}_\theta \\ \vec{a} = -R\dot{\theta}^2\vec{u}_r = -\frac{v_0^2}{R}\vec{u}_r \end{cases}$$

Un mouvement circulaire et uniforme est un mouvement à accélération angulaire nulle : $\ddot{\theta} = 0$ et à vitesse angulaire constante $\dot{\theta} = \omega_0$. L'angle polaire θ évolue linéairement avec t sous la forme $\theta = \omega_0 t + \theta_0$.

La norme de la vitesse est constante et égale à $|v_0|$ mais le vecteur vitesse n'est pas constant puisque sa direction tourne au cours du temps. L'accélération n'est donc pas nulle. Elle est radiale, centripète¹ et de norme $\frac{v_0^2}{R}$.

7.2 Généralisation : mouvement circulaire quelconque

On considère un point M en mouvement circulaire dans le référentiel dans $\mathcal{R}$. Le mouvement est réalisé sur un cercle $\mathcal{C}$ de centre C et de rayon R .

Choix du repère Le mouvement de M est circulaire. On utilise le même repérage polaire que précédemment et la même base polaire de projection $(\vec{u}_r, \vec{u}_\theta)$.

Le mouvement de M étant effectué sur un cercle de rayon R , il est à nouveau entièrement paramétré par l'angle polaire $\theta(t)$ que l'on doit rechercher.

Position, vitesse et accélération Le vecteur position vaut :

$$\overrightarrow{OM} = R\vec{u}_r.$$

On dérive le vecteur $\overrightarrow{OM}$ pour trouver la vitesse :

$$\vec{v} = R\dot{\theta}\vec{u}_\theta.$$

On dérive la vitesse pour trouver l'accélération en faisant attention à $\dot{\theta}$ qui, cette fois, varie au cours du temps :

$$\begin{aligned} \vec{a} &= R\dot{\theta}\frac{d\vec{u}_\theta}{dt} + R\frac{d\dot{\theta}}{dt}\vec{u}_\theta = R\dot{\theta}(-\dot{\theta}\vec{u}_r) + R\ddot{\theta}\vec{u}_\theta \\ &= -R\dot{\theta}^2\vec{u}_r + R\ddot{\theta}\vec{u}_\theta. \end{aligned}$$

Comme dans le cas du mouvement circulaire et uniforme, $R\dot{\theta}^2 = \frac{v^2}{R}$. Par ailleurs, si l'on note $v(t) = R\dot{\theta}$, on obtient $R\ddot{\theta} = \frac{dv}{dt}$ et on en déduit :

$$\vec{a} = -R\dot{\theta}^2\vec{u}_r + R\ddot{\theta}\vec{u}_\theta = -\frac{v^2}{R}\vec{u}_r + \frac{dv}{dt}\vec{u}_\theta.$$

1. Un vecteur radial centripète est dirigé vers le centre de rotation O . Il est selon $-\vec{u}_r$. Un vecteur radial centrifuge fuit le centre de rotation O . Il est selon $+\vec{u}_r$.

On peut remarquer que :

- l'accélération radiale centripète de norme $\frac{v^2}{R}$ est perpendiculaire à la trajectoire et dirigée vers le centre du cercle. Elle est liée au fait que la trajectoire est courbe ;
- l'accélération orthoradiale $R\ddot{\theta} = \frac{dv}{dt}$ est parallèle à la vitesse. Elle est liée aux variations de la norme de la vitesse.

Équation horaire du mouvement L'accélération angulaire est une fonction du temps $\ddot{\theta}(t)$. On trouve alors la vitesse angulaire $\dot{\theta}$ par intégration de $\ddot{\theta}$ par rapport au temps :

$$\int_0^t \ddot{\theta} dt = [\dot{\theta}]_0^t = \dot{\theta}(t) - \dot{\theta}(t=0) \quad \implies \quad \dot{\theta}(t) = \dot{\theta}_{(0)} + \int_0^t \ddot{\theta} dt.$$

On déduit alors θ par intégration de $\dot{\theta}$ par rapport au temps :

$$\int_0^t \dot{\theta} dt = [\theta]_0^t = \theta(t) - \theta(t=0) \quad \implies \quad \theta(t) = \theta_{(0)} + \int_0^t \dot{\theta} dt.$$

On voit apparaître deux constantes d'intégration :

- la position angulaire de M à l'instant $t = 0$: $\theta(t=0)$;
- la vitesse angulaire de M à l'instant $t = 0$: $\dot{\theta}(t=0)$.

Ce sont les conditions initiales du mouvement.

Bilan pour un mouvement circulaire quelconque

$$\begin{cases} \overrightarrow{OM} = R\overrightarrow{u_r} \\ \overrightarrow{v} = R\dot{\theta}\overrightarrow{u_\theta} & = v(t)\overrightarrow{u_\theta} \\ \overrightarrow{a} = -R\dot{\theta}^2\overrightarrow{u_r} + R\ddot{\theta}\overrightarrow{u_\theta} & = -\frac{v^2}{R}\overrightarrow{u_r} + \frac{dv}{dt}\overrightarrow{u_\theta}. \end{cases}$$

8 Interprétation du vecteur accélération

8.1 Le vecteur vitesse et sa norme

Dans le langage courant, la « vitesse » est la grandeur affichée au compteur d'une voiture. Il s'agit en fait de la norme du vecteur vitesse. Le vocabulaire courant est donc souvent lié à cette norme notée v dans ce qui suit.



En physique, la vitesse désigne le vecteur vitesse et on précise lorsque l'on parle de sa norme. Le langage courant et le langage du physicien ne sont donc pas exactement identiques.

CHAPITRE 13 – CINÉMATIQUE DU POINT

8.2 Vecteur accélération et variation de la norme de la vitesse

a) Mouvement uniforme, accéléré ou décéléré

Un mouvement est **uniforme** si la norme du vecteur vitesse est constante :

$$\text{Mouvement uniforme} \iff v = \text{constante} \iff \frac{dv}{dt} = 0.$$

Un mouvement est **accéléré** si la norme du vecteur vitesse augmente :

$$\text{Mouvement accéléré} \iff v \nearrow \iff \frac{dv}{dt} > 0.$$

Un mouvement est **décéléré** si la norme du vecteur vitesse diminue :

$$\text{Mouvement décéléré} \iff v \searrow \iff \frac{dv}{dt} < 0.$$

Si la variation de la norme de la vitesse est uniforme dans le temps c'est-à-dire que $\frac{dv}{dt}$ est une constante non nulle, le mouvement est uniformément varié. Dans ce cas, la norme de la vitesse varie linéairement avec le temps.

Le mouvement est **uniformément accéléré** si la norme de la vitesse augmente proportionnellement au temps, c'est-à-dire $v = at + b$ avec $a > 0$ et $b > 0$.

Le mouvement est **uniformément décéléré** si la norme de la vitesse diminue proportionnellement au temps, c'est-à-dire $v = at + b$ avec $a < 0$ et $b > 0$.

b) Lien avec le vecteur accélération

Cas d'un mouvement rectiligne On se place dans le cas d'un mouvement rectiligne d'axe (Ox) . Les vecteurs vitesse et accélération sont colinéaires et s'écrivent $\vec{v} = \dot{x}\vec{u}_x$ et $\vec{a} = \ddot{x}\vec{u}_x$.

Pour analyser la situation, on suppose que $\dot{x} > 0$. Dans ce cas :

- si $\ddot{x} = 0$ à tout instant alors $\dot{x}$ reste constant. Le mouvement du mobile est uniforme ;
- si $\ddot{x} > 0$ alors $\vec{v}$ et $\vec{a}$ sont de même sens et $\dot{x}$ augmente. Le mouvement du mobile est accéléré ;
- si $\ddot{x} < 0$ alors $\vec{v}$ et $\vec{a}$ sont de sens opposés et $\dot{x}$ diminue. Le mouvement du mobile est décéléré.

On peut généraliser à l'aide du schéma ci-dessous pour un mouvement rectiligne :

- le mouvement est uniforme si $\vec{a} = \vec{0}$ à tout instant ;
- le mouvement est accéléré si $\vec{a}$ et $\vec{v}$ sont dans le même sens ;
- le mouvement est décéléré si $\vec{a}$ et $\vec{v}$ sont de sens opposés.

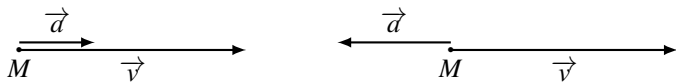


Figure 13.28 – À gauche, $\vec{a}$ et $\vec{v}$ sont dans le même sens, le mouvement rectiligne est accéléré. À droite, $\vec{a}$ et $\vec{v}$ sont de sens opposés, le mouvement rectiligne est décéléré.

Cas d'un mouvement quelconque L'analyse précédente reste bonne mais il ne faut pas considérer l'accélération $\vec{a}$ dans son entier mais uniquement sa projection $\vec{a}_{\parallel}$ sur la trajectoire donc sur le vecteur vitesse $\vec{v}$:

- le mouvement est uniforme si $\vec{a}_{\parallel} = \vec{0}$ à tout instant donc si l'accélération est perpendiculaire à la trajectoire à tout instant ;
- le mouvement est accéléré si $\vec{a}_{\parallel}$ et $\vec{v}$ sont dans le même sens ;
- le mouvement est décéléré si $\vec{a}_{\parallel}$ et $\vec{v}$ sont de sens opposés.

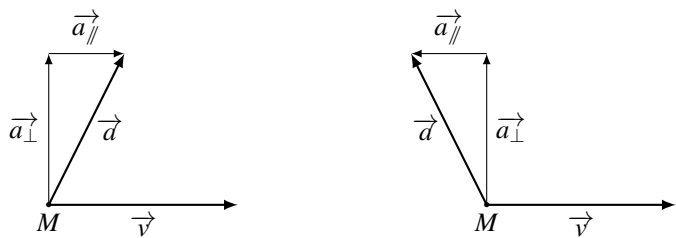


Figure 13.29 – À gauche, $\vec{a}_{\parallel}$ et $\vec{v}$ sont dans le même sens, le mouvement est accéléré. À droite, $\vec{a}_{\parallel}$ et $\vec{v}$ sont de sens opposés le mouvement est décéléré.

Pour justifier ce résultat, on peut remarquer que les variations de la norme de la vitesse sont les même que celles de son carré et en dérivant :

$$\frac{dv^2}{dt} = \frac{d(\vec{v} \cdot \vec{v})}{dt} = 2\vec{v} \cdot \vec{a} = 2\vec{v} \cdot (\vec{a}_{\parallel} + \vec{a}_{\perp}) = 2\vec{v} \cdot \vec{a}_{\parallel}$$

Le signe de la dérivée temporelle de la norme de la vitesse est donc égal au signe de $\vec{v} \cdot \vec{a}_{\parallel}$.



Seuls les mouvements rectilignes et uniformes sont à vecteur accélération nul. Les mouvements uniformes en général ont seulement leur accélération perpendiculaire à la trajectoire. Il faut avoir en tête le cas du mouvement circulaire et uniforme.

8.3 Vecteur accélération et courbure de la trajectoire

Lors de l'étude des mouvements circulaires, le vecteur accélération est de la forme :

$$\vec{a} = -\frac{v^2}{R}\vec{u}_r + \frac{dv}{dt}\vec{u}_{\theta}.$$

CHAPITRE 13 – CINÉMATIQUE DU POINT

- L'accélération radiale centripète de norme $\frac{v^2}{R}$ est perpendiculaire à la trajectoire et dirigée vers le centre du cercle. Elle est liée au fait que la trajectoire est courbe.
- L'accélération orthoradiale de norme $\left| \frac{dv}{dt} \right|$ est parallèle à la vitesse. Elle est liée aux variations de la norme de la vitesse.

Ce résultat se généralise à tout mouvement plan en décomposant le vecteur accélération en :

- une composante $\vec{a}_\perp$ perpendiculaire à la trajectoire, liée à la courbure de la trajectoire et toujours dirigée vers l'intérieur de la concavité de la trajectoire ;
- une composante $\vec{a}_\parallel$ parallèle à la trajectoire, liée à la variation de la norme de la vitesse.

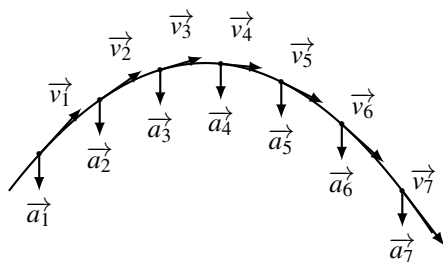


Figure 13.30 – Aux trois premiers points, $\vec{a}_\parallel$ et $\vec{v}$ sont de sens opposés, la norme de la vitesse diminue, le mouvement est décéléré. Aux quatre suivants, $\vec{a}_\parallel$ et $\vec{v}$ sont dans le même sens, la norme de la vitesse augmente, le mouvement est accéléré. Le vecteur accélération pointe vers l'intérieur de la concavité.

9 Étude expérimentale de mouvements

9.1 Généralités

Pour fixer les idées, on étudie les deux mouvements plans présentés en exemple introductif au paragraphe 2.1 : le mouvement parabolique et le mouvement circulaire et uniforme.

La première étape est la réalisation de telles figures. Pour cela, on filme le mouvement avec une caméra ou un appareil photo. On adapte la couleur du fond à la couleur du mobile pour que le mobile ressorte bien sur les clichés. On dispose un mètre sur le fond pour donner une échelle aux photos. On règle alors l'appareil photo et notamment sa fréquence de prise de vue et sa vitesse d'obturation.

La fréquence de prise de vue, qui correspond au nombre de clichés réalisés en une seconde, détermine le nombre de point qui vont matérialiser la trajectoire. Il faut une vingtaine de points pour une trajectoire simple et plus si la trajectoire se complique. La fréquence de prise de vue des appareils courants varie entre 25 et 50 images par seconde, on peut donc enregistrer des mouvements d'environ une demi seconde ou plus. Pour des mouvements plus rapides, il faut un appareil photo plus rapide ou changer de technique².

2. Une autre technique pour obtenir de tels clichés est de se placer dans le noir total, d'éclairer le mouvement à l'aide d'un stroboscope et de réaliser un seul cliché en pose longue. Cela donne de très bons résultats, le traitement d'image est plus simple mais le cliché plus difficile à réaliser.

Le réglage de la vitesse d’obturation permet d’obtenir des prises de vue nettes et d’éviter les « flous de bougé ». En effet, si le déplacement du mobile pendant une ouverture dépasse quelques pixels sur le cliché, la photographie est floue. La vitesse d’obturation par défaut des appareils est rarement assez rapide. Il faut pouvoir régler ce paramètre.



Augmenter la vitesse d’obturation provoque la diminution de la luminosité. Il faut alors éclairer de manière intense pour que le mobile soit visible sur les clichés.

9.2 Étude expérimentale en coordonnées cartésiennes

On étudie le mouvement présenté sur la figure 13.3 du paragraphe 2.1. Cette image a été réalisée par superposition de 27 clichés issus d’un film tourné à 50 images par seconde avec une vitesse d’obturation de $(1/1000)s$. Pour ce type de trajectoires, le repérage cartésien plan s’impose et on cherche les coordonnées x et y du point M à chaque instant. On s’attend à un mouvement d’accélération constante et égale à l’accélération de la pesanteur.

a) Réalisation du film

Dans un premier temps, on pose l’objet immobile devant le fond, on met en place un éclairage suffisant, et on prend un cliché pour vérifier la visibilité du mobile. On passe alors l’appareil en mode « film » et on règle la fréquence de prise de vue et la vitesse d’obturation. On pose l’appareil sur un support fixe et on cadre sur la zone du mouvement que l’on souhaite observer. On réalise alors le film du mouvement que l’on transfère sur ordinateur.



Il ne faut pas oublier de coller un objet de taille connu sur le fond. Cet objet sert à déterminer l’échelle. Ici on a collé un mètre de couture sur le mur blanc du fond.

b) Transformation du film en série d’images

Il faut choisir un logiciel de traitement d’image. Un grand nombre de logiciels existent mais on propose ici d’utiliser ImageJ, logiciel Java du domaine public utilisable sur tout type d’ordinateurs. De bons tutorats et une notice (en anglais) sont téléchargeables pour ce logiciel. La suite des opérations décrites ne prétend pas montrer l’étendue de ses possibilités mais donne un aperçu des opérations à suivre pour obtenir l’image de la figure 13.3 et les caractéristiques du mouvement.

On importe le film (File/Import/AVI ou File/Import/UsingQuickTime)³ et on l’enregistre sous la forme d’une série d’image (File/SaveAs/ImageSequence). On obtient une série de photos prises à intervalle de temps régulier et numérotées⁴.

c) Traitement des images obtenues

On importe la série d’images dans ImageJ (File/Import/ImageSequence) et on affine les réglages (Image/Adjust) pour augmenter le contraste et améliorer la visibilité. On a alors le

3. Pour éviter de saturer la mémoire de son ordinateur, il faut choisir les options Convert to 8-bit grayscale (sauf si la couleur est absolument nécessaire) et Use Virtual Stack. On recadre sur la zone des photos qui nous intéressent (sélectionner un rectangle et Image/Crop) pour réduire encore la taille des fichiers images.

4. On jette les photos sur lesquelles l’objet est hors cadre pour ne garder que les photos intéressantes.

CHAPITRE 13 – CINÉMATIQUE DU POINT

choix entre différents traitements possibles. On a choisi ici de superposer 27 images prises à des intervalles de temps de 20 ms (Image/Stacks/ZProject option MinIntensity) pour obtenir la photo de la figure 13.3.

On définit une échelle sur l’image obtenue en traçant un trait sur l’image superposé au mètre placé en arrière-plan. On indique alors au logiciel à quelle longueur réelle correspond ce trait dans Analyse/SetScale.

Pour finir, on pointe les positions du mobile aux différents instants et on relève les positions (Analyse/Measure) dans un tableau (le logiciel le fait spontanément).

Traitement des résultats On obtient les relevés suivants pour lesquels on calcule les vitesses horizontales et verticales en assimilant la vitesse moyenne entre deux instants consécutifs à la vitesse instantanée par la formule $\dot{x}(t_i) = \frac{x(t_{i+1}) - x(t_i)}{\Delta t}$ où Δt est l’intervalle de temps entre deux photos (ici $(1/50)\text{ s} = 20\text{ ms}$).

$x(\text{mm})$	$y(\text{mm})$	$t(\text{ms})$	$\dot{x}_{(\text{m}\cdot\text{s}^{-1})}$	$\dot{y}_{(\text{m}\cdot\text{s}^{-1})}$	$x(\text{mm})$	$y(\text{mm})$	$t(\text{ms})$	$\dot{x}_{(\text{m}\cdot\text{s}^{-1})}$	$\dot{y}_{(\text{m}\cdot\text{s}^{-1})}$
0,0	0,0	0	1,25	2,15	338,3	211,2	280	1,25	-1,07
25,1	43,0	20	1,16	1,97	363,4	189,7	300	1,25	-1,25
48,3	82,3	40	1,07	1,70	388,4	164,7	320	1,34	-1,43
69,8	116,3	60	1,16	1,52	415,3	136,0	340	1,25	-1,61
93,1	146,8	80	1,25	1,16	440,3	103,8	360	1,16	-1,97
118,1	170,0	100	1,25	1,07	463,6	64,4	380	1,25	-1,97
143,2	191,5	120	1,25	0,81	488,7	25,1	400	1,16	-2,15
168,3	207,6	140	1,25	0,63	511,9	-17,9	420	1,16	-2,51
193,3	220,2	160	1,25	0,36	535,2	-68,0	440	1,16	-2,86
218,4	227,3	180	1,16	0,36	558,5	-125,3	460	1,25	-3,04
241,6	234,5	200	1,34	0,18	583,5	-186,2	480	1,25	-3,40
268,5	238,1	220	1,16	-0,09	608,6	-254,2	500	1,07	-3,49
291,8	236,3	240	1,07	-0,54	630,1	-324,0	520		
313,2	225,5	260	1,25	-0,72					

Tableau 13.1 – Relevé des positions au cours du temps.

Équation de la trajectoire On peut alors tracer la courbe $y = f(x)$ à l’aide d’un tableur et trouver l’équation de la trajectoire (voir figure 13.31).

La trajectoire est une parabole qui justifie l’appellation de mouvement parabolique.

Évolution des vitesses verticales et horizontales On peut également tracer l’évolution temporelle de la vitesse horizontale $\dot{x}$ et celle de la vitesse verticale $\dot{y}$ (voir figure 13.32).

Le signal expérimental sur la vitesse est toujours bruité. Pour cette raison, la dérivée seconde est rarement exploitable sans lissage. Pour obtenir les accélérations horizontales et verticales, on utilise la pente des droites passant par les points de mesure.

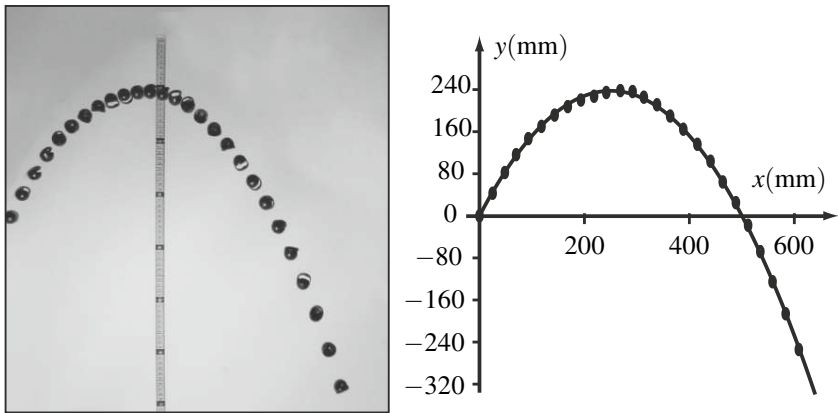


Figure 13.31 – Trajectoire du mobile. À gauche, la trajectoire photographiée ; À droite, la trajectoire reconstituée à partir des pointés réalisés. La courbe passant par les marqueurs ● est une parabole d'équation $y = 1,9x - 0,0038x^2$ avec x et y en mm.

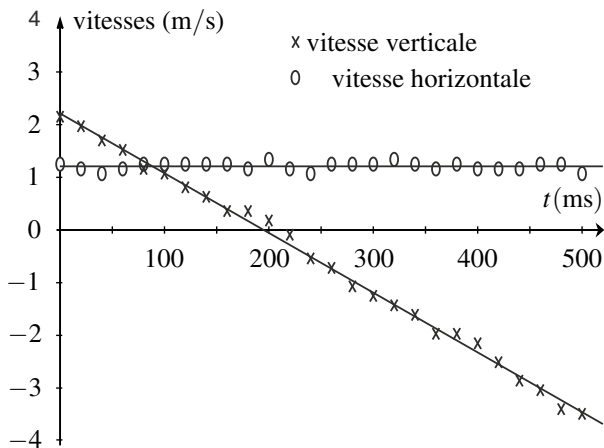


Figure 13.32 – Évolution temporelle de la vitesse horizontale $\dot{x}$ (marqueurs ○) et de la vitesse verticale (marqueurs ×). La droite horizontale passant par les marqueurs ○ a pour équation $\dot{x} = 1,2 \text{ m/s}$. La droite décroissante passant par les marqueurs × a pour équation $\dot{y} = -0,011t + 2,21$ avec t en ms et $\dot{y}$ en $\text{m}\cdot\text{s}^{-1}$.

La vitesse horizontale est quasiment constante. Sa valeur moyenne et son écart type expérimental sur les 26 mesures valent : $\dot{x} = 1,21 \text{ m}\cdot\text{s}^{-1}$ et $s(\dot{x}) = 0,073 \text{ m}\cdot\text{s}^{-1}$. L'incertitude élargie à un niveau de confiance de 95% vaut alors $U(\dot{x}) = 2 \frac{s(\dot{x})}{\sqrt{n}} = 0,03 \text{ m}\cdot\text{s}^{-1}$. On peut donc conclure que le mouvement horizontal est quasiment uniforme à la vitesse :

$$\dot{x} = (1,21 \pm 0,03) \text{ m}\cdot\text{s}^{-1}$$

et que l'accélération horizontale est négligeable.

CHAPITRE 13 – CINÉMATIQUE DU POINT

La vitesse verticale évolue linéairement par rapport au temps $\dot{y} = -0,0112t + 2,21$ avec t en ms et $\dot{y}$ en $\text{m}\cdot\text{s}^{-1}$ soit $\dot{y} = -11,2t + 2,21$ avec t en s et $\dot{y}$ en $\text{m}\cdot\text{s}^{-1}$. L'accélération verticale est donc quasiment constante et égale à :

$$\ddot{y} = (-11,2 \pm 0,4) \text{m}\cdot\text{s}^{-2},$$

où l'intervalle de confiance à 95% sur la pente de la régression linéaire est obtenue à l'aide d'un logiciel de statistique. Le mouvement est bien un mouvement à vecteur accélération constante et égale à :

$$\vec{a} = (-11,2 \pm 0,4) \vec{u}_y \quad \text{en } \text{m}\cdot\text{s}^{-2}.$$

Correction des erreurs de parallaxe On observe une erreur systématique par rapport à la valeur de référence de $\vec{a} = \vec{g} = -9,8 \vec{u}_y$ en $\text{m}\cdot\text{s}^{-2}$. Cette erreur systématique est due à une erreur de parallaxe, courante avec ce type de méthode. Elle peut être corrigée à l'aide des lois de l'optique géométrique.

Pour cela, il faut connaître la distance focale de l'objectif de l'appareil photo (ici $f' = 50 \text{ mm}$), la distance d_1 entre l'objectif et le mètre (ici $d_1 = 2,0 \text{ m}$) et la distance d_2 entre l'objectif et le plan du mouvement de l'objet (ici $d_2 = 1,8 \text{ m}$). Le plan du mouvement n'étant pas confondu avec le plan du mètre, ils sont grossis différemment par l'objectif. d_1 et d_2 étant tous deux grands devant la distance focale de l'objectif, on peut considérer que les images se forment sur le plan focal de l'objectif. Un objet transverse de taille $h_1 = 1 \text{ cm}$, situé dans le plan du mètre, donne une image de taille $h'_1 = h_1 \frac{f'}{d_1}$ sur le capteur de l'appareil photo. Un objet transverse

de taille h_2 , situé dans le plan du mouvement, donne une image de taille $h'_2 = h_2 \frac{f'}{d_2}$ sur le capteur de l'appareil photo. Les images ont la même taille apparente sur la photo lorsque $h'_1 = h'_2$.

Cela signifie qu'une taille $h'_2 = h'_1$ qui correspond à une graduation de la règle sur la photographie correspond à une taille réelle $h_1 = 1 \text{ cm}$ sur le mur et à une taille réelle de $h_2 = h_1 \frac{d_2}{d_1} = 0,9 \text{ cm}$ dans le plan du mouvement. Les déplacements apparents sur la photo sont surestimés dans un rapport $\frac{d_1}{d_2} = 2,0/1,8 = 1,1$. Ici, toutes les mesures de longueurs effectuées doivent être multipliées par 0,9 et par conséquent l'accélération également. On trouve alors que le mouvement est à vecteur accélération constant et égal à :

$$\vec{a} = (-10,1 \pm 0,4) \vec{u}_y \quad \text{en } \text{m}\cdot\text{s}^{-2},$$

avec un niveau de confiance proche de 95%. Cette fois, après correction de l'erreur systématique, on trouve bien une valeur de l'accélération cohérente avec la valeur tabulée.

9.3 Étude expérimentale en coordonnées polaires

On étudie maintenant le mouvement circulaire présenté sur la figure 13.4 au paragraphe 2.1. On s’attend à trouver un mouvement circulaire et uniforme de vitesse angulaire constante. C’est ce qu’on va chercher à vérifier.

Le cliché est obtenu en suivant la même démarche que précédemment et en superposant 23 images prises à des intervalles de temps de 40 ms (25 images par seconde) avec une vitesse d’obturation de (1/500)s.

Le mouvement est clairement circulaire. On pointe cette fois les angles polaires repérant le mobile aux différents instants à l’aide du pointeur d’angle (AngleTool) sur la figure 13.4. On relève les positions angulaires (Analyse/Measure) pour chaque position du mobile. Avec ImageJ, les angles sont donnés en degré sur l’intervalle [0°, 180°]. Il faut corriger les angles compris entre 180° et 360° pour obtenir la coordonnée angulaire θ définie sur la figure 13.4.

On calcule ensuite les vitesses angulaires en assimilant la vitesse angulaire moyenne entre deux instants consécutifs à la vitesse instantanée par la formule $\dot{\theta}_{(t_i)} = \frac{\theta(t_{i+1}) - \theta(t_i)}{\Delta t}$ où Δt est l’intervalle de temps entre deux photos (ici (1/25)s = 40 ms).

$\theta(^{\circ})$	$\theta(\text{rad})$	$t(\text{ms})$	$\dot{\theta}_{(\text{rad}\cdot\text{s}^{-1})}$	$\theta(^{\circ})$	$\theta(\text{rad})$	$t(\text{ms})$	$\dot{\theta}_{(\text{rad}\cdot\text{s}^{-1})}$
7	0,13	0	7,21	197	3,43	480	6,88
24	0,41	40	6,73	212	3,71	520	7,00
39	0,68	80	6,87	228	3,99	560	7,06
55	0,96	120	6,74	245	4,27	600	7,16
70	1,23	160	6,78	261	4,55	640	6,93
86	1,50	200	6,72	277	4,83	680	7,08
101	1,77	240	6,81	293	5,11	720	6,76
117	2,04	280	6,89	309	5,38	760	6,74
133	2,32	320	6,83	324	5,65	800	6,81
148	2,59	360	7,01	340	5,93	840	6,94
164	2,87	400	6,79	356	6,20	880	
180	3,14	440	7,22				

Tableau 13.2 – Relevé des positions angulaires au cours du temps.

La vitesse angulaire est quasiment constante. Sa valeur moyenne et son écart type expérimental sur les 23 mesures valent $\dot{\theta} = 6,91 \text{ rad}\cdot\text{s}^{-1}$ et $s(\dot{\theta}) = 0,16 \text{ rad}\cdot\text{s}^{-1}$. L’incertitude élargie à un niveau de confiance de 95% vaut alors $U(\dot{\theta}) = 2\frac{s(\dot{\theta})}{\sqrt{n}} = 0,07 \text{ rad/s}$ et on peut conclure que le mouvement est circulaire et uniforme à la vitesse angulaire constante et égale à :

$$\dot{\theta} = (6,91 \pm 0,07) \text{ rad}\cdot\text{s}^{-1}$$

avec un niveau de confiance proche de 95%.

CHAPITRE 13 – CINÉMATIQUE DU POINT

SYNTHÈSE

SAVOIRS

- systèmes de coordonnées cartésiennes, polaires, cylindriques et sphériques
- bases de projection associées
- coordonnées d'un point
- composantes d'un vecteur
- composantes du vecteur position dans les différents systèmes de coordonnées
- vitesse moyenne ou instantanée, accélération moyenne ou instantanée
- lien entre l'évolution de la norme du vecteur et la direction du vecteur accélération
- savoir que le vecteur accélération est dirigé dans la concavité de la trajectoire

SAVOIR-FAIRE

- déterminer les vecteurs déplacements infinitésimaux dans les différents systèmes de coordonnées et en déduire le vecteur vitesse
- projeter un vecteur sur une base
- dériver les vecteurs de la base polaire mobile
- déterminer les vecteurs vitesse et accélération instantanées en coordonnées cartésiennes et cylindriques par dérivation du vecteur position
- choisir un système de coordonnées adapté au problème posé
- exprimer la vitesse et la position en fonction du temps et déterminer la trajectoire en coordonnées cartésiennes d'un mouvement de vecteur-accélération constante
- exprimer les vecteurs position, vitesse et accélération en coordonnées polaires planes d'un mouvement circulaire et uniforme

MOTS-CLÉS

- | | | |
|---------------|---------------|----------------------|
| • référentiel | • coordonnées | • mouvement uniforme |
| • repère | • composantes | |
| • base | • trajectoire | |

S'ENTRAÎNER

13.1 Test d'accélération d'une voiture (★)

Une voiture est chronométrée pour un test d'accélération en ligne droite avec départ arrêté (vitesse initiale nulle).

- 1. Elle est chronométrée à 26,6 s au bout d'une distance $D = 180$ m. Déterminer l'accélération (supposée constante) et la vitesse atteinte à la distance D .
- 2. Quelle est alors la distance d'arrêt pour une décélération de $7 \text{ m}\cdot\text{s}^{-2}$.

13.2 Interpellation pour vitesse excessive (★)

Un conducteur roule à vitesse constante v_0 sur une route rectiligne. Comme il est en excès de vitesse à $100 \text{ km}\cdot\text{h}^{-1}$, un gendarme à moto démarre à l'instant où la voiture passe à sa hauteur et accélère uniformément. Le gendarme atteint la vitesse de $90 \text{ km}\cdot\text{h}^{-1}$ au bout de 10 s.

- 1. Quel sera le temps nécessaire au motard pour rattraper la voiture ?
- 2. Quelle distance aura-t-il parcourue ?
- 3. Quelle vitesse aura-t-il alors atteinte ?

13.3 Satellite géostationnaire (★)

Un satellite géostationnaire est en mouvement circulaire uniforme autour de la Terre. Il ressent une accélération $a = g_0 \left(\frac{R}{r}\right)^2$, où $R = 6400 \text{ km}$ est le rayon de la terre, $g_0 = 9,8 \text{ m}\cdot\text{s}^{-2}$ et r est le rayon de l'orbite. La période de révolution du satellite est égale à la période de rotation de la terre sur elle-même.

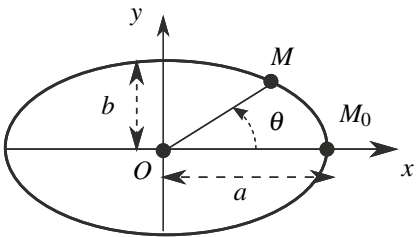
- 1. Calculer la période T de rotation de la Terre en secondes, puis la vitesse angulaire Ω .
- 2. Déterminer l'altitude de l'orbite géostationnaire.
- 3. Déterminer sa vitesse sur sa trajectoire et calculer sa norme.

13.4 Mouvement sur une ellipse (★★)

Un point M se déplace sur une ellipse d'équation cartésienne $\left(\frac{x}{a}\right)^2 + \left(\frac{y}{b}\right)^2 = 1$. On note θ l'angle que fait $\overrightarrow{OM}$ avec l'axe (Ox) . Les coordonnées de M peuvent s'écrire :

$$\begin{cases} x(t) = \alpha \cos(\omega t + \phi) \\ y(t) = \beta \sin(\omega t + \psi), \end{cases}$$

où l'on suppose que ω est une constante.

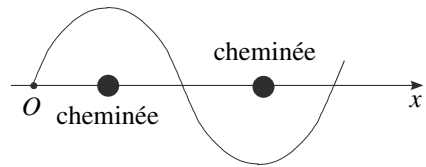


CHAPITRE 13 – CINÉMATIQUE DU POINT

1. À $t = 0$, le mobile est en M_0 . Déterminer α , ϕ et ψ .
2. Des autres données, déduire β .
3. Déterminer les composantes de la vitesse $(\dot{x}, \dot{y})$ et de l'accélération $(\ddot{x}, \ddot{y})$.
4. Montrer que l'accélération est de la forme $\vec{a} = -\omega^2 \overrightarrow{OM}$. Commenter.

13.5 Slalom entre des cheminées ; épisode 2 de la guerre des étoiles (★★)

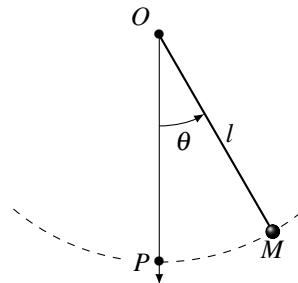
Dans cet épisode de la Guerre des étoiles, on peut assister à une course poursuite de « speeder » entre des cheminées d'usine. On suppose que le véhicule suit une trajectoire sinusoïdale de slalom entre les cheminées alignées selon l'axe (Ox) . Elles sont espacées d'une distance $L = 200$ m.



1. Le véhicule conserve une vitesse v_0 constante selon (Ox) et met $t_t = 12$ s pour revenir sur l'axe après la sixième cheminée. En déduire la vitesse v_0 . Faire l'application numérique.
2. Déterminer l'amplitude de la sinusoïde pour que l'accélération reste inférieure à $10g$ en valeur absolue, avec $g = 9,8 \text{ m}\cdot\text{s}^{-2}$. Que penser des valeurs obtenues ?

13.6 Etude cinématique du pendule simple (★★)

Le mouvement d'un point M accroché à un fil de longueur l dont l'autre extrémité est fixée en un point O s'inscrit sur une portion de cercle de centre O et de rayon l . On repère alors le point M dans le référentiel $\mathcal{R}$ par sa coordonnée angulaire θ définie sur la figure ci-contre et on observe des oscillations pendulaires. Lorsque les oscillations sont de faibles amplitudes, on observe que l'angle polaire θ est tel que $\theta(t) = \theta_0 \sin \omega t$ en choisissant pour origine des temps l'instant où M passe au point P .



1. Définir sur un schéma la base locale de projection polaire. Donner l'expression des vecteurs position, vitesse et accélération sur cette base.
2. Définir la base cartésienne de projection et donner l'expression du vecteur position sur cette base.
3. Lorsque $\theta_0 \ll \frac{\pi}{2}$, on a $\theta(t) \ll \frac{\pi}{2}$ à tout instant et on peut utiliser les approximations suivantes : $\cos \theta \simeq 1$ et $\sin \theta \simeq \theta$.

a. Dans ces conditions, déterminer les composantes des vecteurs vitesses et accélération en coordonnées cartésiennes.

b. En changeant l'origine du repère cartésien pour la placer en P , montrer que l'on obtient alors une relation remarquable entre les vecteurs $\overrightarrow{PM}$ et $\vec{a}$. Quelle est-elle ? Commenter.

13.7 Mouvement hélicoïdal (★)

Un point matériel se déplace le long d’une hélice circulaire. Son mouvement est donné en coordonnées cylindriques par :

$$\begin{cases} r &= R \\ \theta &= \omega t \\ z &= ht. \end{cases}$$

où R, h et ω sont des constantes.

- 1. Donner l’expression de la vitesse.
- 2. En déduire que le module de la vitesse est constant.
- 3. Exprimer l’accélération.

APPROFONDIR

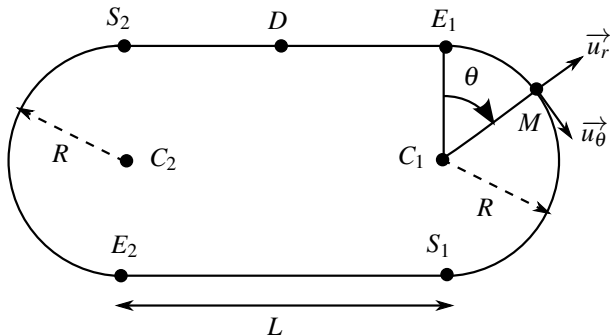
13.8 Courses entre véhicules radio-commandés (★)

Deux modèles réduits de voitures radio-commandées ont des performances différentes : le premier a une accélération de $4,0 \text{ m}\cdot\text{s}^{-2}$, le second de $5,0 \text{ m}\cdot\text{s}^{-2}$. Cependant l’utilisateur de la première voiture a plus de réflexes que celui de la seconde, ce qui lui permet de la faire démarrer $1,0 \text{ s}$ avant le second.

- 1. Déterminer le temps nécessaire au deuxième véhicule pour rattraper l’autre.
- 2. Les deux modèles réduits participent à des courses de 100 m et 200 m . Est-il possible que le perdant du 100 m prenne sa revanche au 200 m ?
- 3. Calculer pour les deux courses la vitesse finale de chacun des véhicules.

13.9 Parcours d’un cycliste sur un vélodrome (★★)

On s’intéresse à un cycliste, considéré comme un point matériel M , qui s’entraîne sur un vélodrome constitué de deux demi-cercles reliés par deux lignes droites (figure ci-dessous). Données : $L = 62 \text{ m}$ et $R = 20 \text{ m}$. Le cycliste part de D avec une vitesse nulle.



CHAPITRE 13 – CINÉMATIQUE DU POINT

1. Il exerce un effort constant ce qui se traduit par une accélération constante a_1 jusqu'à l'entrée E_1 du premier virage. Calculer le temps t_{E_1} de passage en E_1 ainsi que la vitesse V_{E_1} en fonction de a_1 et L .
2. Dans le premier virage, le cycliste a une accélération tangentielle (suivant $\vec{u}_\theta$) constante et égale à a_1 . Déterminer le temps t_{S_1} de passage en S_1 ainsi que la vitesse v_{S_1} en fonction de a_1 , L et R .
3. De même, en considérant l'accélération tangentielle constante tout au long du premier tour et égale à a_1 , déterminer les temps t_{E_2} , t_{S_2} et t_D (après un tour), ainsi que les vitesses correspondantes.
4. La course s'effectue sur quatre tours (1 km) mais on ne s'intéresse donc qu'au premier effectué en $t_1 = 18,155$ s (Temps du britannique Chris Hoy aux Championnats du monde de 2007). Déterminer la valeur de l'accélération a_1 ainsi que la vitesse atteinte en D . La vitesse mesurée sur piste est d'environ $60 \text{ km}\cdot\text{h}^{-1}$. Que doit-on modifier dans le modèle pour se rapprocher de la réalité ?

13.10 Mouvement de l'extrémité d'une barre, d'après ENAC 2003 (★★★)

Dans le référentiel $\mathcal{R}$ de repère $(O, \vec{u}_x, \vec{u}_y, \vec{u}_z)$ défini sur les figures 13.33 et 13.34, une barre rectiligne AB de longueur $2b$ se déplace de sorte que :

- son extrémité A se trouve sur le demi-axe positif (Oz) ;
- son extrémité B décrit le demi-cercle du plan (xOy) de centre $I(0, b, 0)$ et de rayon b , à la vitesse angulaire ω constante et positive. à l'instant $t = 0$, B se trouve en O .

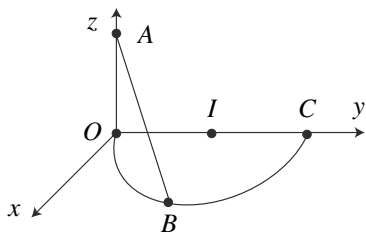


Figure 13.33

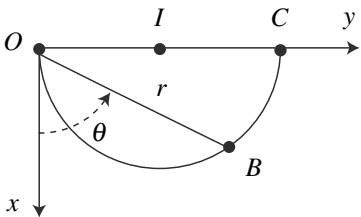


Figure 13.34

L'exercice ne nécessite aucune connaissance de mécanique du solide.

1. Déterminer la durée Δt du mouvement.
2. On note φ l'angle $(\vec{IO}, \vec{IB})$, déterminer une relation simple entre φ et θ .
3. Établir les expressions des coordonnées polaires r et θ de B au cours du temps t (figure 13.34).
4. Déterminer l'angle $\alpha = (\vec{AO}, \vec{AB})$ en fonction de ω et t .
5. Décrire le mouvement de la barre entre l'instant initial et l'instant final.
6. Calculer les coordonnées cartésiennes X , Y et Z du milieu J de la barre.
7. Déterminer la vitesse $\vec{v}$ et l'accélération $\vec{a}$ de J , ainsi que leurs normes.

CORRIGÉS

13.1 Test d'accélération d'une voiture

1. Le mouvement est rectiligne à accélération constante $\vec{a}_0$. On choisit l'axe (Ox) (vecteur $\vec{u}_x$) comme axe du mouvement et le point O comme point de départ. L'intégration de l'accélération donne :

$$\vec{a} = \vec{a}_0 = a_0 \vec{u}_x \Rightarrow \vec{v} = (a_0 t + C) \vec{u}_x$$

or la vitesse initiale ($t = 0$) est nulle donc $C = 0$. L'abscisse x est alors :

$$x = \frac{a_0 t^2}{2} + C' \text{ avec } x(0) = 0 \text{ donc } C' = 0$$

Si on note $t_D = 26,6 \text{ s}$, $a_0 = 2D/t_D^2 = 2 \times 180/26,6^2 = 0,509 \text{ m}\cdot\text{s}^{-2}$. On obtient ensuite la vitesse à $t = t_D$: $v_D = v_0 t_D = 13,5 \text{ m}\cdot\text{s}^{-1}$.

2. Pour la phase de freinage, on peut changer les origines (temps et espace); la nouvelle vitesse initiale est alors v_D . On note $a_f = -7 \text{ m}\cdot\text{s}^{-2}$ l'accélération lors de cette phase ($a_f < 0$ car il y a freinage). Les intégrations de l'accélération amènent à :

$$v = \dot{x} = a_f t + v_D \text{ et } x = \frac{a_f t^2}{2} + v_D t.$$

L'arrêt a lieu pour $v = 0$ soit $t = t_A = v_D/a_f = 1,93 \text{ s}$ d'où une distance de freinage :

$$x(t_A) = \frac{a_f}{2} \left(\frac{v_D}{a_f} \right)^2 + v_D \frac{v_D}{a_f} = 1,5 \frac{v_D^2}{a_f} = 13 \text{ m}.$$

13.2 Interpellation pour vitesse excessive

1. Soit (Ox) l'axe le long duquel ont lieu les mouvements de la voiture et de la moto. La voiture, repérée par son abscisse x_V , est en mouvement rectiligne et uniforme à la vitesse $\dot{x}_V = v_0$. Sa position s'écrit $x_V = v_0 t$ en prenant l'origine O du repère d'espace à la position commune de la voiture et de la moto à $t = 0$.

La moto, repérée par son abscisse x_M , est en mouvement rectiligne et uniformément accélérée avec $\ddot{x}_M = a_0$. Comme dans l'exercice précédent (les conditions initiales sont identiques), on en déduit sa vitesse $\dot{x}_M = a_0 t$ et sa position $x_M = \frac{1}{2} a_0 t^2$ par intégration par rapport au temps.

Le motard rattrape la voiture lorsque les positions de la voiture et de la moto sont à nouveau identiques, ce qui revient à résoudre $x_V(t) = x_M(t)$ soit $v_0 t = \frac{1}{2} a_0 t^2$. Les solutions sont $t = 0$

qui n'est bien évidemment pas la solution cherchée et $t_0 = \frac{2v_0}{a_0} = 22,2 \text{ s}$. Pour effectuer l'application numérique, on transforme la vitesse en $\text{m}\cdot\text{s}^{-1}$ soit $v_0 = 100 \text{ km}\cdot\text{h}^{-1} = 27,8 \text{ m}\cdot\text{s}^{-1}$ et on détermine l'accélération par le fait qu'au bout de $t_1 = 10 \text{ s}$, la moto atteint une vitesse $v_1 = 90 \text{ km}\cdot\text{h}^{-1} = 25 \text{ m}\cdot\text{s}^{-1}$ soit $a_0 = \frac{v_1}{t_1} = 2,5 \text{ m}\cdot\text{s}^{-2}$.

CHAPITRE 13 – CINÉMATIQUE DU POINT

2. La distance parcourue est alors $d = v_0 t_0 = \frac{1}{2} a_0 t_0^2 = 617 \text{ m}$.
3. La moto possède alors une vitesse $v_M = a_0 t_0 = 55 \text{ m}\cdot\text{s}^{-1} = 198 \text{ km}\cdot\text{h}^{-1}$.

13.3 Satellite géostationnaire

1. La Terre effectue un tour sur elle-même en une période d'environ $T = 24 \text{ h}$ soit $T = 86,4 \cdot 10^3 \text{ s}$. Sa vitesse angulaire est $\Omega = 2\pi/T = 7,27 \cdot 10^{-5} \text{ rad}\cdot\text{s}^{-1}$.
2. On note $\overrightarrow{TM} = R\overrightarrow{u}_r$ le vecteur position du satellite en coordonnées polaires (r, θ) d'origine T , le centre de la Terre. La période de révolution du satellite coïncide avec la période de rotation de la Terre sur elle-même donc sa vitesse angulaire $\dot{\theta} = 2\pi/T$ est égale à $\Omega = 2\pi/T$. Il faut donc déterminer $\dot{\theta}$ connaissant l'accélération.
- Le mouvement est circulaire et uniforme de rayon r . L'accélération est :

$$\vec{a} = -\frac{v^2}{r}\overrightarrow{u}_r = -r\dot{\theta}^2\overrightarrow{u}_r.$$

Avec l'expression de l'énoncé, on déduit :

$$\Omega = \dot{\theta} = \sqrt{\frac{g_0 R^2}{r^3}} \Rightarrow r = \left(\frac{g_0 R^2}{\Omega^2}\right)^{1/3} = 42,2 \cdot 10^3 \text{ km}.$$

On retient en général trente six mille kilomètres pour l'orbite géostationnaire ; il s'agit de l'altitude à partir du sol, obtenue en soustrayant le rayon de la Terre au résultat précédent.

3. $\vec{v} = r\Omega\overrightarrow{u}_\theta$ et $\|\vec{v}\| = 42,2 \cdot 10^6 \times 7,27 \cdot 10^{-5} = 3,1 \cdot 10^3 \text{ m}\cdot\text{s}^{-1} = 11 \cdot 10^3 \text{ km}\cdot\text{h}^{-1}$.

13.4 Mouvement sur une ellipse

1. On écrit x et y à $t = 0$ soit : $x(0) = \alpha \cos \phi$ et $y(0) = \beta \cos \psi$.
L'énoncé précise que $x(0) = a$ et $y(0) = 0$; on en déduit $\alpha = a$, $\phi = 0$ et $\psi = 0$.
2. On remplace x et y par leurs expressions dans l'équation de l'ellipse et on obtient :

$$\cos^2 \omega t + \left(\frac{\beta}{b}\right)^2 \sin^2 \omega t = 1.$$

La seule possibilité pour que l'équation soit vérifiée à chaque instant est $\beta = b$.

3. Pour obtenir la vitesse et l'accélération on dérive par rapport au temps, sachant que ω est constante :

$$\begin{cases} \dot{x} = -a\omega \sin \omega t \\ \dot{y} = b\omega \cos \omega t \end{cases} \quad \text{et} \quad \begin{cases} \ddot{x} = -a\omega^2 \cos \omega t \\ \ddot{y} = -b\omega^2 \sin \omega t \end{cases}$$

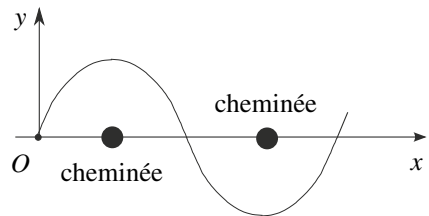
4. On remarque que $\ddot{x} = -\omega^2 x$ et $\ddot{y} = -\omega^2 y$, donc :

$$\vec{a} = -\omega^2 \overrightarrow{OM}.$$

Il s'agit de l'équation caractéristique d'un oscillateur harmonique, mais à deux dimensions dans le plan (xOy) .

13.5 Slalom entre des cheminées ; épisode 2 de la guerre des étoiles

On choisit un repère cartésien (O, x, y) et on note $\vec{u}_x$ et $\vec{u}_y$ les vecteurs de base. L'équation de la trajectoire, telle qu'elle apparaît sur la figure, s'écrit $y = a \sin\left(\frac{2\pi x}{L}\right)$. Il faut déterminer a .



- 1. Selon (Ox) , $\dot{x} = v_0$, soit $x = v_0 t$ puisqu'à $t = 0$, $x = 0$. Le véhicule revient sur l'axe après la sixième cheminée, ayant parcouru une distance $D = 6L = 1200$ m, la vitesse v_0 est donc $v_0 = 6L/t = 100 \text{ m}\cdot\text{s}^{-1}$ ou $360 \text{ km}\cdot\text{h}^{-1}$. La vitesse « au compteur » est encore plus grande car le speeder ne se déplace pas en ligne droite.
- 2. En remplaçant x par son expression dans y , on obtient $y = a \sin\left(\frac{2\pi v_0 t}{L}\right)$. L'accélération a pour composantes $\ddot{x} = \dot{v}_0 = 0$ et $\ddot{y}$. Il vient :

$$\dot{y} = a \frac{2\pi v_0}{L} \cos \frac{2\pi v_0 t}{L} \text{ et } \ddot{y} = -a \left(\frac{2\pi v_0}{L}\right)^2 \sin \frac{2\pi v_0 t}{L}$$

Le sinus est compris dans l'intervalle $[-1, 1]$, donc pour que l'accélération reste inférieure à $10g$ en valeur absolue, il faut :

$$a \left(\frac{2\pi v_0}{L}\right)^2 \leq 10g \Rightarrow a \leq 10g \left(\frac{L}{2\pi v_0}\right)^2$$

L'application numérique donne : $a = 9,9$ m.
Il faut être excellent pilote pour passer aussi près des cheminées à cette vitesse. Il faut supporter de plus une accélération importante ; à titre de comparaison lors du catapultage d'un avion depuis un porte-avion, l'accélération est environ $5g$. Mais dans la cas étudié, il s'agit de « Chevaliers Jedi » !

13.6 Etude cinématique du pendule simple

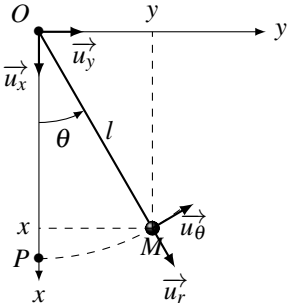
On étudie le mouvement de M dans le référentiel $\mathcal{R}$ dans lequel la droite OP et le plan du mouvement sont fixes, le vecteur position s'écrit différemment selon la base de projection.

- 1. $\vec{OM} = l\vec{u}_r$ sur la base polaire définie ci-contre. On tire alors :

$$\begin{cases} \vec{v} &= l\dot{\theta}\vec{u}_\theta = l\omega\theta_0 \cos \omega t \vec{u}_\theta \\ \vec{a} &= -l\dot{\theta}^2\vec{u}_r + l\ddot{\theta}\vec{u}_\theta \\ &= -l\omega^2\theta_0^2 \cos^2 \omega t \vec{u}_r + l\omega^2\theta_0 \sin \omega t \vec{u}_\theta. \end{cases}$$

- 2. Sur la base cartésienne définie ci-contre, on a alors :

$$\vec{OM} = x\vec{u}_x + y\vec{u}_y = l \cos \theta \vec{u}_x + l \sin \theta \vec{u}_y,$$



CHAPITRE 13 – CINÉMATIQUE DU POINT

soit :

$$\overrightarrow{OM} = l \cos(\theta_0 \sin \omega t) \vec{u}_x + l \sin(\theta_0 \sin \omega t) \vec{u}_y.$$

Cette expression est délicate à dériver, ce qui explique que les coordonnées polaire soient en général préférables pour étudier ce problème.

3. Pour les oscillations de faible amplitude, on obtient :

a.

$$\begin{cases} \overrightarrow{OM} &= l \theta_0 \sin \omega t \vec{u}_y + l \vec{u}_x \\ \vec{v} &= l \omega \theta_0 \cos \omega t \vec{u}_y \\ \vec{a} &= -l \omega^2 \theta_0 \sin \omega t \vec{u}_y. \end{cases}$$

b. Pour déplacer l'origine du repère cartésien en P , on utilise la relation de Chasles $\overrightarrow{PM} = \overrightarrow{PO} + \overrightarrow{OM} = -l \vec{u}_x + \overrightarrow{OM}$ et on remarque que le vecteur $\overrightarrow{PO}$ est un vecteur constant dont les dérivées temporelles sont identiquement nulles. On obtient alors :

$$\begin{cases} \overrightarrow{PM} = Y \vec{u}_y &= l \theta_0 \sin \omega t \vec{u}_y \\ \vec{v} = \dot{Y} \vec{u}_y &= l \omega \theta_0 \cos \omega t \vec{u}_y \\ \vec{a} = \ddot{Y} \vec{u}_y &= -l \omega^2 \theta_0 \sin \omega t \vec{u}_y. \end{cases}$$

Le mouvement du point M vérifie alors $\ddot{Y} + \omega^2 Y = 0$. Il s'agit du mouvement d'un oscillateur harmonique d'axe (Py) , centré sur P , d'amplitude $l\theta_0$ et de pulsation ω .

13.7 Mouvement helicoidal

1. On utilise l'expression générale de la vitesse en coordonnées cylindriques :

$$\vec{v} = \dot{r} \vec{u}_r + r \dot{\theta} \vec{u}_\theta + \dot{z} \vec{u}_z = R \omega \vec{u}_\theta + h \vec{u}_z.$$

2. En calculant la norme de l'expression précédente, on obtient $v = \sqrt{(R\omega)^2 + h^2}$.

3. On utilise l'expression générale de l'accélération en coordonnées cylindriques :

$$\vec{a} = (\ddot{r} - r \dot{\theta}^2) \vec{u}_r + (2\dot{r} \dot{\theta} + r \ddot{\theta}) \vec{u}_\theta + \ddot{z} \vec{u}_z = -R \omega^2 \vec{u}_r.$$

13.8 Courses entre véhicules radio-commandés

1. Pour le premier, on a une accélération a_1 constante soit après une double intégration l'équation horaire suivante $x_1 = \frac{1}{2} a_1 t^2$.

Pour le second, on a de même $x_2 = \frac{1}{2} a_2 (t - t_0)^2$ en tenant compte du retard t_0 au départ.

On cherche l'instant t pour lequel on a $x_1(t) = x_2(t)$. La résolution de l'équation du second degré $(a_1 - a_2)t^2 - 2a_2 t_0 t + a_2 t_0^2 = 0$ donne $t = \frac{a_2 \pm \sqrt{a_1 a_2}}{a_2 - a_1} t_0$.

L'application numérique donne $t = 9,5$ s et $t = 0,5$ s. La seconde solution n'est pas possible puisque l'un des deux n'est pas encore parti !

Il faut donc 8,5 s pour que le second rattrape le premier.

2. Il suffit de reporter la valeur de t dans l'une ou l'autre des équations horaires. On trouve $x = 181$ m. On peut vérifier qu'on trouve bien la même valeur avec les deux équations horaires. La deuxième voiture arrive après la première au 100 m mais la dépasse avant l'arrivée au 200 m.

3. Quant aux vitesses, elles sont $v_1 = 38 \text{ m}\cdot\text{s}^{-1}$ et $v_2 = 42,5 \text{ m}\cdot\text{s}^{-1}$.

13.9 Parcours d'un cycliste sur un vélodrome

1. De D à E_1 , le mouvement est rectiligne uniformément accéléré. À l'instant t , la vitesse est $v = a_1 t$ puisque la vitesse initiale est nulle et la distance parcourue est $d = a_1 t^2 / 2$, d'où :

$$t_{E_1} = \sqrt{\frac{2DE_1}{a_1}} = \sqrt{\frac{L}{a_1}} \quad \text{et} \quad v_{E_1} = a_1 t_{E_1} = \sqrt{a_1 L}.$$

2. Dans le virage, le mouvement est circulaire, donc l'expression de l'accélération est :

$$\vec{a} = -R\dot{\theta}^2 \vec{u}_r + R\ddot{\theta} \vec{u}_\theta.$$

D'après l'énoncé, on a $R\ddot{\theta} = a_1$, on en déduit en intégrant, $\dot{\theta} = a_1 t / R + cte$. On prend une nouvelle origine des temps en E_1 , alors $cte = \dot{\theta}_{E_1} = v_{E_1} / R = \sqrt{a_1 L} / R$. En prenant $\theta = 0$ en E_1 , on détermine :

$$\theta(t) = \frac{a_1 t^2}{2R} + \frac{\sqrt{a_1 L}}{R} t.$$

On obtient alors le temps de sortie du virage pour $\theta = \pi$, ce qui revient à résoudre l'équation du second degré :

$$t^2 + 2\sqrt{\frac{L}{a_1}} t - \frac{2\pi R}{a_1} = 0.$$

La racine positive est : $t_1 = -\sqrt{\frac{L}{a_1}} + \sqrt{\frac{L+2\pi R}{a_1}}$, et la vitesse en S_1 est $v_{S_1} = R\dot{\theta}(S_1)$ soit :

$$v_{S_1} = a_1 t_1 + \sqrt{a_1 L} = \sqrt{(L+2\pi R)a_1}.$$

Le temps t_{S_1} demandé dans l'énoncé est à partir de l'origine en D soit :

$$t_{S_1} = t_{E_1} + t_1 = \sqrt{\frac{L+2\pi R}{a_1}}.$$

3. La distance entre S_1 et E_2 est égale à L . On change de nouveau l'origine des temps ($t = 0$ en S_1) et on note x l'abscisse à partir de S_1 . L'intégration de $\ddot{x} = a_1$ avec une vitesse initiale en S_1 donne : $x = \frac{a_1 t^2}{2} + v_{S_1} t$. On doit donc résoudre l'équation suivante pour calculer t_{E_2} :

$$t^2 + 2\frac{v_{S_1}}{a_1} t - \frac{2L}{a_1} = 0$$

La racine positive est : $t'_1 = -\sqrt{\frac{L+2\pi R}{a_1}} + \sqrt{\frac{3L+2\pi R}{a_1}}$ d'où :

$$t_{E_2} = t'_1 + t_{S_1} = \sqrt{\frac{3L+2\pi R}{a_1}},$$

CHAPITRE 13 – CINÉMATIQUE DU POINT

et $v_{E_2} = a_1 t'_1 + v_{S_1}$ soit :

$$v_{E_2} = \sqrt{(3L + 2\pi R)a_1}.$$

En ce qui concerne t_{S_2} et v_{S_2} , le raisonnement est le même que pour le premier virage, seule la constante d'intégration pour θ change soit $\theta = \frac{a_1 t^2}{2R} + \frac{v_{E_2} t}{R}$, en prenant l'origine des temps et l'origine des angles en S_2 . L'équation donnant le temps de sortie en S_2 est (pour $\theta = \pi$) :

$$t^2 + 2\frac{v_{E_2}}{a_1}t - \frac{2R\pi}{a_1} = 0.$$

La solution est : $t'_2 = -\sqrt{\frac{3L + 2\pi R}{a_1}} + \sqrt{\frac{3L + 4\pi R}{a_1}} \Rightarrow t_{S_2} = t_{E_2} + t'_2 = \sqrt{\frac{3L + 4\pi R}{a_1}}$ et $v_{S_2} = \sqrt{(3L + 4\pi R)a_1}$.

Enfin pour le calcul de t_D le raisonnement est le même que pour t_{E_2} , en changeant L en $L/2$ et v_{S_1} en v_{S_2} ce qui conduit à l'équation (en prenant les origines de x et t en S_2) :

$$t^2 + 2\frac{v_{S_2}}{a_1}t - \frac{L}{a_1} = 0,$$

dont la solution positive est : $t'_3 = -\sqrt{\frac{3L + 4\pi R}{a_1}} + 2\sqrt{\frac{L + \pi R}{a_1}} \Rightarrow t_D = t_{S_2} + t'_3 = 2\sqrt{\frac{L + \pi R}{a_1}}$ et $v_D = 2\sqrt{(L + \pi R)a_1}$.

4. L'accélération a_1 est donnée par $a_1 = 4(L + \pi R)/t_D^2$ soit $a_1 = 1,515 \text{ m.s}^{-2}$. On obtient $v_D = 27,5 \text{ m.s}^{-1}$ ce qui fait 99 km.h^{-1} au lieu de 60 km.h^{-1} . On peut penser que l'accélération n'est pas constante et qu'elle diminue au cours du temps.

13.10 Mouvement de l'extrémité d'une barre, d'après ENAC 2003

1. Le point a un mouvement circulaire uniforme de centre I . Il décrit un tour en une période $T = \frac{2\pi}{\omega}$ et un demi-tour de O à C en une demi-période soit $\Delta t = \frac{T}{2} = \frac{\pi}{\omega}$.

2. La relation entre φ et θ est une propriété des cercles : $\varphi = 2\theta$. On peut la démontrer en considérant le triangle isocèle OIB . Si l'on note β les angles (IOB) et (IBO) , alors dans le triangle OIB , $2\beta + \varphi = \pi$, or $\theta = \pi/2 - \beta$, d'où le résultat.

3. L'angle φ correspond à l'angle polaire du mouvement circulaire de centre I donc $\dot{\varphi} = \omega$. Puisque ω est constant, on a $\varphi = \omega t + cte$, or à $t = 0$, B est en O donc $\varphi(0) = 0$ et $cte = 0$. On en déduit $\theta = \omega t/2$.

En ce qui concerne $r = OB$, dans le triangle OIB , on peut écrire $OB = 2b \sin(\varphi/2)$, soit $r = 2b \sin \frac{\omega t}{2}$.

4. Les coordonnées de A dans la base $(\vec{u}_r, \vec{u}_\theta, \vec{u}_z)$ sont $(0, 0, 2b \cos \alpha)$; celles de B sont $(r, 0, 0)$ d'où $\vec{AB} = r\vec{u}_r - 2b \cos \alpha \vec{u}_z$, or la norme de $\vec{AB}$ est $2b$. On en déduit :

$$4b^2 \sin^2 \frac{\omega t}{2} + 4b^2 \cos^2 \alpha = 4b^2 \Rightarrow \cos^2 \alpha = 1 - \sin^2 \frac{\omega t}{2}$$

comme $\alpha \in [0, \pi/2]$, alors $\alpha = \frac{\omega t}{2}$.

5. Initialement la barre est suivant l'axe Oz et dans l'état final suivant l'axe Oy .

6. Coordonnées de J :

$$\begin{cases} X &= \frac{1}{2}OB \cos \theta = \frac{1}{2}r \cos \theta = b \sin \frac{\omega t}{2} \cos \frac{\omega t}{2} = \frac{1}{2}b \sin \omega t \\ Y &= \frac{1}{2}OB \sin \theta = b \sin^2 \frac{\omega t}{2} \\ Z &= \frac{1}{2}2b \cos \alpha = b \cos \frac{\omega t}{2}. \end{cases}$$

7. On dérive les coordonnées de J soit :

$$\vec{v} \begin{cases} \dot{X} &= \frac{\omega}{2}b \cos \omega t \\ \dot{Y} &= \omega b \sin \frac{\omega t}{2} \cos \frac{\omega t}{2} = \frac{\omega}{2}b \sin \omega t \\ \dot{Z} &= -\frac{\omega}{2}b \sin \frac{\omega t}{2}, \end{cases} \quad \text{puis} \quad \vec{a} \begin{cases} \ddot{X} &= -\frac{\omega^2}{2}b \sin \omega t \\ \ddot{Y} &= \frac{\omega^2}{2}b \cos \omega t \\ \ddot{Z} &= -\frac{\omega^2}{4}b \cos \frac{\omega t}{2}. \end{cases}$$

On en déduit les normes :

$$\|\vec{v}\| = b \frac{\omega}{2} \sqrt{1 + \sin^2 \frac{\omega t}{2}} \quad \text{et} \quad \|\vec{a}\| = b \frac{\omega^2}{2} \sqrt{1 + \frac{1}{4} \cos^2 \frac{\omega t}{2}}.$$

CHAPITRE 13 – CINÉMATIQUE DU POINT

Cinématique du solide

14

Lorsque l'on observe le mouvement d'un solide, deux cas retiennent particulièrement l'attention : les mouvements au cours desquels le solide est en rotation autour d'un axe fixe et les mouvements de translation. Un ascenseur est en translation alors qu'un tambour de machine à laver est en rotation autour d'un axe fixe. Dans ce chapitre, on définit, on décrit et on paramètre ces deux types de mouvements. Une grande partie de ce chapitre doit être mise en correspondance avec le cours de SII du premier semestre.

1 Repérage d'un solide

Jusqu'à présent, on s'est intéressé aux systèmes dont on pouvait négliger l'extension spatiale et que l'on assimilait à des points. On s'occupe maintenant des systèmes pour lesquels ce n'est plus possible. La description du mouvement de tels systèmes requiert plus de paramètres.

1.1 Définition d'un solide

Un **solide** est système matériel dont les points restent à distance constante les uns des autres.

On oppose les solides aux systèmes déformables dont les points peuvent se déplacer les uns par rapport aux autres. Les déformations ou les ruptures du solide sont exclues de cette étude.

Exemple

Un ressort peut être étiré ou comprimé. C'est un système déformable. Une boule de billard est un solide indéformable.

1.2 Repérage d'un solide dans l'espace

On considère un référentiel auquel on associe un repère d'espace orthonormé direct $\mathcal{R} = (O, x, y, z)$. À chaque solide, on peut fixer un repère d'espace $\mathcal{R}_1 = (O_1, x_1, y_1, z_1)$ également orthonormé direct. Les coordonnées d'un point quelconque du solide sont alors constantes dans ce repère $\mathcal{R}_1$. Repérer le solide dans l'espace revient donc à situer le repère lié au solide $\mathcal{R}_1$ par rapport au repère $\mathcal{R}$ lié au référentiel d'étude (voir figure 14.1).

CHAPITRE 14 – CINÉMATIQUE DU SOLIDE

Pour cela, dans le cas général, on a besoin de six paramètres :

- les trois coordonnées d'un point du solide : par exemple celles de l'origine O_1 du repère attaché au solide ;
- trois angles qui définissent l'orientation des axes du repère lié au solide par rapport au référentiel d'étude.

Le repérage du point O_1 a été étudié dans le chapitre de cinématique du point et requiert l'utilisation de trois paramètres. La figure 14.2 permet de se convaincre de la nécessité de définir trois angles pour repérer l'orientation du cube et donc l'orientation de $\mathcal{R}_1$ par rapport à $\mathcal{R}$.

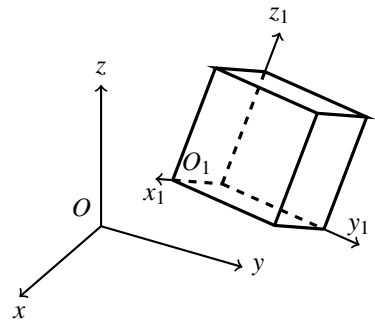


Figure 14.1 – Repérage d'un cube dans l'espace.

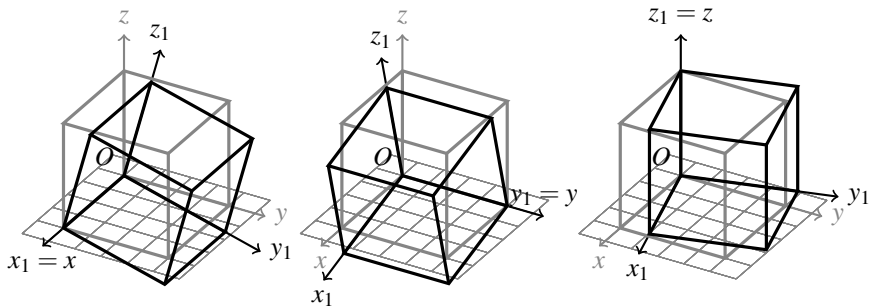


Figure 14.2 – Cube de sommet O fixe repéré dans l'espace. À gauche, le cube a tourné autour l'axe (Ox) ; au centre, autour de (Oy) et à droite, autour de (Oz) .

2 Mouvement de translation

2.1 Définition

Un solide est en **translation** lorsque les directions du repère lié au solide sont fixes par rapport au référentiel d'étude.

2.2 Mouvement d'un point d'un solide en translation

Les directions du repère lié au solide étant fixes dans le référentiel d'étude, on peut faire coïncider les vecteurs de base de $\mathcal{R}_1 : (\vec{u}_{x_1}, \vec{u}_{y_1}, \vec{u}_{z_1})$ avec ceux de $\mathcal{R} : (\vec{u}_x, \vec{u}_y, \vec{u}_z)$. La position d'un point P fixé au solide est repérée par rapport au solide par ses coordonnées (x_P, y_P, z_P) telles que :

$$\overrightarrow{O_1P} = x_P \vec{u}_{x_1} + y_P \vec{u}_{y_1} + z_P \vec{u}_{z_1}.$$

Le point P étant fixé au solide, ses coordonnées (x_P, y_P, z_P) sont constantes. La base $(\vec{u}_{x_1}, \vec{u}_{y_1}, \vec{u}_{z_1})$ étant confondue avec la base $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$, elle est constante dans $\mathcal{R}$. Le vecteur $\vec{O_1P}$ est donc lui aussi constant dans $\mathcal{R}$. Un vecteur quelconque lié au solide est invariant dans $\mathcal{R}$. Le solide ne tourne pas sur lui-même. La position de P dans $\mathcal{R}$ se déduit de celle de O_1 par une translation de vecteur $\vec{O_1P}$ constant. Le mouvement suivi par un point quelconque d'un solide en translation est alors identique à celui de O_1 .

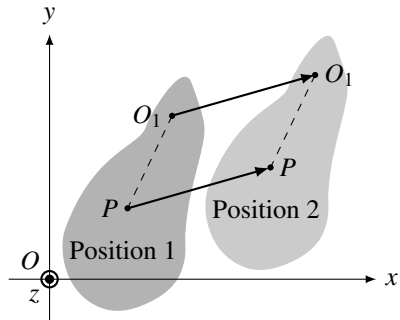


Figure 14.3 – Solide en translation dans $\mathcal{R}$.

2.3 Conséquences

Les vecteurs déplacement infinitésimal, vitesse et accélération de chacun des points du solide sont identiques à chaque instant donc :

- cela a un sens de parler de déplacement, vitesse et accélération du solide ;
- le mouvement du solide est parfaitement défini par la donnée du mouvement d'un de ses points. On choisit souvent O_1 ou le centre d'inertie G du solide.

Tous les points d'un solide en translation ont le même mouvement. Le mouvement du solide est complètement décrit par le mouvement d'un de ses points.

Il suffit de trois coordonnées pour décrire la position d'un solide en translation et on est ramené au cas de la cinématique du point.

Remarque

En lien avec le cours de cinématique du solide de SII, un mouvement de translation d'un solide est caractérisé par la nullité de son vecteur instantané de rotation. Le torseur cinématique est de la forme couple.

2.4 Deux mouvements de translations remarquables

a) Translation rectiligne

Lorsque le mouvement de O_1 est un mouvement rectiligne, tous les points du solide en translation ont un mouvement rectiligne et la translation du solide est qualifiée de **translation rectiligne**. On peut noter que chaque point du solide décrit une portion de droite de même direction mais que l'origine de chacune de ces portions de droite est différente.

Exemple

Un ascenseur est en translation rectiligne le long d'une ligne verticale par rapport au référentiel lié au sol.

CHAPITRE 14 – CINÉMATIQUE DU SOLIDE

b) Translation circulaire

Lorsque le mouvement de O_1 est un mouvement circulaire, tous les points du solide en translation ont un mouvement circulaire et la translation du solide est qualifiée de **translation circulaire**. On peut noter que chaque point du solide décrit un arc de cercle de même rayon mais que le centre de chacun de ces arcs de cercles est différent.

Exemple

Les nacelles d'un manège de type « grande roue » sont en translation circulaire par rapport au référentiel lié au sol.

3 Solides en rotation autour d'un axe fixe

3.1 Définition

Un solide est en rotation autour d'un axe Δ fixe dans $\mathcal{R}$, s'il existe une unique droite Δ immobile à la fois par rapport au solide et au référentiel $\mathcal{R}$.

3.2 Mouvement d'un point d'un solide en rotation

On choisit les origines $O = O_1$ sur l'axe Δ . On fait coïncider l'axe (Oz) du repère d'étude $\mathcal{R}$ et l'axe (O_1z_1) du repère $\mathcal{R}_1$ lié au solide à l'axe fixe Δ . On oriente l'axe de rotation Δ par le vecteur unitaire $\vec{u}_z$. La position du solide est alors entièrement repérée par l'angle θ que fait l'axe (Ox_1) avec l'axe (Ox) (voir figure 14.4).

Un point P lié au solide est repéré par le vecteur $\vec{OP}$ que l'on décompose en une composante parallèle à Δ et une composante qui lui est perpendiculaire. Pour cela, on introduit le point H , projeté orthogonal de P sur Δ et on écrit :

$$\vec{OP} = \vec{OH} + \vec{HP}.$$

Le vecteur $\vec{OH}$ est constant puisqu'il appartient à Δ . La distance HP reste constante puisque H et P appartiennent au solide. Le vecteur $\vec{HP}$ fait un angle constant avec le vecteur $\vec{u}_{x_1}$ car tous deux appartiennent au solide. Le vecteur $\vec{HP}$ est donc entraîné en rotation avec le solide et le point P décrit un cercle de centre H , de rayon HP à la vitesse angulaire $\dot{\theta}$.

Dans $\mathcal{R}$, le mouvement de P est un mouvement circulaire de centre H dans le plan (Hxy) passant par H , orthogonal à $\vec{u}_z$ et orienté par $\vec{u}_z$.

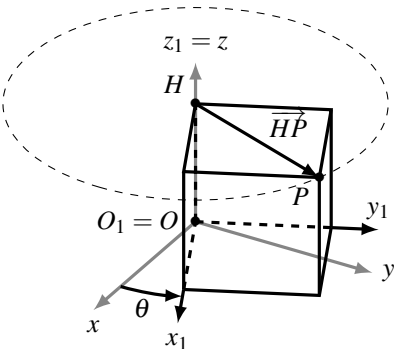


Figure 14.4 – Solide en rotation autour de l'axe (O_1z) fixe dans $\mathcal{R}$.

On peut donc définir sa position par ses coordonnées polaires de centre H . On pose $\|\overrightarrow{HP}\| = r$, $\overrightarrow{u}_r = \frac{\overrightarrow{HP}}{r}$ et $\overrightarrow{u}_\theta = \overrightarrow{u}_z \wedge \overrightarrow{u}_r$ le vecteur du plan (Hxy) directement orthogonal à $\overrightarrow{u}_r$. On obtient alors :

$$\overrightarrow{HP} = r\overrightarrow{u}_r \quad \text{et} \quad \overrightarrow{v}_{P/\mathcal{R}} = r\dot{\theta}\overrightarrow{u}_\theta.$$

La vitesse angulaire $\dot{\theta}$ est commune à tous les points du solide et mesure la vitesse de rotation du solide autour de son axe de rotation. Elle s'exprime en $\text{rad}\cdot\text{s}^{-1}$.

Le mouvement d'un point P lié à un solide en rotation autour d'un axe fixe Δ est un mouvement circulaire situé dans le plan perpendiculaire à Δ contenant P . La trajectoire est caractérisée par

- son centre H , projeté orthogonal P sur l'axe de rotation Δ ,
- son rayon r , distance de P à l'axe de rotation Δ ,
- sa vitesse angulaire $\dot{\theta}$ égale à la vitesse de rotation du solide autour de Δ .

Dans le système de coordonnées cylindriques d'origine O et d'axe (Oz) , la vitesse instantanée du point P s'écrit :

$$\overrightarrow{v}_{P/\mathcal{R}} = r\dot{\theta}\overrightarrow{u}_\theta.$$

Il suffit d'une coordonnée angulaire θ pour décrire la position du solide et on est ramené au cas du mouvement circulaire de la cinématique du point.

Remarque

- L'axe de rotation n'est pas forcément situé à l'intérieur du solide.
- La vitesse de rotation est parfois donnée en $\text{tour}\cdot\text{min}^{-1}$, qui n'est pas une unité du Système International.

3.3 Conséquences

Le déplacement angulaire infinitésimal $d\theta$, la vitesse angulaire $\dot{\theta}$, ainsi que l'accélération angulaire $\ddot{\theta}$ sont identiques pour chacun des points du solide donc :

- cela a un sens de parler de rotation et de vitesse de rotation du solide ;
- le mouvement du solide est parfaitement défini par la donnée du déplacement angulaire d'un de ses points.

Remarque

En lien avec le cours de cinématique du solide de SII, le torseur cinématique pour ce type de mouvement est de la forme glisseur. En prenant $\overrightarrow{u}_z$ orienté selon l'axe de rotation, le vecteur $\overrightarrow{\Omega} = \dot{\theta}\overrightarrow{u}_z$ est le vecteur instantané de rotation du solide par rapport à $\mathcal{R}$. La vitesse de P s'exprime alors $\overrightarrow{v}_{P/\mathcal{R}} = \overrightarrow{\Omega} \wedge \overrightarrow{OP} = \dot{\theta}\overrightarrow{u}_z \wedge r\overrightarrow{u}_r = r\dot{\theta}\overrightarrow{u}_\theta$ où le vecteur $\overrightarrow{u}_\theta = \overrightarrow{u}_z \wedge \overrightarrow{u}_r$ est tel que $(\overrightarrow{u}_r, \overrightarrow{u}_\theta, \overrightarrow{u}_z)$ est une base orthonormée directe.

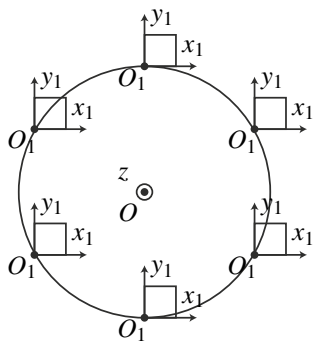
3.4 Quelques exemples de rotation autour d'un axe fixe

Il existe de nombreux exemples de systèmes en rotation autour d'un axe fixe par rapport à un référentiel $\mathcal{R}$. On peut citer :

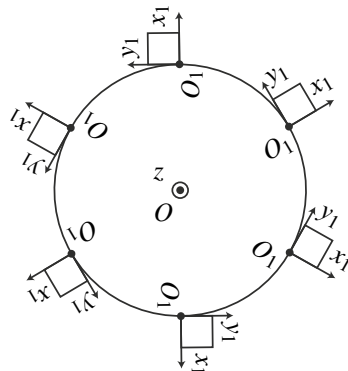
- la Terre dans son mouvement de rotation propre autour de son axe Nord-Sud en un jour sidéral de durée $T = 23\text{ h } 56\text{ min } 4\text{ s} = 86\,164\text{ s}$ par rapport au référentiel géocentrique. La vitesse de rotation est alors constante et égale à $\dot{\theta} = \frac{2\pi}{T} = 7,3 \cdot 10^{-5}\text{ rad}\cdot\text{s}^{-1}$;
- le rotor d'un moteur électrique par rapport au référentiel lié au stator. La vitesse de rotation peut alors varier fortement ;
- une voiture qui prend un virage circulaire par rapport au référentiel lié au sol. Dans ce cas, l'axe de rotation n'est pas dans la voiture mais au centre du virage ;
- les roues d'une voiture se déplaçant en ligne droite par rapport au référentiel lié au châssis de la voiture. L'axe de rotation est alors l'axe du moyeu de la roue. Il est en mouvement par rapport au sol mais fixe par rapport au châssis de la voiture.



Une erreur courante consiste à confondre le mouvement de rotation autour d'un axe fixe et le mouvement de translation circulaire. La figure 14.5 illustre la différence entre ces deux mouvements : tous les points d'un solide en translation circulaire décrivent une trajectoire circulaire de même rayon mais de centres différents et les axes (O_1x_1) et (O_1y_1) restent fixes ; tous les points d'un solide en rotation autour d'un axe fixe décrivent une trajectoire circulaire de même centre mais de rayons différents et les axes (O_1x_1) et (O_1y_1) sont en rotation.



Translation circulaire.



Rotation autour de (Oz) .

Figure 14.5 – Différence entre le mouvement de translation circulaire et le mouvement de rotation pour un solide carré.

SYNTHÈSE

SAVOIRS

- définir un solide
- mouvement de translation
- mouvement de rotation autour d'un axe fixe
- vitesse angulaire d'un solide en rotation

SAVOIR-FAIRE

- différencier un solide d'un système déformable
- reconnaître et décrire une translation rectiligne, une translation circulaire
- décrire la trajectoire d'un point quelconque d'un solide en rotation autour d'un axe fixe.
Exprimer sa vitesse en fonction de sa distance à l'axe et de la vitesse angulaire

MOTS-CLÉS

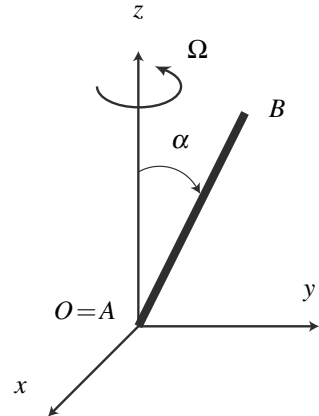
- | | |
|----------------|---------------|
| • solide | • translation |
| • indéformable | • rotation |

S'ENTRAÎNER

14.1 Rotation d'une barre autour d'un axe fixe (★)

Une barre rigide AB de longueur L est mise en rotation uniforme à la vitesse angulaire Ω autour d'un axe fixe (Oz) dans un référentiel $\mathcal{R}$ lié au repère (O, x, y, z) . La barre est située dans un plan vertical et son extrémité A est confondu avec O . Elle fait un angle α avec l'axe (Oz).

1. Décrire le mouvement du point A dans $\mathcal{R}$.
2. Décrire celui de B puis exprimer ses vecteurs position, vitesse et accélération. On s'aidera d'un schéma sur lequel on définira une base adaptée au problème.



14.2 Mouvement d'une nacelle de grande roue (★)

La grande roue installée place de la Concorde à Paris est haute de $H = 60$ m. Chaque nacelle, suspendue au bout d'un bras d'une longueur $d = 2,5$ m, réalise un tour en $\Delta t = 10$ min. On note O , le centre de la grande roue.

1. Décrire le mouvement d'une nacelle. On donnera en particulier les vitesses (linéaire et angulaire) moyennes.



14.3 Face cachée de la Lune (★★)

Dans le référentiel géocentrique, la Lune effectue une révolution circulaire centrée sur la Terre en 27,3 jours. La distance du centre de la Terre au centre de la Lune est environ égale à $D_{TL} = 3,84 \cdot 10^5$ km. Au cours de cette révolution, la Lune montre toujours la même face à la Terre.

1. Décrire le mouvement de la Lune dans le référentiel géocentrique. On s'attachera particulièrement à distinguer s'il s'agit-il d'un mouvement de translation circulaire ou de rotation.
2. En déduire la vitesse angulaire $\dot{\theta}_0$ de la révolution du centre de la Lune sur sa trajectoire.
3. Déterminer la vitesse et l'accélération du centre de la Lune dans le référentiel géocentrique. Calculer numériquement la norme de sa vitesse.
4. Décrire le mouvement de la Lune dans le repère sélénocentrique qui a les mêmes axes de références que le référentiel géocentrique mais pour origine le centre L de la Lune.
5. Déterminer la vitesse angulaire $\dot{\theta}_p$ de la rotation de la Lune sur elle-même.

APPROFONDIR

14.4 Swing de golf (★★)

On considère le mouvement d'un club de golf lors d'un swing représenté sur la chronophotographie ci-jointe avec une prise de vue toutes les 20 ms.

1. Peut-on décrire simplement le mouvement du club de golf ?
2. On se restreint aux moments qui entourent l'impact de la tête du club avec la balle, c'est-à-dire aux 5 positions du club numérotées sur la photographie. Peut-on alors décrire simplement le mouvement du club de golf ?
3. Déterminer les vitesses angulaires moyennes entre chaque prise de vue. Le mouvement du club est-il uniforme ?
4. L'impact avec la balle a approximativement lieu au point 3. Peut-on donner une explication au ralentissement de la vitesse du club observé par la suite ?

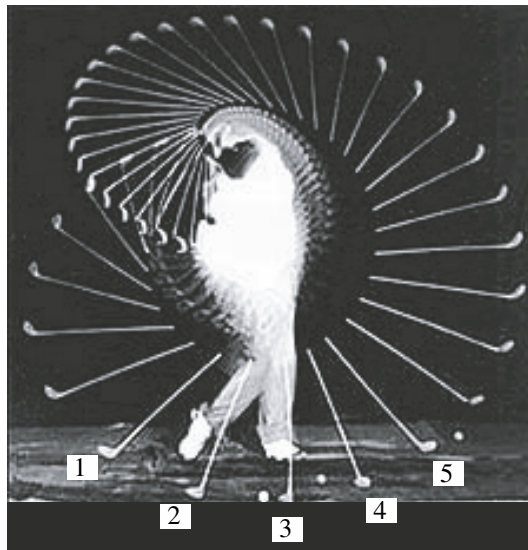


Figure 14.6 – Swing de golf. Source : <http://www.clubmaker-online.com/bj003.gif>.

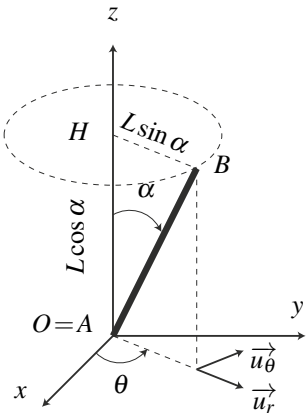
CORRIGÉS

14.1 Rotation d'une barre autour d'un axe fixe

1. Le point A appartient à l'axe de rotation (Oz) fixe dans $\mathcal{R}$. Il est donc fixe.
2. La barre est un solide en rotation autour de l'axe fixe (Oz) à la vitesse angulaire Ω . B décrit donc une trajectoire circulaire et uniforme à la vitesse angulaire Ω . Le centre de la trajectoire de B est son projeté orthogonal H sur l'axe de rotation (Oz) .

La trajectoire de B est circulaire de centre H de coordonnées $(0, 0, z = L \cos \alpha)$ et de rayon $HB = r = L \sin \alpha$ dans le plan perpendiculaire à (Oz) passant par H . On définit la base cylindrique qui suit la barre dans son mouvement de rotation et on applique les résultats du cours de *Cinématique du point* en coordonnées cylindriques pour trouver :

$$\begin{cases} \overrightarrow{OB} &= r\vec{u}_r + z\vec{u}_z = L \sin \alpha \vec{u}_r + L \cos \alpha \vec{u}_z \\ \vec{v} &= L \sin(\alpha) \Omega \vec{u}_\theta \\ \vec{a} &= -L \sin(\alpha) \Omega^2 \vec{u}_r. \end{cases}$$



14.2 Mouvement d'une nacelle de grande roue

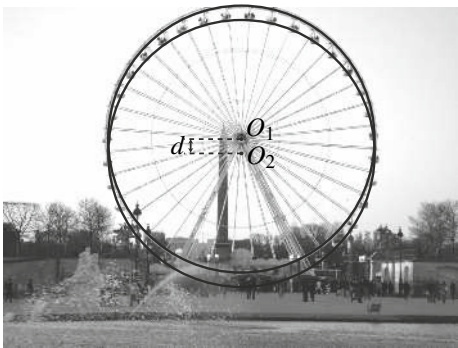
1. Une nacelle est en translation circulaire de rayon $R = \frac{H}{2}$ et parcourt le cercle en une durée $\Delta t = 10$ min.

Sa vitesse angulaire moyenne est donc $\langle \dot{\theta} \rangle = \frac{2\pi}{\Delta t} = \frac{2 \times 3,1415}{10 \times 60} = 0,010 \text{ rad} \cdot \text{s}^{-1}$.

Sa vitesse linéaire moyenne est :

$$\langle R\dot{\theta} \rangle = \frac{2\pi R}{\Delta t} = 0,31 \text{ m} \cdot \text{s}^{-1}$$

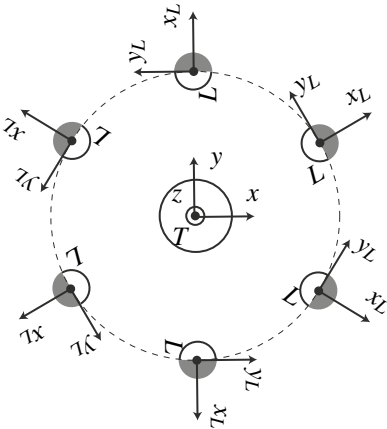
Tous les points de la nacelle décrivent un cercle de même rayon avec la même vitesse moyenne. Par contre, comme on le voit sur la figure ci-contre, la trajectoire n'est pas centrée sur le même point. Le bas de la nacelle décrit un cercle centré sur O_2 alors que le point d'attache de la nacelle sur la roue décrit un cercle centré sur O_1 situé à $d = 2,5$ m au dessus.



14.3 Face cachée de la Lune

1. La Lune montrant toujours la même face à la Terre, on peut représenter son mouvement dans le référentiel géocentrique comme sur le schéma ci-contre où la face cachée est grisée. Dans le référentiel géocentrique, la Lune a un mouvement de rotation uniforme dans le sens direct autour de l'axe $(T, \vec{u}_z)$.

2. Dans ce référentiel, la Lune retrouve sa position initiale après une durée $\Delta T = 27,3$ jours = $2,359.10^6$ s. La vitesse angulaire de la rotation de la Lune autour de la Terre est donc $\dot{\theta}_0 = \frac{2\pi}{\Delta T} = \frac{2 \times 3,1415}{2,359.10^6} = 2,66.10^{-6} \text{ rad}\cdot\text{s}^{-1}$.



3. Dans ce référentiel, tous les points fixes de la Lune décrivent une trajectoire circulaire et uniforme à la vitesse angulaire $\dot{\theta}_0$. Le centre de la Lune ne déroge pas à cette règle. Le centre de sa trajectoire est T et son rayon D_{TL} . On peut alors appliquer les résultats de la cinématique du point avec $(\vec{u}_{x_L}, \vec{u}_{y_L})$ la base mobile des coordonnées polaires de centre T :

$$\vec{TL} = D_{TL}\vec{u}_{x_L} \quad ; \quad \vec{v}(L) = D_{TL}\dot{\theta}_0\vec{u}_{y_L} \quad ; \quad \vec{a}(L) = -D_{TL}\dot{\theta}_0^2\vec{u}_{x_L}.$$

Numériquement, $\|\vec{v}(L)\| = 3,84.10^8 \times 2,66.10^{-6} = 1,02.10^3 \text{ m}\cdot\text{s}^{-1}$

4. Dans le repère sélénocentrique, la Lune a un mouvement de rotation dans le sens direct autour de l'axe $(L, \vec{u}_z)$. Elle fait un tour sur elle-même en $\Delta T = 27,3$ jours.

5. La vitesse angulaire $\dot{\theta}_p$ de la rotation de la Lune sur elle-même est donc la même que la vitesse angulaire $\dot{\theta}_0$ de la révolution de la Lune autour de la Terre. On parle de rotation synchrone.

14.4 Swing de golf

Le club de golf est un solide indéformable tenu par un joueur dont le corps peut se déformer (au sens de la mécanique).

1. Le mouvement du club de golf est un mouvement compliqué qui ne peut pas être décrit simplement dans le cadre de la mécanique du solide. La tête suit une courbe à trois dimensions à laquelle on ne peut pas accéder à l'aide de photographies.

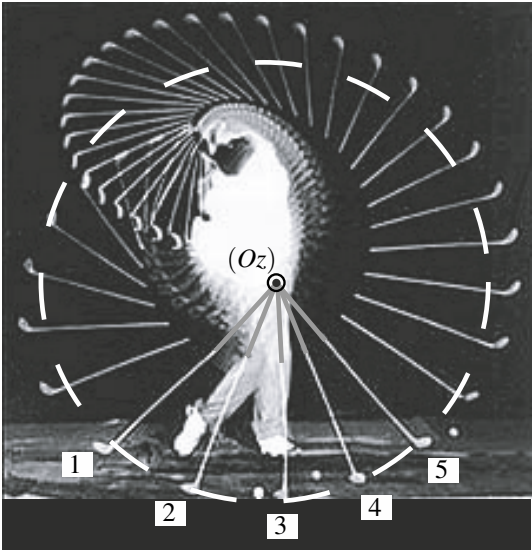
2. Autour de l'impact, le mouvement du club de golf peut raisonnablement être assimilé à un mouvement de rotation autour de l'axe (Oz) représenté sur la chronophotographie ci-jointe.

3. On détermine alors la vitesse moyenne entre deux prises de vue en mesurant l'angle de rotation $\Delta\theta$ du club et en le divisant par la durée entre deux clichés $\Delta t = 20$ ms. On regroupe les résultats dans le tableau suivant :

CHAPITRE 14 – CINÉMATIQUE DU SOLIDE

	rotation 1-2	rotation 2-3	rotation 3-4	rotation 4-5
$\Delta\theta$	24°	24°	21°	18,5°
$\langle \dot{\theta} \rangle$	21 rad·s ⁻¹	21 rad·s ⁻¹	18 rad·s ⁻¹	16 rad·s ⁻¹

4. L'impact avec la balle a approximativement lieu au point 3. Le choc du club sur la balle projette la balle vers l'avant. En réaction, le club de golf est freiné, ce qui explique le ralentissement après l'impact visible directement par le fait que l'écart angulaire entre la position du club est manifestement maximal juste avant l'impact. Les mesures effectuées confirment cette impression visuelle.



Principes de la dynamique newtonienne

15

Dans le chapitre de cinématique du point, on ne s'est intéressé qu'à l'analyse du mouvement et à la description des liens entre les divers vecteurs décrivant ce dernier (position, vitesse et accélération). On n'a jamais cherché à en déterminer les causes. C'est ce point qui va être envisagé ici. Cela conduit à introduire la notion de force ainsi qu'à énoncer les trois lois de Newton. Avant cela, il est nécessaire de définir les éléments cinétiques du système.

1 Éléments cinétiques d'un point matériel

1.1 Masse

D'un point de vue cinématique, un point matériel est décrit à l'aide des vecteurs position, vitesse et accélération. Cependant quelques exemples de la vie courante mettent en évidence que le comportement d'un corps ne dépend pas uniquement de ces paramètres cinématiques. Ainsi un joueur de ping-pong sera dans l'impossibilité de renvoyer la balle si celle-ci est remplacée par une boule de billard. De même, si un enfant joue avec des boules de pétanque, il ne pourra pas lancer celles-ci très loin. Il est donc nécessaire d'introduire une grandeur physique mesurant la capacité du corps à résister au mouvement qu'on souhaite lui imposer. Cette propriété s'appelle l'**inertie** du système.

La grandeur introduite est fondamentale au niveau dynamique et s'appelle **masse inerte** ou **masse inertielle** du corps. Il s'agit d'un scalaire positif qui est d'autant plus grand que le corps s'oppose au mouvement. Son unité légale est le kilogramme, de symbole kg. On constate expérimentalement que cette grandeur est proportionnelle à la quantité de matière composant le corps. C'est une grandeur additive, c'est-à-dire que la masse de l'ensemble formé par deux corps est égale à la somme des masses de chacun d'entre eux.

La mesure de cette grandeur s'obtient en pratique à l'aide d'une balance en utilisant la mesure du poids du corps : $\vec{P} = m\vec{g}$ où m est alors la **masse pesante**. Plusieurs problèmes apparaissent avec cette méthode. Tout d'abord, le champ de pesanteur $\vec{g}$ varie à la surface de la Terre. Il est donc *a priori* nécessaire d'effectuer des réglages à chaque déplacement de la

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

balance grâce à une référence. Le deuxième point est l'hypothèse selon laquelle masse inertielle et masse pesante (celle qui apparaît dans l'expression du poids) ne sont qu'une seule et même grandeur. Cette égalité, posée en principe appelé « principe d'équivalence » est vérifiée expérimentalement avec une précision relative de 10^{-12} . Par conséquent, on ne fera pas de différence entre masse inerte et masse pesante.

Remarque

La référence universelle de masse est un étalon cylindrique de platine iridié conservé au Pavillon de Breteuil au sein du *Bureau International des Poids et Mesures*. À l'heure actuelle, la masse est la dernière grandeur dont l'unité est définie par un étalon.

Les premières expériences pour établir l'équivalence entre masse inertielle m_i et masse gravitationnelle m_g sont dues à Galilée. Cette équivalence a été confirmée par différentes expériences. En 1890, le hongrois L. von Eötvös a obtenu une précision relative de 10^{-6} à l'aide d'une adaptation de la balance de H. Cavendish. Dans les années 1960, R. H. Dicke et V. B. Brazinsky ont obtenu une précision relative de 10^{-12} avec une méthode basée sur l'étude de l'accélération du Soleil.

La grandeur qui mesure la capacité d'un corps à résister à la mise en mouvement est sa masse mesurée en kg. La masse est un scalaire d'autant plus grand que le corps est inerte. C'est une grandeur extensive (additive) et intrinsèque (liée uniquement au corps considéré).

1.2 Quantité de mouvement

a) Quantité de mouvement d'un point matériel

La **quantité de mouvement** est une grandeur introduite par Isaac Newton pour formuler les lois de la mécanique portant son nom. Elle est définie par le vecteur :

$$\vec{p} = m \vec{v},$$

pour un point M de masse m et de vecteur vitesse $\vec{v}$ par rapport à un référentiel $\mathcal{R}$. Contrairement à la masse, cette quantité dépend du référentiel dans lequel on travaille puisqu'elle est fonction de la vitesse dans ce référentiel.

b) Quantité de mouvement d'un système de points matériel

Lorsque le système est constitué de plusieurs points matériels, on peut définir sa quantité de mouvement comme la somme des quantités de mouvement de chacun des points qui le constituent. Pour fixer les idées, on se place dans le cas où le système est constitué de deux points matériels M_1 et M_2 de masses respectives m_1 et m_2 et de vitesses respectives $\vec{v}_1$ et $\vec{v}_2$ dans le référentiel $\mathcal{R}$. La quantité de mouvement $\vec{p}$ du système est définie comme la somme des quantités de mouvement de chacun des deux points :

$$\vec{p} = m_1 \vec{v}_1 + m_2 \vec{v}_2.$$

La position du centre de gravité G du système de point est définie par la relation barycentrique

$$(m_1 + m_2) \vec{OG} = m_1 \vec{OM}_1 + m_2 \vec{OM}_2,$$

d'où l'on déduit en dérivant terme à terme :

$$(m_1 + m_2) \frac{d\vec{OG}}{dt} = m_1 \frac{d\vec{OM}_1}{dt} + m_2 \frac{d\vec{OM}_2}{dt} \iff (m_1 + m_2) \vec{v}_G = m_1 \vec{v}_1 + m_2 \vec{v}_2.$$

En posant $m = m_1 + m_2$ la masse du système, la quantité de mouvement du système est ainsi égale à la quantité de mouvement de son centre de gravité G affecté de la masse totale m et animé de la vitesse $\vec{v}_G$:

$$\vec{p} = m\vec{v}_G.$$

Il s'agit d'un résultat général, qui reste vrai pour un système constitué de plus de deux points matériels et par exemple pour un solide de masse m . On peut alors revenir sur notre définition initiale du point matériel en mécanique : on peut assimiler un solide à un point matériel s'il suffit d'étudier le mouvement de son centre de gravité pour comprendre son mouvement.

2 Les trois lois de Newton

Ces trois lois constituent les fondements de la mécanique classique.

2.1 Première loi de Newton : Principe d'inertie

a) Point matériel isolé

Un point matériel est **isolé** s'il n'est soumis à aucune interaction avec l'extérieur.

Il s'agit d'un cas limite utilisé en mécanique. En pratique, un point matériel est considéré comme isolé lorsque l'on peut négliger les forces auxquelles il est soumis.

Exemple

Les sondes Pioneer 10 et 11 et Voyager 1 et 2 qui se dirigent vers les confins du système solaire sont approximativement des systèmes isolés.

b) Énoncé du principe d'inertie

Il existe une classe de référentiels privilégiés appelés **référentiels galiléens** dans lesquels **tout** point matériel isolé est animé d'un mouvement rectiligne et uniforme.

Le principe d'inertie formule l'existence de référentiels particuliers, les référentiels galiléens, dont il fournit une définition à partir du mouvement des points matériels isolés. On constate la différence essentielle apportée par la dynamique vis-à-vis de la cinématique : les référentiels ne jouent plus tous le même rôle.

Cependant, si l'on considère deux référentiels $\mathcal{R}_1$ et $\mathcal{R}_2$ en translation rectiligne et uniforme l'un par rapport à l'autre, tout point matériel animé d'un mouvement rectiligne et uniforme par rapport à $\mathcal{R}_1$ sera également en translation rectiligne et uniforme par rapport à $\mathcal{R}_2$. Lorsque l'un est galiléen alors l'autre l'est également.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

Les **référentiels galiléens** sont en translation rectiligne et uniforme les uns par rapport aux autres.

c) Détermination pratique d'un référentiel galiléen

Tout comme la notion de point isolé, à partir de laquelle elle est définie, la notion de référentiel galiléen est un cas limite. En pratique, on utilise principalement trois référentiels en fonction de la durée du phénomène étudié :

- le **référentiel terrestre** ou **référentiel du laboratoire** est lié à la Terre. Son origine est située au point de la surface du globe où se déroule l'expérience et ses axes sont fixes par rapport à la Terre. Il est considéré comme galiléen lorsque l'on peut négliger la rotation de la Terre autour de l'axe des ses pôles. Il est adapté à l'étude des mouvements se déroulant sur Terre et dont la durée est faible devant la durée d'un jour ;
- le **référentiel géocentrique** a son origine au centre de la Terre et des axes pointant vers des étoiles lointaines fixes. Dans ce référentiel, la Terre tourne sur elle-même autour de l'axe de ses pôles. Il est considéré comme galiléen lorsque l'on peut négliger le mouvement orbital de la Terre dont la durée caractéristique est un an. Il est adapté à l'étude du mouvement des satellites autour de la Terre ;
- le **référentiel héliocentrique** a son origine au centre du Soleil et des axes pointant vers des étoiles lointaines fixes. Il est considéré comme galiléen tant que l'on peut négliger le mouvement orbital du Soleil autour du centre de notre galaxie, la Voie lactée, dont la durée caractéristique est estimée à environ 230 millions d'années. Il est adapté à l'étude du mouvement des planètes autour du Soleil.

2.2 Deuxième loi de Newton : Principe fondamental de la dynamique

On étudie le mouvement d'un point matériel M de masse m dans un référentiel galiléen $\mathcal{R}$. À l'instant t , on note $\overrightarrow{OM}$, $\vec{v}$, $\vec{p} = m\vec{v}$ et $\vec{a}$ les vecteurs position, vitesse, quantité de mouvement et accélération de M dans $\mathcal{R}$. On s'intéresse aux causes des mouvements et/ou de leur modification, c'est-à-dire aux interactions mécaniques entre le système et le milieu extérieur.

a) Notion de force

Un système peut être mis en mouvement ou, s'il est déjà en mouvement, ce dernier peut être modifié. Les causes de cette « modification » doivent être recherchées dans ses interactions avec l'extérieur. Cela conduit à définir la notion de forces.

Une **force** est une grandeur vectorielle décrivant l'interaction capable de modifier et/ou de produire un mouvement ou une déformation du système.

La force est décrite par un vecteur ; le caractère vectoriel de la force apparaît dans des expériences simples. Par exemple, lorsque l'on tire sur un ressort, on constate que ce dernier s'allonge le long d'une direction et dans un sens qui sont ceux de l'effort ou de la force qu'on exerce sur lui. Son allongement est d'autant plus grand que la force est importante. Trois paramètres : direction, sens et intensité interviennent pour déterminer l'action exercée sur

le ressort. Un vecteur défini par ces trois mêmes paramètres décrit la force. Il faudra donc donner la direction, le sens et la norme d'une force pour la connaître parfaitement.

On distingue deux grandes catégories de forces : les forces à distance et les forces de contact. On abordera plus tard en détail ces deux types d'interactions.

b) Énoncé du principe fondamental de la dynamique

Ayant admis l'existence d'un référentiel galiléen et de forces caractérisant les interactions du système avec l'extérieur, on se place dans ce référentiel et on modélise les effets des forces extérieures sur le système.

Dans un référentiel galiléen, la dérivée par rapport au temps du vecteur quantité de mouvement du système est égale à la somme des forces extérieures s'exerçant sur le système. Mathématiquement, cela se traduit par la relation :

$$\frac{d\vec{p}}{dt} = \sum_i \vec{f}_i. \tag{15.1}$$

Interprétation Ce principe relie la dérivée d'un terme cinétique, le vecteur quantité de mouvement, à un terme dynamique, la somme des forces extérieures traduisant les interactions subies par le point matériel. Cette relation peut être utilisée dans les deux sens :

- obtenir la description cinématique du mouvement connaissant les forces subies par le point matériel,
- déterminer la somme des forces s'exerçant sur le point matériel à partir de la connaissance du mouvement.

Remarque

Le principe fondamental de la dynamique est également appelé « relation fondamentale de la dynamique » ou « théorème de la quantité de mouvement ».

La somme des forces extérieures est également appelée « résultante des forces ».

c) Première conséquence : cas des systèmes pseudo-isolés

Les systèmes **pseudo-isolés** sont définis comme les systèmes pour lesquels la somme des forces extérieures appliquées est nulle :

$$\sum_i \vec{f}_i = \vec{0} \implies \frac{d\vec{p}}{dt} = \vec{0}.$$

Pour de tels systèmes, le principe fondamental de la dynamique montre que la quantité de mouvement $\vec{p}$, et donc la vitesse $\vec{v}$, sont constantes au cours du temps. Ces systèmes sont animés d'un mouvement rectiligne uniforme dans un référentiel galiléen. Ceci explique leur dénomination : tout se passe comme s'ils étaient isolés.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

d) Cas particulier d'un système à masse constante :

Dans ce cas, il est possible de « sortir » la masse de la dérivée de la quantité de mouvement puisqu'il s'agit d'une constante : $\frac{d\vec{p}}{dt} = \frac{d(m\vec{v})}{dt} = m\frac{d\vec{v}}{dt} = m\vec{a}$.

Le principe fondamental de la dynamique s'écrit alors sous la forme :

$$m\vec{a} = \sum_i \vec{f}_i.$$

Remarque

Il faut se méfier de cette formulation qui ne s'applique qu'aux systèmes de masse constante. Le cas d'une fusée qui consomme du carburant et voit sa masse diminuer ou encore celui d'une goutte d'eau qui tombe dans une atmosphère contenant de la vapeur d'eau et grossit au cours du mouvement ne peuvent pas être traités avec cette relation. Il est donc fortement conseillé de retenir le principe fondamental de la dynamique sous la forme générale de l'équation (15.1).

e) Cas particulier d'un système à l'équilibre :

La quantité de mouvement des systèmes à l'équilibre est nulle à tout instant :

$$\vec{p} = \vec{0} \implies \frac{d\vec{p}}{dt} = \vec{0}.$$

Le principe fondamental de la dynamique s'écrit alors :

$$\sum_i \vec{f}_i = \vec{0},$$

que l'on appelle parfois principe fondamental de la statique.

2.3 Troisième loi de Newton : principe des actions réciproques

a) Énoncé du principe des actions réciproques

Ce principe, qui constitue la troisième loi de Newton, est également appelé « principe de l'action et de la réaction ».

Si le milieu extérieur exerce la force $\overrightarrow{f_{ext \rightarrow M}}$ sur M , alors M exerce la force $\overrightarrow{f_{M \rightarrow ext}}$ sur le milieu extérieur telle que $\overrightarrow{f_{M \rightarrow ext}} = -\overrightarrow{f_{ext \rightarrow M}}$. Les actions de M sur l'extérieur et de l'extérieur sur M sont opposées donc de même norme et de sens différents.

Exemple

Lorsque l'on réalise un coup au tennis, les cordes de la raquette appliquent une force sur la balle pour la mettre en mouvement. Par réaction, la balle applique une force sur les cordes. À la longue, ces forces provoquent l'usure des cordes qui finissent par casser.

b) Cas de l'interaction avec un milieu extérieur réduit à un point

Si le système et le milieu extérieur avec lequel a lieu l'interaction sont tous deux ponctuels et nommés M et M_{ext} , les forces $\overrightarrow{f_{M \rightarrow ext}}$ et $\overrightarrow{f_{ext \rightarrow M}}$ s'exercent de plus sur la même droite d'action qui relie les deux points (MM_{ext}). En effet, dans ce cas, le problème présente une symétrie d'axe (MM_{ext}). Les forces doivent présenter la même symétrie en vertu du principe de Curie¹. Elles sont donc portées par le même axe.

3 Limite de validité de la mécanique classique

3.1 Qu'est-ce qu'un principe ?

Les trois lois de Newton constituent la base de la dynamique newtonienne et de la mécanique classique. Il s'agit de principes au sens où ces lois sont postulées et non démontrées théoriquement. Elles sont considérées comme valables tant que l'expérience ne les a pas contredit. Ce sont les écarts entre certaines observations expérimentales et les prédictions théoriques issues des principes de Newton, qui ont conduit à l'édification des théories de la mécanique quantique et de la mécanique relativiste, en posant des principes différents de ceux de I. Newton. Pour les problèmes envisagés ici, ceux-ci sont bien évidemment valables mais on va détailler les limites de validité de ces principes.

3.2 Les hypothèses de la mécanique classique

Les hypothèses de la mécanique ne sont pas souvent formulées car elles paraissent évidentes. Cependant, la physique moderne limite la portée de la mécanique classique en remettant en cause ces évidences. Il n'est donc pas inutile de les énoncer :

- le temps est universel, il ne dépend pas de l'observateur ;
- l'espace est euclidien ;
- le comportement des systèmes mécaniques est déterministe : deux objets identiques placés dans des conditions identiques (mêmes conditions initiales et forces appliquées identiques) suivent un mouvement identique ;
- l'espace et le temps sont des grandeurs continues.

3.3 Les limites de la mécanique classique

À la fin du XIX^e siècle, certaines observations ne sont pas explicables par la théorie classique. Au début du XX^e siècle, de nouvelles théories voient le jour. Ces nouvelles théories ne remettent pas en cause les résultats de la mécanique classique mais elles en fixent les limites de validité et la complètent.

- La théorie classique est valide tant que le système ne va pas trop vite. Pour quantifier la notion de « trop vite », on compare la vitesse v du système à celle de la lumière c . Si $v \ll c$, la théorie classique est valide. Sinon, il faut tenir compte de la théorie de la relativité restreinte énoncée par A. Einstein en 1905. Le temps perd son caractère absolu, il devient relatif et dépend du référentiel dans lequel on réalise l'étude.
- La théorie classique est valide tant que l'on n'est pas trop proche d'un objet massif. Pour

1. Ce principe stipule que les effets présentent au minimum les mêmes symétries que les causes.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

quantifier la notion de « trop proche », on compare la distance R séparant le système et l'objet massif à une distance caractéristique R_S appelée rayon de Schwartzchild². Si $R \gg R_S$, la théorie classique s'applique. Sinon, il faut tenir compte de la relativité générale énoncée par A. Einstein en 1916. L'espace perd son caractère euclidien, il peut être courbé par une très grande masse (étoile massive, trou noir).

- La théorie classique est valide pour des objets de taille suffisante. Pour quantifier la notion de « taille suffisante », on compare la taille l du système à la longueur d'onde de de Broglie λ définie dans le chapitre *Introduction au monde quantique*. Si $l \gg \lambda$, la théorie classique s'applique. Sinon, il faut tenir compte de la théorie de la mécanique quantique énoncée par N. Bohr, L. de Broglie, P. Dirac, A. Einstein, W.K. Heisenberg, M. Planck, E. Schrodinger et bien d'autres dans la première moitié du XX^e siècle. Le déterminisme classique doit être abandonné au profit d'une approche probabiliste.

4 Premières applications : détermination d'une loi de force

On a vu dans le chapitre de cinématique, qu'on est capable d'observer le mouvement d'un point et de déterminer son accélération. À partir de là, le principe fondamental de la dynamique va nous permettre de déterminer l'expression des forces usuelles.

4.1 Détermination dynamique d'une force : mesure de g

Lors de l'étude cinématique du mouvement, on a observé, dans le référentiel terrestre, le mouvement d'un corps de masse $m = (50 \pm 1)\text{g}$ soumis à la pesanteur terrestre. On a trouvé que l'accélération était verticale, dirigée vers le bas et de norme :

$$a = (10,1 \pm 0,4)\text{m}\cdot\text{s}^{-2}.$$

Les incertitudes sur les valeurs ci-dessus sont à prendre au sens d'incertitudes élargies à un niveau de confiance identique de 95%. On suppose que le référentiel terrestre est galiléen et qu'on peut négliger toutes les forces autres que le poids. Le système est soumis uniquement à son poids $\vec{P}$ et en lui appliquant le principe fondamental de la dynamique, on trouve :

$$m\vec{a} = \vec{P} = m\vec{g} \quad \Longleftrightarrow \quad \vec{a} = \vec{g},$$

qui correspond à un mouvement à accélération constante. On en déduit que l'accélération de la pesanteur $\vec{g}$ est verticale, dirigée vers le bas et de norme $g = a$. On trouve l'incertitude sur le poids $P = mg$ en sommant les incertitudes relatives sur m et sur g :

$$\frac{\Delta P}{P} = \sqrt{\left(\frac{\Delta m}{m}\right)^2 + \left(\frac{\Delta g}{g}\right)^2} = \sqrt{\left(\frac{1}{50}\right)^2 + \left(\frac{0,4}{10,1}\right)^2} = 5\%.$$

On en déduit que la norme du poids qui s'applique sur le système est, avec un niveau de confiance de 95% :

$$P = mg = (0,505 \pm 0,025)\text{N}.$$

2. Le rayon de Schwartzchild d'un corps céleste de masse M est défini par la relation $R_S = 2\mathcal{G}M/c^2$ où $\mathcal{G} = 6,67 \cdot 10^{-11} \text{N}\cdot\text{m}^2\cdot\text{kg}^{-2}$ est la constante de gravitation universelle.

Remarque

Cette démarche permet à I. Newton de déterminer l'expression de la force d'attraction gravitationnelle en déterminant l'accélération des planètes autour du Soleil grâce aux observations de J. Kepler.

4.2 Détermination statique d'une force

Une fois que l'on connaît une force (par exemple le poids), on peut déterminer les autres forces par comparaison. C'est la méthode de mesure statique qui permet, à partir de situations d'équilibre, de déterminer des égalités de forces.

Pour clarifier le problème, on étudie la situation suivante : on attache un ressort à un support et on le laisse évoluer vers un état d'équilibre. Dans un premier temps, l'extrémité libre est laissée à vide et on mesure une longueur à vide l_0 . Dans un deuxième temps, on accroche une masse connue m à son extrémité libre et on mesure une longueur à l'équilibre l_{eq} . Lorsque la masse est à l'équilibre dans le référentiel terrestre, son accélération est nulle. Or elle est soumise à deux forces :

- son poids $\vec{P} = m \vec{g}$;
 - la force de rappel exercée par le ressort $\vec{T}_{eq}$.
- Le principe fondamental de la dynamique donne :

$$\vec{a} = \vec{0} = \vec{P} + \vec{T}_{eq} \quad \Leftrightarrow \quad \vec{T}_{eq} = -m \vec{g}.$$

En faisant varier la masse m suspendue au ressort, on observe alors que plus m est grand, plus le ressort est étiré à l'équilibre. On dispose ainsi d'une méthode qui permet d'étudier le lien entre la force exercée par le ressort $\vec{T}_{eq}$ et son allongement $\Delta l = l_{eq} - l_0$. On peut alors montrer que la force exercée par le ressort est proportionnelle à son allongement (voir paragraphe 5.3).

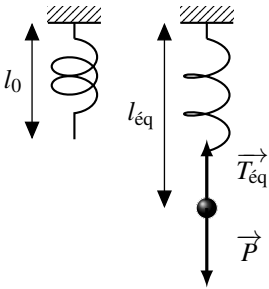


Figure 15.1 –
Allongement d'un
ressort auquel on
suspend une masse.

Remarque

La mesure de force en statique est actuellement utilisée pour mesurer des forces à l'aide d'un dynamomètre. Historiquement, ce type de démarche permet à H. Cavendish de déterminer la valeur de la constante de gravitation $\mathcal{G}$ par comparaison avec la force connue exercée par un pendule de torsion.

5 Classification des forces

En mécanique, au niveau macroscopique, on observe deux grands types de forces : les **forces de contact** et les **forces à distances**. Dans cette partie, on va décrire les forces usuelles que le milieu extérieur peut exercer sur le système étudié. Toutes ces forces sont des modèles traduisant au niveau macroscopique les effets de quatre **interactions fondamentales** que l'on va très brièvement présenter.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

5.1 Les quatre interactions fondamentales

Interaction	Particules concernées	Portée	Principale manifestation
Forte	quarks et dérivés	$\simeq 10^{-15}$ m	cohésion des noyaux
Faible	constituants des noyaux	$\simeq 10^{-18}$ m	radioactivité β
Electromagnétique	particules chargées	infinie	existence des atomes réactions chimiques
Gravitationnelle	particules massiques	infinie	existence et trajectoires des astres, des planètes

Tableau 15.1 – Les interactions fondamentales et leurs principales manifestations.

Au niveau des particules élémentaires de la matière, les interactions sont au nombre de quatre. Le tableau 15.1 les décrit très brièvement. Pour lire ce tableau, on rappelle que les particules dérivées des quarks les plus connues sont les protons et les neutrons (les électrons et les neutrinos font partie de la famille des leptons). Par ailleurs, l’interaction faible est dénommée ainsi car elle est 10^5 fois plus faible que l’interaction forte.

Remarque

Le souhait d’obtenir une théorie unique pour rendre compte de l’ensemble de ces quatre interactions est très ancien. Dans les années 1960, S.H. Glashow, A. Salam et S. Weinberg (prix Nobel en 1979) ont proposé un modèle électrofaible unifiant les interactions faible et électromagnétique. Cette théorie a été vérifiée expérimentalement en 1983. À ce jour, aucune théorie unificatrice n’a été établie malgré de nombreuses recherches. Les interactions forte et faible sont négligeables au niveau macroscopique. Dans la suite, on n’en tiendra pas compte puisqu’on se placera à cette échelle. Il suffit ici de savoir que ces interactions existent et permettent d’expliquer des phénomènes du noyau incompréhensibles à l’aide des interactions électromagnétique et gravitationnelle.

5.2 Forces à distance

Ces forces s’appliquent sans qu’il y ait de contact entre le milieu extérieur et le système.

a) Force gravitationnelle

Cette force s’applique au centre de gravité G du système de masse m situé dans un champ gravitationnel $\vec{g}_{\text{gravi}}$ créé en M par le milieu extérieur. On se restreint au cas le plus simple où le champ gravitationnel est créé par une masse ponctuelle m_p placée en P et où le système est réduit à un point $M = G$.

La force gravitationnelle entre deux masses m_p et m placées en P et M est :

- attractive,
- de norme proportionnelle à m et m_p , ainsi qu'à l'inverse du carré de la distance PM .

Elle s'écrit :

$$\vec{F}_{P \rightarrow M} = -\vec{F}_{M \rightarrow P} = -\mathcal{G} m_p m \frac{\vec{PM}}{PM^3} = -\frac{\mathcal{G} m_p m}{PM^2} \vec{u}_{PM}$$

où $\vec{u}_{PM}$ est le vecteur unitaire dirigé de P vers M et $\mathcal{G} = 6,67 \cdot 10^{-11} \text{ N} \cdot \text{m}^2 \cdot \text{kg}^{-2}$ est la **constante de gravitation universelle**.

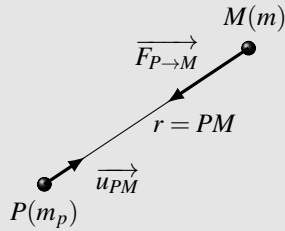


Figure 15.2 – Force gravitationnelle entre deux masses m_p et m placées en P et M .

Remarque

On peut aussi écrire $\vec{F}_{P \rightarrow M} = m \left(-\frac{\mathcal{G} m_p}{PM^2} \vec{u}_{PM} \right) = m \vec{g}_{\text{gravi}}$ en introduisant le champ gravitationnel créé en M par la masse m_p placée en P : $\vec{g}_{\text{gravi}} = -\frac{\mathcal{G} m_p}{PM^2} \vec{u}_{PM}$. Cette notation présente l'avantage de séparer ce qui dépend de la masse m de ce qui dépend du milieu extérieur. De ce fait le champ gravitationnel $\vec{g}_{\text{gravi}}$ existe en M même si la masse m n'est pas présente au point M . Cela permet de généraliser au cas où le milieu extérieur est plus compliqué qu'une masse ponctuelle placée en P . Le champ gravitationnel a alors une expression différente.

b) Le poids

Lorsqu'on étudie des mouvements dans le référentiel terrestre, on tient compte d'une force à distance appelée le **poids**. Cette force est définie dans le référentiel terrestre local comme la force exercée sur un système de masse m par la Terre. Elle s'applique au centre de gravité du système de masse m lorsqu'on étudie son mouvement au voisinage de la Terre dans le référentiel terrestre. La verticale est définie par la direction du poids : pour la repérer, on utilise un fil à plomb qui se place verticalement à l'équilibre. Le poids est vertical, dirigé vers le bas et de norme mg .

La force de pesanteur terrestre exercée sur un système de masse m au voisinage de la Terre s'applique au centre de gravité du système et vaut :

$$\vec{P} = m \vec{g}$$

où $\vec{g}$ est l'accélération de la pesanteur, verticale vers le bas. $\|\vec{g}\| = 9,8 \text{ m} \cdot \text{s}^{-2}$.

Modélisation du poids Dans ce qui suit, on se restreint au cas le plus simple où le système est réduit à un point matériel M de masse m qui coïncide alors avec son centre de gravité. En admettant que la force gravitationnelle exercée par la Terre sur M est la même que si

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONNIENNE

toute la masse de la Terre était concentrée en son centre P , la force gravitationnelle exercée par la Terre est alors $\overrightarrow{F_{\text{Terre} \rightarrow M}} = m \overrightarrow{g}_{\text{gravi, Terre}} \simeq m \left(-\frac{\mathcal{G} m_{\text{Terre}}}{R_T^2} \overrightarrow{u_{PM}} \right)$ au voisinage du sol. En gardant trois chiffres significatifs pour :

- $R_T = 6371 \text{ km} \simeq 6,37 \cdot 10^6 \text{ m}$, le rayon terrestre moyen ;
- $\mathcal{G} = 6,67 \cdot 10^{-11} \text{ N} \cdot \text{m}^2 \cdot \text{kg}^{-2}$, la constante gravitationnelle ;
- $M_T = 5,97 \cdot 10^{24} \text{ kg}$, la masse de la Terre,

on trouve $\|\overrightarrow{g}_{\text{gravi, Terre}}\| = 6,67 \cdot 10^{-11} \frac{5,97 \cdot 10^{24}}{(6,37 \cdot 10^6)^2} = 9,81 \text{ m} \cdot \text{s}^{-2}$.

Ce modèle explique le champ de pesanteur avec une précision relative de l'ordre de 10^{-3} .



En référentiel terrestre, on tient compte du poids qui inclue l'attraction gravitationnelle terrestre. Il ne faut jamais compter les deux forces en même temps.

Remarque

Pour modéliser l'accélération de la pesanteur g avec une précision relative supérieure 10^{-3} , il faut tenir compte de la rotation de la Terre sur elle-même autour de son axe Nord-Sud en 1 jour qui a deux effets d'intensité comparable. En premier lieu, la rotation entraîne un aplatissement de la Terre au niveau des pôles. L'aplatissement est de 21 km : son rayon équatorial vaut 6378 km et son rayon polaire 6357 km. En second lieu, la rotation provoque l'apparition d'une « force centrifuge » que l'on doit ajouter vectoriellement à la force gravitationnelle. Au final, g varie en fonction de la latitude λ selon une formule adoptée en 1967 par l'*Union Géodésique et Géophysique Internationale* :

$$g = g_0 (1 + k_1 \sin^2(\lambda) + k_2 \sin^2(2\lambda))$$

avec $g_0 = 9,780318 \text{ m} \cdot \text{s}^{-2}$, $k_1 = 5,3024 \cdot 10^{-3}$, $k_2 = -5,8 \cdot 10^{-6}$.

Par ailleurs, g diminue de $3 \cdot 10^{-6} \text{ m} \cdot \text{s}^{-2}$ par mètre d'altitude .

c) Force électrostatique

Cette force s'applique au point matériel M de charge q placé dans un champ électrostatique $\overrightarrow{E}_{\text{ext}}$ créé en M par le milieu extérieur. Elle est proportionnelle à la charge du point matériel M est s'annule lorsque le point M n'est pas chargé :

$$\overrightarrow{F} = q \overrightarrow{E}_{\text{ext}}.$$

Cas de l'interaction coulombienne Dans le cas où le milieu extérieur est réduit à une charge ponctuelle, on parle d'**interaction coulombienne**.

La force coulombienne exercée par une charge q_p placée en P sur la charge q placée en M est :

- attractive si q et q_p sont de signes opposés ;
- répulsive si q et q_p sont de même signe ;
- de norme proportionnelle à q et q_p , ainsi qu'à l'inverse du carré de la distance PM .

Elle s'écrit :

$$\overrightarrow{F_{P \rightarrow M}} = \frac{qq_p}{4\pi\epsilon_0} \frac{\overrightarrow{PM}}{PM^3} = \frac{qq_p}{4\pi\epsilon_0} \frac{\overrightarrow{u_{PM}}}{PM^2}$$

où $\overrightarrow{u_{PM}}$ est le vecteur unitaire dirigé de P vers M et $\epsilon_0 = 8,85 \cdot 10^{-12} \text{ F} \cdot \text{m}^{-1}$ la **permittivité du vide**.

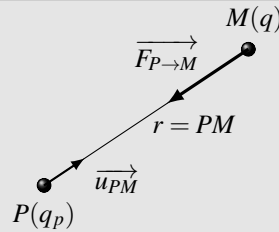


Figure 15.3 –
Interaction
coulombienne entre
deux charges q_p et q
placées en P et M . Sur
le schéma, q_p et q sont
de signes opposés.

Remarque

On peut aussi écrire $\overrightarrow{F_{P \rightarrow M}} = q \left(\frac{q_p}{4\pi\epsilon_0 PM^2} \overrightarrow{u_{PM}} \right) = q \vec{E}_{\text{ext}}$ en introduisant le champ électrique créé en M par la charge q_p placée en P : $\vec{E}_{\text{ext}} = \frac{q_p}{4\pi\epsilon_0 PM^2} \overrightarrow{u_{PM}}$. Cette notation présente l'avantage de séparer ce qui dépend de la charge q de ce qui dépend du milieu extérieur. De ce fait le champ électrostatique $\vec{E}_{\text{ext}}$ existe même si la charge q n'est pas présente au point M . On peut alors généraliser cette expression au cas où le milieu extérieur est plus compliqué qu'une charge ponctuelle placée en P . Le champ électrique a alors une expression différente qui sera donné.

d) Généralisation : force électromagnétique ou force de Lorentz

Cette force s'applique au point matériel M de charge q placé dans un champ électromagnétique $(\vec{E}, \vec{B})$ créé par le milieu extérieur. Cette force généralise le cas précédent dans le cas où il y a à la fois un champ électrique et un champ magnétique. Elle fera l'objet d'une étude spécifique dans un prochain chapitre. À ce niveau du cours, les champs $\vec{E}$ et $\vec{B}$ seront donnés.

La force de Lorentz exercée sur un système M de charge q par le milieu extérieur s'écrit :

$$\vec{F} = q \left(\vec{E} + \vec{v} \wedge \vec{B} \right),$$

où $\vec{v}$ est la vitesse de M , $\vec{E}$ le champ électrostatique créé en M , $\vec{B}$ le champ magnétique créé en M .

 $\vec{v} \wedge \vec{B}$ est le produit vectoriel de $\vec{v}$ par $\vec{B}$ défini dans l'appendice mathématique.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

5.3 Forces de contact

Par opposition aux forces étudiées au paragraphe précédent qui avaient une action à distance, on s’intéresse maintenant aux actions qui n’existent que lors d’un contact entre systèmes.

a) Définition des forces de contact

Lorsqu’un point matériel n’est soumis qu’à des forces à distance, on dit qu’il est libre : sa trajectoire n’est pas astreinte à rester dans une zone plus ou moins confinée de l’espace. C’est le cas par exemple d’un corps qui tombe dans le champ de pesanteur en l’absence de tout frottement. On s’intéresse maintenant aux forces qui n’interviennent qu’en présence d’un contact du point matériel avec un solide ou un fluide. On aura alors essentiellement des forces de liaison et des forces de frottement.

Il n’existe pas de théorie permettant de déterminer les forces de contact à l’aide de considérations microscopiques que l’on pourrait étendre à l’échelle macroscopique. Tout ce qui va suivre ne sera donc qu’une approche phénoménologique, c’est-à-dire obtenue expérimentalement, et valable uniquement dans le même contexte.

b) Tension d’un fil

Un fil de masse négligeable prend une forme rectiligne dès qu’il est tendu. Il exerce alors sur un objet accroché à une de ses extrémités une **tension** notée $\vec{T}$. La direction de cette force est celle du fil. Elle est toujours dirigée d’une extrémité du fil vers l’autre : un fil peut tirer un objet mais pas le repousser. La norme de la tension est *a priori* indéterminée, sa valeur dépend des autres forces. Dans le cas où le fil n’est pas tendu, la tension est nulle.

La force de tension exercée par un fil tendu sur un objet accroché à l’une de ses extrémités vaut :

$$\vec{T} = -T\vec{u}_{\text{ext}} \quad \text{où :}$$

- $\vec{u}_{\text{ext}}$ est un vecteur unitaire parallèle au fil tendu, toujours orienté vers l’extérieur du fil,
- $T > 0$ est la norme de cette tension.

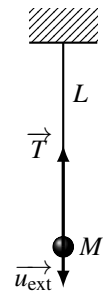


Figure 15.4
– Le point M accroché à un fil tendu subit la tension $\vec{T}$ du fil.

c) Force de rappel élastique exercée par un ressort

Les ressorts envisagés sont supposés idéaux : on peut négliger leur masse devant celle du point M attaché à leur extrémité et, après une elongation ou une compression, ils reprennent leur forme et leur longueur initiales. La force qu’un ressort exerce sur un objet auquel il est accroché s’applique au point d’attache, le long du ressort, dans une direction opposée à son « allongement » (qui pourra être un réel allongement ou une compression) et son intensité est proportionnelle à l’allongement.

L'allongement $\Delta l = l - l_0$ est une grandeur algébrique :

- elle est positive si le ressort est étiré : $\Delta l = l - l_0 > 0$;
- elle est négative si le ressort est comprimé : $\Delta l = l - l_0 < 0$.

La force s'oppose à la déformation du ressort par rapport à sa forme au repos (c'est-à-dire sans déformation) : c'est une force de rappel.

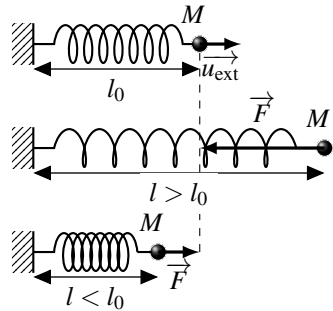


Figure 15.5 – Tension d'un ressort.

La force de rappel élastique vaut : $\vec{F} = -k\Delta l \vec{u}_{\text{ext}}$ où :

- $\vec{u}_{\text{ext}}$ est un vecteur unitaire parallèle au ressort, orienté vers l'extérieur du ressort,
- $\Delta l = l - l_0$ est l'allongement du ressort en notant respectivement l et l_0 la longueur du ressort et sa longueur à vide,
- k est une constante de proportionnalité caractéristique du ressort utilisé et appelée **raideur du ressort**, exprimée en $\text{N}\cdot\text{m}^{-1}$ dans les unités du système international.

d) Action exercée par un support solide

Cette force s'applique au point matériel M de masse m en contact avec le milieu extérieur par un contact solide. Le point M peut être posé sur un support ou enfilé sur une tige.



Figure 15.6 – À gauche, la force $\vec{f}$ est insuffisante pour mettre le solide en mouvement sur son support horizontal. Elle est compensée par les frottements solides $\vec{R}_T$. À droite, la force de poussée $\vec{f}$ est plus grande et on a dépassé le seuil d'adhérence. Le solide glisse sur son support.

Exemple introductif On étudie un solide posé sur un plan horizontal que l'on pousse avec une force $\vec{f}$ (figure 15.6). Lorsque la poussée est faible le solide ne se met pas en mouvement : la force de frottement solide l'empêche de se déplacer. Si l'on pousse plus fort, la force de frottement solide s'adapte et empêche le mouvement. Lorsque l'on pousse encore plus fort, on arrive à une force seuil au delà de laquelle le solide se met à glisser. On parle de seuil d'adhérence. Au delà de ce seuil, la force de frottement solide est constante et le solide glisse.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

Paramétrisation de la force de contact On considère que le contact entre l'objet M et le support solide permet de définir un point de contact I et un plan de contact $\mathcal{P}$. On caractérise le plan de contact par son vecteur normal unitaire $\vec{n}$ que l'on oriente du support vers l'objet.

On décompose alors la force de contact $\vec{R}$ en deux :

$$\vec{R} = \vec{R}_T + \vec{R}_N,$$

où les composantes tangentielle $\vec{R}_T$ et normale $\vec{R}_N$ de la force de contact ont les propriétés suivantes :

- $\vec{R}_N$ est la réaction normale du support. Sa norme est indéterminée *a priori* et dépend des autres forces. Sa direction et son sens sont ceux de $\vec{n}$: elle s'oppose à la pénétration de l'objet dans le support. Elle est orientée du support vers le système.
- $\vec{R}_T$ est la réaction tangentielle du support. Elle appartient au plan de contact. Elle correspond à une force de frottement solide entre les surfaces en contact. Comme toute force de frottement, elle s'oppose au mouvement.

Pour déterminer la composante normale de la force, on utilise le principe fondamental de la dynamique en projection sur $\vec{n}$. Quant aux forces de frottement, elles seront données si nécessaire.

 Les forces de frottement dépendent de la nature des matériaux et de l'état de rugosité des surfaces. Lorsqu'il n'y a pas de frottements, la force de contact est réduite à $\vec{R}_N$. Elle est donc perpendiculaire au support.

Remarque

La force de contact est une force de liaison entre le mobile et le support. Pour un anneau enfilé sur une tige par exemple, c'est elle qui guide le mobile le long de la tige. Pour un mobile se déplaçant sur un plan horizontal, c'est elle qui assure la planéité du mouvement en s'opposant aux autres forces verticales telles que le poids.

 Lorsque l'on cherche à montrer que le contact est rompu entre le système et son support, on peut déterminer la **rupture de contact** par le fait que $\vec{R} = \vec{0}$.

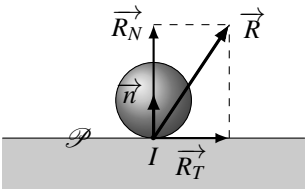


Figure 15.7 – Force de contact entre un système et un support. Le point de contact I , le plan de contact $\mathcal{P}$ et sa normale $\vec{n}$ sont représentés, ainsi que la force de contact $\vec{R} = \vec{R}_N + \vec{R}_T$.

e) Action exercée par un fluide

Il existe principalement trois forces exercées par un fluide sur un solide plongé dans ce fluide : **la poussée d'Archimède, la force de traînée, la force de portance**. Ces forces sont toutes dues au contact du fluide environnant sur le corps immergé. La poussée d'Archimède, qui s'exerce sur les solides plongés dans un fluide, existe même si le solide est au repos par rapport au fluide. Les forces de traînée et de portance n'existent que lorsque le solide est en mouvement par rapport au fluide. Elles dépendent de la vitesse relative $\vec{v}$ du solide par rapport au fluide.

La poussée d’Archimède Tout corps plongé dans un fluide au repos, entièrement ou partiellement immergé dans celui-ci, subit une force appelée « poussée d’Archimède ». Cette force verticale, dirigée de bas en haut est de norme égale au poids du volume de fluide déplacé. En posant $V_{\text{immergé}}$, le volume du corps immergé dans le fluide de masse volumique ρ_{fluide} , elle s’exprime alors :

$$\vec{\Pi} = -\rho_{\text{fluide}} V_{\text{immergé}} \vec{g}.$$

La poussée d’Archimède est la résultante des forces de pression qui agissent sur le corps immergé.
Si système est entièrement immergé, le volume immergé $V_{\text{immergé}}$ est égal au volume V du corps. Il est alors souvent intéressant d’exprimer la poussée d’Archimède en fonction du poids du système.

Pour cela, on exprime le volume $V_{\text{immergé}}$ en fonction de la masse volumique moyenne ρ du système :

$$V_{\text{immergé}} = V = \frac{m}{\rho} \implies \vec{\Pi} = -\frac{\rho_{\text{fluide}}}{\rho} m \vec{g}.$$

On peut alors sommer le poids et la poussée d’Archimède pour obtenir :

$$\vec{\Pi} + \vec{P} = m \left(1 - \frac{\rho_{\text{fluide}}}{\rho} \right) \vec{g} = m \vec{g}'$$

avec $\vec{g}' = \left(1 - \frac{\rho_{\text{fluide}}}{\rho} \right) \vec{g}$ la gravité réduite dans le fluide.

On peut alors tenir compte du poids et de la poussée d’Archimède en remplaçant $\vec{g}$ par $\vec{g}'$.
Pour un solide plein plongé dans l’eau, la masse volumique du fluide vaut $\rho_{\text{eau}} = 10^3 \text{ kg}\cdot\text{m}^{-3}$ et celle du système est comprise entre 10^3 et $10^4 \text{ kg}\cdot\text{m}^{-3}$ d’où $\vec{g}' \neq \vec{g}$. La poussée d’Archimède est importante.

Exemple

Le corps humain étant essentiellement constitué d’eau, sa masse volumique est voisine de $10^3 \text{ kg}\cdot\text{m}^{-3}$. Plongé dans l’eau, on ressent une gravité réduite quasiment nulle, c’est-à-dire un état d’apesanteur.

Pour un solide plein plongé dans l’air, la masse volumique du fluide vaut $\rho_{\text{air}} = 1 \text{ kg}\cdot\text{m}^{-3}$ et celle du système est comprise entre 10^3 et $10^4 \text{ kg}\cdot\text{m}^{-3}$ d’où $\vec{g}' \simeq \vec{g}$. La poussée d’Archimède est négligeable.

Exemple

Pour flotter dans l’air à l’aide de la poussée d’Archimède, il faut réussir à annuler $\vec{g}'$ donc à obtenir une masse volumique moyenne ρ du système égale à la densité de l’air. Il faut donc éviter le solide pour obtenir un solide creux de grand volume et de faible masse. C’est le principe de fonctionnement des montgolfières et des dirigeables.

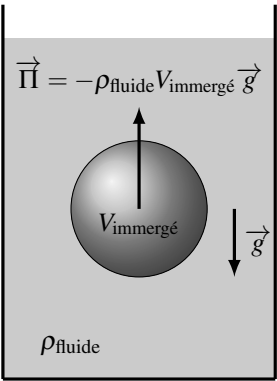


Figure 15.8 –
Poussée
d’Archimède.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

Forces de traînée et de portance Un corps en mouvement relatif par rapport à un fluide subit des forces aérodynamiques (dans l'air) ou hydrodynamiques (dans l'eau) : la force de traînée et la force de portance.

- La traînée est opposée au mouvement du mobile par rapport au fluide. C'est la force principale contre laquelle on lutte quand on roule en vélo. Elle est d'autant plus grande que la vitesse relative est grande : plus on roule vite, plus il faut pédaler fort.
- La portance est perpendiculaire au mouvement du mobile par rapport au fluide. C'est grâce à elle que les avions volent.

La force de traînée exercée sur M par le milieu fluide environnant a pour expression :

- $\vec{f} = -k_1 \vec{v}$ avec $k_1 > 0$ à faible vitesse.
 k_1 est ici un **coefficient de frottement fluide** qui s'exprime en $\text{kg} \cdot \text{s}^{-1}$.
- $\vec{f} = -k_2 \|\vec{v}\| \vec{v}$ avec $k_2 > 0$ à haute vitesse.
 k_2 est ici un **coefficient de frottement fluide** qui s'exprime en $\text{kg} \cdot \text{m}^{-1}$.

Les coefficients de frottement k_1 et k_2 dépendent de l'écoulement autour du système. Ils sont déterminés expérimentalement parfois à l'aide d'essais en soufflerie. On peut aussi les évaluer grâce à des simulations numériques. *A priori*, dans les gaz, les frottements sont proportionnels au carré de la vitesse. Dans les liquides, on peut assez souvent utiliser des frottements proportionnels à la vitesse. Dans le doute, on peut toujours essayer un modèle avec des frottements proportionnels à la vitesse qui a l'avantage d'être linéaire. Si les résultats obtenus par ce modèle linéaire ne sont pas satisfaisants, on doit utiliser la deuxième forme. Dans ce cas, il faut souvent envisager une résolution numérique du problème.

6 Résolution d'un problème de mécanique du point

Pour résoudre un problème de mécanique du point, il est judicieux de suivre les étapes suivantes. Tous les points doivent apparaître et la succession de ces étapes dans cet ordre permet de suivre une démarche de démonstration scientifique rigoureuse qui consiste à faire des hypothèses, vérifier que l'on est dans les conditions d'application d'un principe, d'une loi physique ou d'un théorème, puis appliquer cette loi et en tirer des conclusions. Pour cela :

- la première étape consiste à **définir le système** et à vérifier que l'on peut l'assimiler à un point matériel ;
- lors de la deuxième étape, on **choisit le référentiel** qui sera utilisé et on vérifie que l'on peut bien le considérer comme galiléen ;
- on commence ensuite à **établir ou compléter un schéma** de la situation qui va permettre de **choisir le système de coordonnées adapté** à l'étude du mouvement ;
- on établit alors un **bilan des forces** précis et complet ;
- la dernière étape consiste à **choisir la méthode de résolution**. Pour l'instant, on ne dispose que d'une méthode : le principe fondamental de la dynamique. On disposera par la suite de deux lois qui en découlent et il faudra être capable de choisir la méthode qui permet de résoudre le problème avec le plus de facilité. Cette étape est forcément la dernière puisqu'elle requiert de connaître la réponse aux points précédents.

7 Chute libre dans le champ de pesanteur

7.1 Mise en équation

On étudie le mouvement d'un point matériel M de masse m qu'on lâche sans vitesse initiale d'une hauteur h dans le champ de pesanteur $\vec{g}$.

Dans un premier temps, on néglige la résistance de l'air au cours du mouvement, ce qui revient à considérer que la chute se fait dans le vide. On tiendra compte par la suite de frottements proportionnels à la vitesse puis de frottements proportionnels au carré de la vitesse. Dans tous les cas, le mouvement est un mouvement de chute verticale et on utilise une base de projection cartésienne réduite à un axe. On choisit l'axe (Oy) vertical ascendant et l'étude cinématique du mouvement rectiligne selon l'axe (Oy) donne :

$$\overrightarrow{OM} = y\vec{u}_y, \quad \vec{v} = \dot{y}\vec{u}_y \quad \text{et} \quad \vec{a} = \ddot{y}\vec{u}_y.$$

On établit le schéma de la figure 15.9 en faisant apparaître le point M à $t = 0$ avec les conditions initiales du mouvement. On représente également le point M à l'instant t avec les forces qui s'y appliquent.

Le choix du système de coordonnées respecte les symétries du problème et en l'occurrence des conditions initiales et des forces qui s'appliquent sur M . On suit alors la procédure de résolution d'un problème de mécanique décrite ci-dessus :

- définition du système : le point matériel M ;
- choix du référentiel : le référentiel terrestre que l'on considère galiléen ;
- bilan des forces : les forces qui s'exercent sur le point matériel sont le poids $\vec{P} = m\vec{g}$ et les forces de frottements fluides $\vec{f}$;
- choix de la méthode de résolution : on utilise le principe fondamental de la dynamique :

$$m\vec{a} = \sum_i \vec{f}_i = m\vec{g} + \vec{f}$$

soit

$$\vec{a} = \vec{g} + \frac{\vec{f}}{m}. \tag{15.2}$$

7.2 Chute libre dans le vide

Lorsque l'on néglige la résistance de l'air, on parle de chute libre dans le vide. Dans ces conditions, le principe fondamental de la dynamique se réduit à :

$$\vec{a} = \vec{g}. \tag{15.3}$$

Le mouvement est donc caractérisé par une accélération constante. Ce type de mouvement a

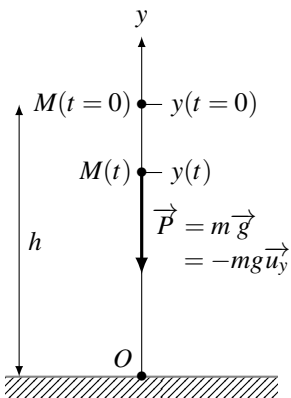


Figure 15.9 – Situation de chute libre sous l'action du poids.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

été étudié dans le chapitre de cinématique du point. Il est remarquable de noter que la masse n'intervient pas dans les équations du mouvement. Il faut cependant avoir conscience que cela n'est vrai que parce que l'on a négligé la résistance de l'air.

a) Équations horaires du mouvement

Comme $\vec{a} = \ddot{y}\vec{u}_y$ et $\vec{g} = -g\vec{u}_y$, l'équation du mouvement obtenue par projection de l'équation (15.3) sur l'axe $\vec{u}_y$ se traduit par :

$$\ddot{y} = -g.$$

Par intégrations successives, et en tenant compte des conditions initiales (vitesse nulle et altitude h), on trouve :

$$\dot{y} = -gt \quad \text{puis} \quad y = -\frac{1}{2}gt^2 + h.$$

Le mouvement est un mouvement rectiligne uniformément accéléré le long de l'axe vertical (Oy) dirigé vers le bas à la vitesse de norme :

$$v = -\dot{y} = gt.$$

b) Caractéristiques du mouvement

On peut étudier quelques caractéristiques du mouvement.

- La durée de chute T est obtenue en cherchant l'instant d'arrivée au sol donc en résolvant l'équation $y(T) = 0$:

$$-\frac{1}{2}gT^2 + h = 0 \quad \Longleftrightarrow \quad T = \sqrt{\frac{2h}{g}}.$$

- La vitesse du point matériel à l'impact au sol est la vitesse à l'instant $t = T$ soit :

$$v = gT = \sqrt{2gh}.$$

Cette dernière relation montre que la vitesse à l'impact au sol augmente indéfiniment en fonction de la hauteur de chute h . Pour une hauteur de chute de 1000 m, on trouve une vitesse à l'impact de $140 \text{ m}\cdot\text{s}^{-1} \simeq 500 \text{ km}\cdot\text{h}^{-1}$. Or les sauteurs en parachute qui pratiquent la chute libre ne dépassent pas $250 \text{ km}\cdot\text{h}^{-1}$ lors d'une chute d'une hauteur supérieure à 1000 m. **Un phénomène limite la vitesse qui n'augmente pas au delà d'une valeur limite.** Grossièrement, lors d'une chute libre, la vitesse du parachutiste augmente pendant 600 m qu'il parcourt en 20 s pour atteindre $60 \text{ m}\cdot\text{s}^{-1}$. Après l'ouverture du parachute, la vitesse diminue jusqu'à environ $7 \text{ m}\cdot\text{s}^{-1}$. La voile du parachute sert à augmenter la résistance de l'air et permet de diminuer la vitesse de chute. Nous allons étudier le phénomène qui limite la vitesse à l'aide de deux modèles de résistance de l'air : un modèle de frottement proportionnel à la vitesse et un modèle de frottement proportionnel au carré de la vitesse.

7.3 Chute libre avec frottements proportionnels à la vitesse

On ne néglige plus la force de frottement $\vec{f}$ et on la prend sous la forme $\vec{f} = -k_1 \vec{v}$. Le principe fondamental de la dynamique s'écrit alors :

$$m \vec{a} = -k_1 \vec{v} + m \vec{g}.$$

La chute ayant lieu vers le bas, y est négatif. La norme de la vitesse étant toujours positive, $v = -\dot{y}$. Le vecteur vitesse $\vec{v}$ vaut donc $\vec{v} = \dot{y} \vec{u}_y = -v \vec{u}_y$ et on en déduit que :

$$\vec{a} = \dot{y} \vec{u}_y = -\frac{dv}{dt} \vec{u}_y.$$

Comme $\vec{g} = -g \vec{u}_y$, le principe fondamental de la dynamique projeté sur l'axe (Oy) donne l'équation du mouvement :

$$m \frac{dv}{dt} + k_1 v = mg.$$

a) Analyse de l'équation différentielle du mouvement

Deux forces sont en présence :

- une force motrice, le poids, orientée vers le bas et responsable du mouvement de chute ;
- une force de freinage, le frottement fluide, opposée au mouvement dont la norme augmente lorsque la vitesse augmente.

À partir d'une vitesse initiale nulle, le point M chute de plus en plus vite et la force de freinage proportionnelle à la vitesse augmente. Au bout d'un moment, la force motrice et la force de freinage se compensent. Dans ces conditions, la somme des forces s'annule, le système est pseudo-isolé et son mouvement est par la suite rectiligne et uniforme à la vitesse limite v_l qui assure l'égalité du poids et de la force de frottement :

$$k_1 v = mg \quad \Longleftrightarrow \quad v = v_l = \frac{mg}{k_1}.$$

Très souvent, il est intéressant d'introduire les paramètres physiques que l'on a obtenu de cette manière dans l'équation différentielle pour faciliter sa résolution. Dans le cas présent, on trouve :

$$\frac{m}{k_1} \frac{dv}{dt} + v = \frac{mg}{k_1} \quad \Longleftrightarrow \quad \frac{v_l}{g} \frac{dv}{dt} + v = v_l$$

de la forme

$$\tau \frac{dv}{dt} + v = v_l \quad \text{avec} \quad \tau = \frac{v_l}{g} = \frac{m}{k_1}.$$

On reconnaît une équation différentielle linéaire du premier ordre dont le temps caractéristique τ correspond à l'échelle de temps nécessaire pour atteindre un régime permanent caractérisé par $v = v_l$.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

b) Résolution de l'équation différentielle du mouvement

La solution de cette équation est la somme de la solution générale de l'équation homogène associée $v_h(t)$ et d'une solution particulière $v_p(t)$.

La résolution de l'équation homogène associée conduit à $v_h(t) = \lambda \exp\left(-\frac{t}{\tau}\right)$ où λ est une constante. On recherche une solution particulière sous la forme d'une constante puisque le second membre est constant d'où $v_p(t) = g\tau = v_l$. La solution générale est alors :

$$v(t) = C \exp\left(-\frac{t}{\tau}\right) + v_l.$$

On détermine la constante à l'aide des conditions initiales à savoir $v(t=0) = 0 = \lambda + v_l$ soit $\lambda = -v_l$ et on obtient :

$$v(t) = v_l \left(1 - \exp\left(-\frac{t}{\tau}\right)\right).$$

La chute ayant lieu vers le bas, $\dot{y}$ est négatif. La norme de la vitesse v étant toujours positive, $v = -\dot{y}$. La position y se déduit alors par intégration par rapport au temps de $\dot{y} = -v$:

$$\dot{y} = -v = -v_l \left(1 - \exp\left(-\frac{t}{\tau}\right)\right) \implies y = -v_l t + v_l \tau \exp\left(-\frac{t}{\tau}\right) + K.$$

La constante K s'obtient en tenant compte de la position initiale : $y(t=0) = h = -v_l \tau + K$ soit $K = h + v_l \tau$ et finalement :

$$y(t) = -v_l t + v_l \tau \left(1 - \exp\left(-\frac{t}{\tau}\right)\right) + h.$$

c) Retour sur l'analyse initiale

Pour analyser l'expression de la vitesse en fonction du temps : $v(t) = v_l \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$, on définit une vitesse et un temps sans dimension en notant $v^* = \frac{v}{v_l}$ et $t^* = \frac{t}{\tau}$. On obtient alors :

$$v^*(t) = 1 - \exp(-t^*),$$

et on trace l'évolution de v^* en fonction de t^* qui sur la figure 15.10.

On remarque que v^* augmente de 0 à 1 qu'elle atteint pratiquement lorsque t^* atteint environ 5. Cela signifie que la vitesse du mobile augmente vers sa vitesse limite v_l qui est atteinte après quelques τ . La vitesse évolue en deux phases : une phase transitoire qui dure quelques τ et permet au mobile d'accélérer jusqu'à v_l puis une phase permanente lors de laquelle la vitesse est constante et égale à v_l :

- $v_l = \frac{m}{k_1} g$ représente la vitesse finale du mobile ;
- $\tau = \frac{v_l}{g} = \frac{m}{k_1}$ est le temps caractéristique nécessaire pour que v_l soit atteinte.

Ces deux paramètres sont indépendants des conditions initiales et déterminent entièrement l'évolution temporelle de la vitesse de chute du système. Les paramètres m , g et k_1 n'interviennent pas indépendamment mais seulement sous la forme de combinaisons formant v_l et τ .

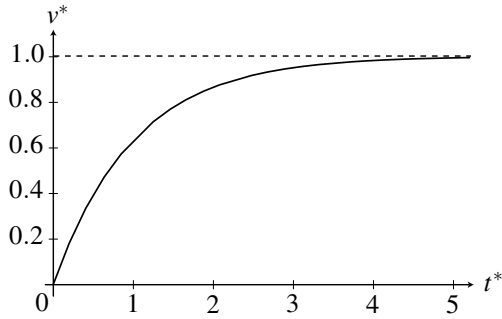


Figure 15.10 – Évolution de la vitesse en fonction du temps lors d’une chute libre avec frottements proportionnels à la vitesse.

7.4 Chute libre avec frottements proportionnels au carré de la vitesse

On considère à présent que la force de frottement $\vec{f}$ se met sous la forme $\vec{f} = -k_2 v \vec{v}$, en notant, ici encore, v la norme de la vitesse. Le principe fondamental de la dynamique s’écrit alors :

$$m \frac{d\vec{v}}{dt} = m \vec{g} - k_2 v \vec{v}. \tag{15.4}$$

Comme précédemment, la chute a lieu vers le bas, $\vec{v} = -v \vec{u}_y$, $\vec{a} = -\frac{dv}{dt} \vec{u}_y$ et $\vec{g} = -g \vec{u}_y$. L’équation du mouvement est obtenue par projection du principe fondamental de la dynamique sur l’axe (Oy) :

$$m \frac{dv}{dt} + k_2 v^2 = mg.$$

Cette équation n’est pas linéaire, on ne peut donc pas la résoudre comme la précédente.

a) Analyse physique de l’équation différentielle du mouvement

Pour les mêmes raisons qu’en présence de frottements proportionnels à la vitesse, le point matériel atteint une vitesse limite v_l lorsque le poids est compensé par la force de frottement d’où :

$$k_2 v^2 = mg \quad \Longleftrightarrow \quad v = v_l = \sqrt{\frac{mg}{k_2}}.$$

On introduit alors la vitesse limite dans l’équation différentielle qui devient :

$$\frac{m}{k_2} \frac{dv}{dt} + v^2 = v_l^2 \quad \Longleftrightarrow \quad \frac{m}{k_2 v_l^2} \frac{dv}{dt} = 1 - \left(\frac{v}{v_l}\right)^2 \quad \Longleftrightarrow \quad \frac{1}{g} \frac{dv}{dt} = 1 - \left(\frac{v}{v_l}\right)^2.$$

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

puisque $\frac{m}{k_2 v_l^2} = \frac{1}{g}$. On aboutit à une équation différentielle à variables séparables :

$$\frac{dv}{1 - \left(\frac{v}{v_l}\right)^2} = g dt \quad \text{soit} \quad \frac{d\left(\frac{v}{v_l}\right)}{1 - \left(\frac{v}{v_l}\right)^2} = \frac{g}{v_l} dt.$$

On définit alors la vitesse adimensionnée $v^* = \frac{v}{v_l}$ et le temps caractéristique $\tau = \frac{v_l}{g}$ dont on se sert pour établir un temps adimensionné $t^* = \frac{t}{\tau}$. On obtient alors l'équation :

$$\frac{dv^*}{1 - v^{*2}} = dt^* \implies \int_{v^*(t=0)=0}^{v^*(t)} \frac{dv^*}{1 - v^{*2}} = \int_0^{t^*} dt^*.$$

On sait calculer cette intégrale analytiquement ou numériquement. Cela permet de représenter la courbe de la figure 15.11 sur laquelle on remarque que v^* augmente de 0 à 1 qu'elle atteint pratiquement lorsque t^* atteint environ 3. Cela signifie que la vitesse du mobile augmente vers sa vitesse limite v_l qui est atteinte après quelques τ .

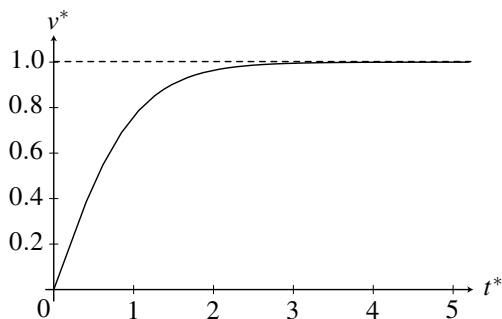


Figure 15.11 – Évolution de la vitesse en fonction du temps lors d'une chute libre avec frottements proportionnels au carré de la vitesse.

Remarque

La solution analytique, obtenue à l'aide de la méthode de décomposition en éléments simples vue en mathématique, est $v^* = \frac{\exp(t^*) - \exp(-t^*)}{\exp(t^*) + \exp(-t^*)} = \frac{\text{sh}(t^*)}{\text{ch}(t^*)} = \text{th}(t^*)$ où sh, ch et th sont les fonctions trigonométriques hyperboliques définies dans l'annexe mathématique.

b) Retour sur l'analyse initiale

Comme précédemment, la vitesse du mobile tend vers la vitesse limite lorsque le temps tend vers l'infini et évolue en deux phases : une phase de régime transitoire qui dure quelques τ et

permet au mobile d'accélérer jusqu'à v_l puis une phase de régime permanent lors de laquelle la vitesse est constante et égale à v_l :

- $v_l = \sqrt{\frac{mg}{k_2}}$ représente la vitesse finale du mobile ;
- $\tau = \frac{v_l}{g} = \sqrt{\frac{m}{k_2 g}}$ est le temps caractéristique nécessaire pour que v_l soit atteinte.

Ces deux paramètres sont indépendants des conditions initiales et déterminent entièrement l'évolution temporelle de la vitesse de chute du système. Les paramètres m , g , k_2 n'interviennent pas indépendamment mais seulement sous la forme de combinaisons formant v_l et τ .

7.5 Comparaison des deux modèles de frottements

Dans les deux cas, le mouvement comporte une phase d'accélération d'une durée de quelques τ suivie d'une évolution à vitesse constante et égale à la vitesse limite ce qui est bien en accord avec les observations de chutes libres de parachutistes. Cependant, l'expression de cette vitesse n'est pas la même :

$$\begin{cases} v_l = \frac{mg}{k_1} & \text{pour un frottement proportionnel à la vitesse,} \\ v_l = \sqrt{\frac{mg}{k_2}} & \text{pour un frottement proportionnel au carré de la vitesse.} \end{cases}$$

En pratique, c'est la mesure de ces vitesses limites qui permet de déterminer la valeur des constantes k_1 et k_2 et on ne peut pas discriminer ces modèles par cette seule étude. Pour déterminer lequel de ces deux modèles est le plus proche de la réalité, il est nécessaire de faire des essais en soufflerie ou de raisonner sur l'énergie. On obtient alors les résultats suivants :

- un frottement proportionnel à la vitesse convient bien lorsque les vitesses sont faibles, ou que la chute a lieu dans un liquide visqueux ;
- lorsque les vitesses deviennent plus importantes et que la chute a lieu dans un gaz comme l'air, un frottement proportionnel au carré de la vitesse est plus adapté.

La méthode employée dans le cas d'un frottement proportionnel au carré de la vitesse est intéressante : l'introduction d'une valeur limite pour la vitesse suivie de la séparation des variables v et t permet de transformer une équation différentielle *a priori* compliquée à résoudre en une intégrale « classique ». L'idée d'introduire un paramètre limite pour simplifier la résolution est souvent une idée à tester. Le fait de rendre les équations sans dimension avant de chercher à les résoudre également.

8 Tir d'un projectile dans le champ de pesanteur

8.1 Mise en équation

On étudie toujours le mouvement d'un point matériel M de masse m dans le champ de pesanteur $\vec{g}$. On s'intéresse maintenant au cas du vol balistique pour lequel le point matériel est initialement lancé avec une vitesse $\vec{v}_0$ faisant un angle α avec l'horizontale.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

Par rapport au mouvement de chute libre sans vitesse initiale vu précédemment, seules les conditions initiales sont modifiées. La position initiale du mobile est un point M_0 et son mouvement a lieu dans un plan vertical $\mathcal{P}$ contenant M_0 , $\vec{g}$ et $\vec{v}_0$.

Sans rien enlever à la généralité du problème, on choisit l'origine du repère O en M_0 , un axe (Oy) vertical ascendant et l'axe (Ox) contenu dans $\mathcal{P}$ et orthogonal à (Oy) et on utilise une base de projection cartésienne réduite à un plan (xOy) .

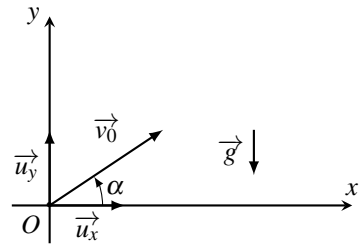


Figure 15.12 –
Conditions initiales
d'un tir dans le vide.

L'étude cinématique du mouvement dans le plan (xOy) en coordonnées cartésiennes donne :

$$\overrightarrow{OM} = \begin{pmatrix} x \\ y \end{pmatrix}, \quad \vec{v} = \begin{pmatrix} \dot{x} \\ \dot{y} \end{pmatrix}, \quad \vec{a} = \begin{pmatrix} \ddot{x} \\ \ddot{y} \end{pmatrix}.$$

Dans un premier temps, on néglige la résistance de l'air. On tiendra compte par la suite d'une force proportionnelle au carré de la vitesse qui modélise bien les frottements dus à un gaz.

8.2 Tir dans le vide

On néglige la résistance de l'air au cours du mouvement. Cela revient à considérer que le vol balistique se fait dans le vide. L'équation (15.3) établie au paragraphe 7.2 est inchangée :

$$\vec{a} = \vec{g}. \tag{15.5}$$

Le mouvement est à nouveau un mouvement à accélération constante étudié en cinématique du point. On la projette cette fois sur les axes (Ox) et (Oy) pour obtenir :

$$\begin{cases} \ddot{x} = 0 \\ \ddot{y} = -g. \end{cases}$$

Par intégrations successives en tenant compte des conditions initiales

$$\begin{cases} \dot{x}(t=0) = v_0 \cos \alpha \\ \dot{y}(t=0) = v_0 \sin \alpha \end{cases} \quad \text{et} \quad \begin{cases} x(t=0) = 0 \\ y(t=0) = 0, \end{cases}$$

on obtient :

$$\begin{cases} \dot{x} = v_0 \cos \alpha \\ \dot{y} = -gt + v_0 \sin \alpha \end{cases} \quad \text{puis} \quad \begin{cases} x = v_0 t \cos \alpha \\ y = -\frac{1}{2}gt^2 + v_0 t \sin \alpha. \end{cases}$$

L'équation de la trajectoire s'obtient en éliminant le temps dans les équations horaires.

De l'expression de x en fonction du temps, on tire : $t = \frac{x}{v_0 \cos \alpha}$, puis en reportant dans l'expression de y :

$$y = -\frac{g}{2v_0^2 \cos^2 \alpha} x^2 + x \tan \alpha.$$

La trajectoire est donc une parabole.

Remarque

On peut également intégrer l'équation vectorielle (15.5) et obtenir : $\vec{v} = \vec{g}t + \vec{v}_0$ puis $\vec{OM} = \frac{1}{2}\vec{g}t^2 + \vec{v}_0t + \vec{OM}_0$ ce qui permet de ne projeter que pour étudier la trajectoire.

Hauteur maximale atteinte Elle correspond au sommet de la parabole et s'obtient en résolvant l'équation : $\frac{dy}{dx} = \frac{-gx}{v_0^2 \cos^2 \alpha} + \tan \alpha = 0$ soit $x_m = \frac{v_0^2 \sin \alpha \cos \alpha}{g}$.

L'altitude maximale atteinte est alors : $y_m = \frac{v_0^2 \sin^2 \alpha}{2g}$.

Portée du tir La portée du tir correspond à l'endroit où le point matériel retombe sur le sol qu'on suppose plan et horizontal d'altitude $y = 0$. On obtient l'expression de la portée en résolvant l'équation :

$$0 = y = -\frac{g}{2v_0^2 \cos^2 \alpha} x^2 + x \tan \alpha = x \left(\tan \alpha - \frac{g}{2v_0^2 \cos^2 \alpha} x \right).$$

On obtient deux solutions : $x = 0$ et $x = \frac{v_0^2 \sin 2\alpha}{g}$.

La solution $x = 0$ correspond à la position d'origine du tir qui appartient naturellement à la trajectoire. Il est normal de la trouver mais elle ne présente pas d'intérêt ici. La portée du tir est donc :

$$x_p = \frac{v_0^2 \sin 2\alpha}{g}.$$

On note qu'elle est maximale pour $\alpha = \frac{\pi}{4}$.

Exemple

C'est l'angle avec lequel sautent les grenouilles pour retomber le plus loin possible.

Tir tendu ou tir en cloche On s'interroge maintenant sur la possibilité d'atteindre un point à la même altitude que celle où s'effectue le tir et situé à une distance donnée d de l'origine du tir (figure 15.13). Le seul paramètre variable est l'angle de tir α . D'après le paragraphe précédent, il faut chercher les valeurs de α vérifiant l'équation :

$$d = \frac{v_0^2 \sin 2\alpha}{g}.$$

La résolution de cette équation conduit à deux solutions dans l'intervalle $\left[0, \frac{\pi}{2}\right]$:

$$\alpha = \frac{1}{2} \arcsin \frac{gd}{v_0^2} \quad \text{ou} \quad \alpha = \frac{\pi}{2} - \frac{1}{2} \arcsin \frac{gd}{v_0^2}.$$

On obtient deux angles de tir possibles : la valeur la plus faible correspond au tir tendu tandis que la valeur la plus grande est celle du tir en cloche (voir figure 15.13).

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

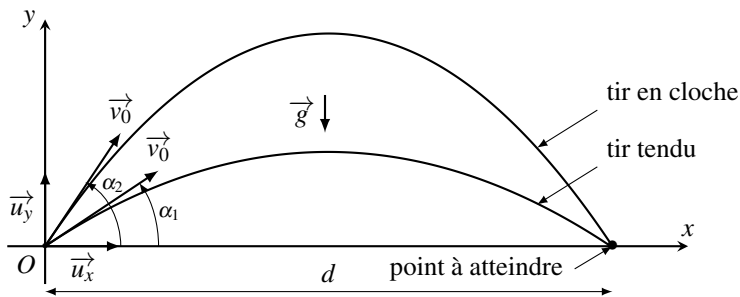


Figure 15.13 – Les deux trajectoires possibles d'un tir dans le vide.

Parabole de sûreté On cherche à déterminer l'ensemble des points (x_0, y_0) qu'il est possible d'atteindre pour une vitesse v_0 donnée. Ce sera le cas si l'équation d'inconnue α :

$$y_0 = -\frac{gx_0^2}{2v_0^2 \cos^2 \alpha} + x_0 \tan \alpha$$

admet des solutions. En utilisant le fait que $1 + \tan^2 \alpha = \frac{1}{\cos^2 \alpha}$, on aboutit à l'équation en $X = \tan \alpha$ suivante :

$$\frac{gx_0^2}{2v_0^2} X^2 - x_0 X + \frac{gx_0^2}{2v_0^2} + y_0 = 0.$$

Cette équation du second degré n'admet des solutions réelles (seules acceptables ici) que si son discriminant est positif donc si :

$$\Delta = x_0^2 - 4 \frac{gx_0^2}{2v_0^2} \left(\frac{gx_0^2}{2v_0^2} + y_0 \right) \geq 0 \iff y_0 \leq \frac{v_0^2}{2g} - \frac{gx_0^2}{2v_0^2}.$$

Un point ne peut donc être atteint que s'il se trouve dans le plan (xOz) sous la parabole d'équation :

$$y = \frac{v_0^2}{2g} - \frac{gx^2}{2v_0^2}.$$

Cette parabole est appelée **parabole de sûreté**. C'est l'enveloppe des trajectoires obtenues en faisant varier l'angle de lancement pour une vitesse initiale de module donné. La figure 15.14 représente les trajectoires pour différents angles de tir avec une vitesse initiale de $10 \text{ m}\cdot\text{s}^{-1}$ ainsi que la parabole de sûreté en trait plus épais.

Pour un point situé sous la parabole de sûreté, on a deux solutions pour l'angle α , donc deux angles de tir possibles permettant d'atteindre le point considéré. L'un correspond au tir tendu, l'autre au tir en cloche.

Lorsque le point considéré appartient à la parabole de sûreté, le discriminant est nul. Dans ce cas, il est possible de l'atteindre mais il n'y aura qu'une seule valeur possible pour l'angle de tir, à savoir $\arctan\left(\frac{v_0^2}{g}\right)$.

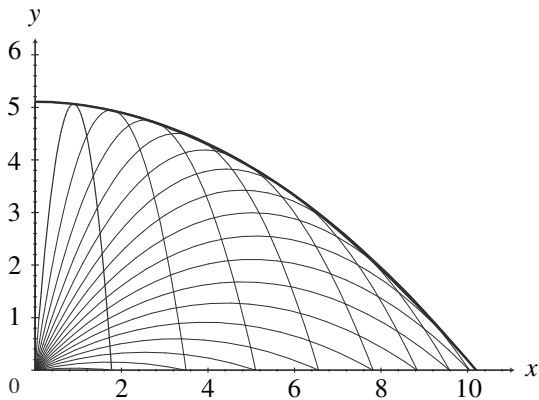


Figure 15.14 – Parabole de sûreté pour une vitesse initiale de 10 m·s⁻¹.

8.3 Tir en tenant compte de la résistance de l'air

Dans l'étude précédente, on n'a pas tenu compte de la résistance de l'air. L'air étant un gaz, on modélise sa résistance par une force de frottement proportionnelle au carré de la vitesse. L'équation du mouvement est la même que l'équation (15.4) établie au paragraphe 7.4 :

$$m \vec{a} = -k_2 v \vec{v} + m \vec{g}.$$

Lorsque la vitesse initiale n'est pas nulle, on ne peut résoudre cette équation différentielle que numériquement et pour cela, il est intéressant de la rendre adimensionnelle en introduisant des grandeurs caractéristiques.

a) Équation différentielle adimensionnelle

Il s'agit de la même équation différentielle que celle que l'on a étudiée au paragraphe 7.4. On introduit la même vitesse limite $v_l = \sqrt{\frac{mg}{k_2}}$ et le même temps caractéristique $\tau = \frac{v_l}{g}$. On pose $t = \tau t^*$ où t^* est un temps adimensionné et $\vec{v} = v_l \vec{v}^*$ où $\vec{v}^*$ est une vitesse adimensionnée. On déduit $\vec{a} = \frac{d\vec{v}}{dt} = \frac{d(v_l \vec{v}^*)}{d(\tau t^*)} = \frac{v_l}{\tau} \frac{d\vec{v}^*}{dt^*} = g \frac{d\vec{v}^*}{dt^*}$. On reporte ces expressions dans

l'équation différentielle précédente et on obtient : $\frac{d\vec{v}^*}{dt^*} = -v^* \vec{v}^* - \vec{u}_y.$

On introduit alors l'échelle de longueur $l = v_l \tau$ pour définir les longueurs adimensionnelles $x^* = \frac{x}{l}$ et $y^* = \frac{y}{l}$. Par projection sur les axes (Ox) et (Oy), on obtient les équations sans dimensions :

$$\begin{cases} \ddot{x}^* &= -x^* \sqrt{x^{*2} + y^{*2}} \\ \ddot{y}^* &= -y^* \sqrt{x^{*2} + y^{*2}} - 1, \end{cases}$$

qui ne peuvent être résolues que par des méthodes numériques.
Cette écriture montre que le type de comportement du mobile n'est déterminé que par sa

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

vitesse initiale adimensionnée $\vec{v}_0^* = \vec{v}_0/v_l$ dont on peut faire varier la norme et la direction. On va s'intéresser ici à l'influence de sa norme en fixant l'angle que fait $\vec{v}_0^*$ avec l'horizontale à 45° et traçant les trajectoires pour différentes valeurs de $\|\vec{v}_0^*\|$. On rappelle que v_l est la vitesse pour laquelle la norme des frottements est égale à la norme du poids. Ainsi,

- lorsque $v \gg v_l \Leftrightarrow v^* \gg 1$, le poids est négligeable ;
- lorsque $v \ll v_l \Leftrightarrow v^* \ll 1$, les frottements sont négligeables.

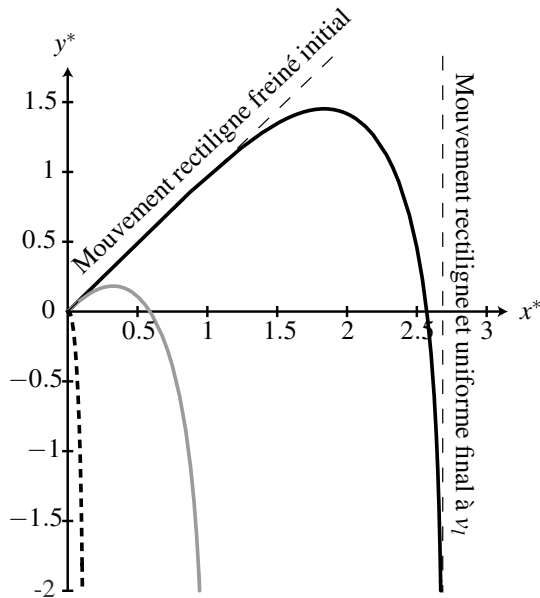


Figure 15.15 – Évolution de la trajectoire pour différentes vitesses initiales. L'angle de tir est fixé à 45° et la vitesse initiale prend les valeurs de $0, 1v_l$ (tirets), v_l (trait continu gris) et $10v_l$ (trait continu noir).

Selon la norme de $\vec{v}_0^*$, on distingue trois différents régimes pour le mouvement de M :

- lorsque $\|\vec{v}_0^*\| \ll 1$, la vitesse initiale est négligeable et le mouvement ressemble au mouvement rectiligne de chute libre traité au paragraphe 7.4 avec un faible décalage vers les x^* positifs ;
- lorsque $\|\vec{v}_0^*\|$ est de l'ordre de 1, les deux forces ont une importance équivalente et on ne peut pas simplifier la résolution : la trajectoire est courbe et amène le mobile sur une trajectoire verticale rectiligne et uniforme parcourue à la vitesse v_l ;
- lorsque $\|\vec{v}_0^*\| \gg 1$, le poids est négligeable dans un premier temps et le mouvement initial est rectiligne freiné. Puis la vitesse diminue et on est ramené au cas précédent : le mobile change de direction pour évoluer verticalement vers le bas. Enfin, la vitesse étant verticale, on est ramené au premier cas pour finir avec un mouvement vertical rectiligne et uniforme à la vitesse v_l .

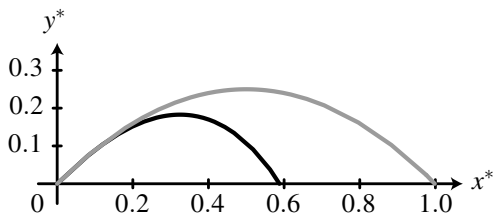


Figure 15.16 – Comparaison entre une trajectoire avec frottement pour laquelle $v_0 = v_l$ (en noir) et une trajectoire sans frottement avec la même vitesse initiale (en gris). Les frottements diminuent la portée sans changer fondamentalement la forme de la trajectoire. L'angle de tir est de 45° par rapport à l'horizontale.

Exemple

Lors d'un échange de badminton, le volant part de la raquette avec une vitesse de 200 à 300 km/h et on peut évaluer la vitesse limite à environ 15 à 30 km/h. La vitesse initiale est très grande devant la vitesse limite : la trajectoire ressemble à la trajectoire de type 3 avec une phase initiale rectiligne et un volant qui retombe presque à la verticale.

Lors d'un dégagement du gardien, un ballon de football part avec une vitesse de 100 à 150 km/h et on peut évaluer la vitesse limite à environ 80 à 100 km/h. La vitesse initiale et la vitesse limite sont comparables : la trajectoire ressemble à la trajectoire du type 2. On peut voir sur la figure 15.16 que la différence principale par rapport à une trajectoire parabolique obtenue sans frottement est que la portée diminue.

9 Le pendule simple

On considère un solide de petite taille et de masse m attaché à un fil. L'autre extrémité du fil est liée à un point fixe. À l'instant initial, on lâche le solide sans vitesse d'une position faisant un angle θ_0 avec la verticale. On observe que le mouvement ultérieur est plan.

9.1 Modélisation

On modélise le solide par un point matériel M de masse m et le fil par un fil inextensible, sans masse et sans rigidité (fil idéal). On étudie son mouvement dans le référentiel terrestre supposé galiléen. On néglige les frottements dus à l'air.

9.2 Équation du mouvement

On va suivre la procédure de résolution d'un problème de mécanique décrite précédemment :

- définition du système : le solide de masse m assimilé à un point matériel M ,
- choix du référentiel : le référentiel terrestre que l'on considère galiléen,
- bilan des forces : ne pas considérer la résistance de l'air revient à négliger tout phénomène de frottements. Le point matériel est soumis à :
 - son poids : $\vec{P} = m \vec{g}$;

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

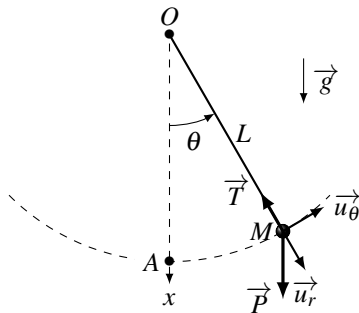


Figure 15.17 – Le pendule simple

- la tension du fil : $\vec{T} = -T\vec{u}_r$ où $T > 0$.
- choix de la méthode de résolution : on utilise le principe fondamental de la dynamique qui s'écrit :

$$m\vec{a} = \sum_i \vec{f}_i = m\vec{g} + \vec{T}.$$

Une étape importante consiste à réfléchir à la forme géométrique de la trajectoire du mobile pour choisir correctement le système de coordonnées. Ici, le mouvement étant plan et le fil inextensible, le point M se déplace sur une portion de cercle de centre O et de rayon L . Le repérage polaire s'impose et, en utilisant le cours de cinématique sur les mouvements circulaires, on établit l'expression de l'accélération en coordonnées polaires définis sur le schéma de la figure 15.17 :

$$\vec{OM} = L\vec{u}_r \quad ; \quad \vec{v} = L\dot{\theta}\vec{u}_\theta \quad ; \quad \vec{a} = -L\dot{\theta}^2\vec{u}_r + L\ddot{\theta}\vec{u}_\theta.$$

Par ailleurs, on projette tous les vecteurs sur la base polaire $(\vec{u}_r, \vec{u}_\theta)$ et en particulier la force $\vec{P}$ puisque que $\vec{T} = -T\vec{u}_r$ est déjà écrite sur cette base. Pour cela, on s'aide du schéma de la figure 15.17 et on trouve :

$$\vec{P} = mg \cos \theta \vec{u}_r - mg \sin \theta \vec{u}_\theta.$$

Il ne reste plus qu'à projeter le principe fondamental de la dynamique sur $\vec{u}_r$ et $\vec{u}_\theta$:

$$\begin{cases} -mL\dot{\theta}^2 &= mg \cos \theta - T \\ mL\ddot{\theta} &= -mg \sin \theta \end{cases}$$

La première équation permet d'établir l'expression de la tension du fil T qui est une inconnue de ce mouvement, à condition de connaître l'évolution temporelle de θ . La deuxième équation est l'équation du mouvement. Sa résolution donne l'expression de $\theta(t)$. En posant $\omega_0 = \sqrt{g/L}$, on peut l'écrire :

$$\ddot{\theta} + \omega_0^2 \sin \theta = 0. \tag{15.6}$$

Cette équation n'est pas linéaire et on ne peut pas la résoudre simplement. Deux approches sont possibles : la résolution numérique ou la linéarisation en vue de trouver une solution analytique.

9.3 Résolution numérique

Cette équation est entièrement déterminée par la connaissance de ω_0 . Les paramètres g , m et L n'interviennent pas indépendamment mais uniquement sous la forme de la combinaison $\omega_0 = \sqrt{g/L}$ et la masse m n'intervient pas dans le problème. Le seul paramètre libre est l'angle initial θ_0 dont on va étudier l'influence en traçant les solutions $\theta(t)$ en fonction du temps pour θ_0 variant entre 0,1 et 1,5 radians (oscillations de grande amplitude) ou entre 0,01 et 0,15 radians (oscillations de faible amplitude). Les résultats sont représentés sur la figure 15.18.

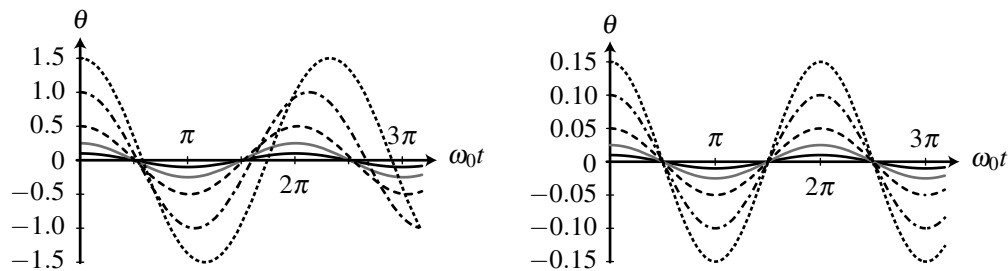


Figure 15.18 – Évolution temporelle de l'angle θ pour différentes conditions initiales. À gauche, l'amplitude des oscillations est comprise entre 0,1 et 1,5 rad ; à droite, elle est comprise entre 0,01 et 0,15 rad.

Dans tous les cas, on observe des oscillations périodiques. Pour les oscillations de grande amplitude, la période de ces oscillations dépend de θ_0 donc des conditions initiales. Lorsque l'amplitude des oscillations reste petite (inférieure à 0,4 rad environ), les oscillations ont une période commune. On parle d'**isochronisme des petites oscillations**. On va chercher à établir ce résultat en linéarisant l'équation différentielle (15.6) pour les petits angles.

9.4 Cas des oscillations de faibles amplitudes

a) Linéarisation et résolution analytique

Dans le cas où θ reste petit devant 1, on peut assimiler $\sin \theta$ à θ et linéariser l'équation (15.6). On obtient alors l'équation différentielle linéaire :

$$\ddot{\theta} + \omega_0^2 \theta = 0. \tag{15.7}$$

La variable angulaire θ vérifie alors l'équation différentielle d'un oscillateur harmonique que l'on a étudié au premier chapitre de cet ouvrage. Les solutions pour $\theta(t)$ sont de la forme

$$\theta = A \cos(\omega_0 t) + B \sin(\omega_0 t),$$

où A et B sont des constantes. Les conditions initiales $\theta(t = 0) = \theta_0$ et $\dot{\theta}(t = 0) = 0$ s'écrivent :

$$\theta_0 = A \quad \text{et} \quad 0 = B \omega_0,$$

soit $A = \theta_0$ et $B = 0$. Finalement : $\theta = \theta_0 \cos(\omega_0 t).$

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

$\theta(t)$ présente des oscillations sinusoïdales de pulsation ω_0 et de période $T_0 = \frac{2\pi}{\omega_0}$ indépendante de θ_0 . Dans la limite des oscillations de faible amplitude, on vient de montrer l'isochronisme des petites oscillations avec une période :

$$T_0 = 2\pi\sqrt{\frac{L}{g}}.$$

En pratique, l'isochronisme est assuré avec une précision de 1% lorsque θ reste inférieur à environ 0,4 rad soit environ 25°.

b) Portrait de phase

On peut intégrer l'équation (15.7) pour obtenir une intégrale première du mouvement et en déduire le portrait de phase. La méthode est détaillée dans l'annexe mathématique de ce livre. On commence par écrire l'équation (15.7) sous la forme

$$\ddot{\theta}(t) = -\omega_0^2 \theta(t),$$

et on multiplie cette égalité par $\dot{\theta}$:

$$\ddot{\theta}\dot{\theta} = -\omega_0^2 \theta(t)\dot{\theta} \iff \frac{d(\dot{\theta}^2/2)}{dt} = -\omega_0^2 \frac{d(\theta^2/2)}{dt} \iff \frac{d(\dot{\theta}^2 + \omega_0^2 \theta^2)}{dt} = 0.$$

On déduit alors que : $\dot{\theta}^2 + \omega_0^2 \theta^2 = C$, où C est une constante d'intégration.

Les conditions initiales $\theta(0) = \theta_0$ et $\dot{\theta}(0) = 0$ permettent de déterminer C :

$$\dot{\theta}^2 + \omega_0^2 \theta^2 = \omega_0^2 \theta_0^2,$$

que l'on peut également écrire : $\left(\frac{\dot{\theta}}{\omega_0 \theta_0}\right)^2 + \left(\frac{\theta}{\theta_0}\right)^2 = 1.$

On reconnaît l'équation d'un cercle de centre O et de rayon 1 dans le plan de phase adimensionné $\left(\frac{\theta}{\theta_0}, \frac{\dot{\theta}}{\omega_0 \theta_0}\right)$ que l'on trace sur la figure 15.19.

Remarque

On appelle **intégrale première du mouvement** une grandeur conservée au cours du mouvement obtenue en intégrant une fois le principe fondamental de la dynamique. Ici,

la quantité $\left(\frac{\dot{\theta}}{\omega_0 \theta_0}\right)^2 + \left(\frac{\theta}{\theta_0}\right)^2 = 1$ est une intégrale première du mouvement.

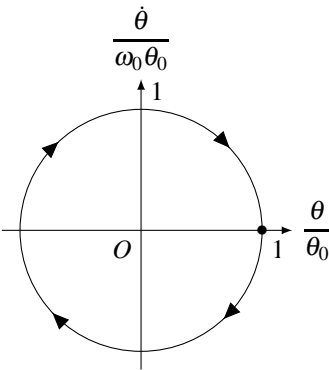


Figure 15.19 – Diagramme de phase du pendule simple dans la limite des petits angles.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

SYNTHÈSE

SAVOIRS

- définition d'une force
- principe des actions réciproques
- principe de l'inertie
- définition d'un référentiel galiléen
- définition d'une quantité de mouvement
- principe fondamental de la dynamique
- expression des forces usuelles
- définition d'un portrait de phase

SAVOIR-FAIRE

- définir un système mécanique et un référentiel d'étude
- établir un schéma clair du problème
- établir un bilan des forces qui s'exercent sur un système
- choisir un système de coordonnées et une base de projection adaptés au problème
- projeter des équations vectorielles
- reconnaître une équation différentielle
- établir et intégrer l'équation différentielle du mouvement dans le cas du mouvement dans le champ de pesanteur terrestre en négligeant les frottements puis en tenant compte des frottements
- analyser des résultats d'intégration numérique présentés sous la forme de courbes
- établir et linéariser l'équation différentielle du pendule simple
- déterminer l'intégrale première de l'équation différentielle précédente pour établir son portrait de phase
- lire et exploiter un portrait de phase

MOTS-CLÉS

- | | | |
|-------------------------|---------------------------|-----------------------|
| • masse | • conditions initiales | • pendule simple |
| • force | • équation différentielle | • frottements fluides |
| • quantité de mouvement | • portrait de phase | |
| • équation du mouvement | • chute libre | |

S'ENTRAÎNER

15.1 Bond sur la Lune (★)

Dans l'album de Tintin, « On a marché sur la Lune », le Capitaine Haddock est étonné de pouvoir faire un bond beaucoup plus grand que sur terre : c'est le sujet de cet exercice.

On assimile le mouvement du Capitaine Haddock à celui de son centre de gravité M de masse m . Il saute depuis le sol lunaire avec une vitesse initiale v_0 faisant un angle $\alpha = 30^\circ$ avec le sol. On note g_l l'accélération de la pesanteur à la surface de la lune. En l'absence d'atmosphère, on peut considérer qu'il n'y a aucune force de frottement.

1. Pour ce premier exercice, on détaille la mise en équation du problème. Pour les exercices suivants, les quatre questions qui suivent sont sous-entendues. Il est malgré tout nécessaire de les traiter systématiquement.

- a. Définir le système et le référentiel dans lequel on étudie son mouvement.
- b. Etablir un bilan des forces détaillé et réaliser un schéma à un instant quelconque en y faisant figurer les différentes forces.
- c. Définir le repère adapté à l'étude de ce mouvement et rappeler l'expression des vecteurs position, vitesse et accélération dans ce type de repérage.
- d. En utilisant les questions précédentes, établir les équations du mouvement.

2. Déterminer l'expression de la distance horizontale parcourue au cours du saut en fonction de v_0 , α et g_l .

3. Sur la Lune, la pesanteur est environ six fois moins forte que sur la Terre. Quelle sera la distance horizontale parcourue par le sauteur sur la Lune, si elle est de $d = 1,50$ m sur la Terre ?

15.2 Oscillations d'une masse suspendue à un ressort (★)

On s'intéresse au mouvement d'un petit objet de masse m attaché à un ressort dont l'autre extrémité est accrochée à un bâti fixé au sol. Le ressort étant initialement dans sa situation de repos pour laquelle sa longueur est égale à sa longueur à vide, on lâche l'objet sans lui donner de vitesse initiale. Le mouvement qui suit est vertical et on veut l'étudier.

On idéalise le comportement du ressort en l'assimilant à un ressort parfaitement élastique, sans masse, de raideur k et de longueur à vide l_0 . On repère la position de l'objet sur un axe (Ox) vertical descendant par son abscisse x (figure 15.20). L'origine O du repère est située à l'extrémité fixe du ressort. On néglige les frottements dus à l'air.

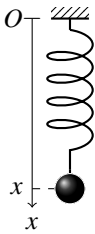


Figure 15.20

1. Montrer que l'équation du mouvement s'écrit :

$$\ddot{x} + \omega_0^2 x = \omega_0^2 x_{eq},$$

où ω_0 et x_{eq} sont des constantes à déterminer en fonction de l_0 , g , m et k .

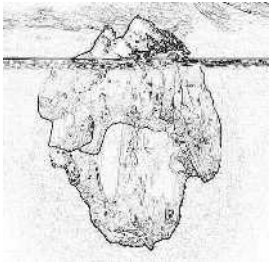
2. Que représente la position M_{eq} d'abscisse $x = x_{eq}$?

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

3. Pour résoudre l'équation du mouvement, on déplace l'origine du repère en M_{eq} .
- Montrer que cela revient à faire le changement d'inconnue $X = x - x_{eq}$.
 - En déduire que l'on obtient alors une équation différentielle du mouvement connue.
 - Résoudre en tenant compte des conditions initiales.
 - Représenter l'évolution temporelle de l'abscisse x .

15.3 Partie immergée de l'iceberg (★)

On considère un iceberg dont on peut voir un dessin sur la figure ci-contre. La ligne horizontale représente la surface de l'eau. On note V le volume total de l'iceberg, V_I son volume immergé, $\rho_G = 0,92 \cdot 10^3 \text{ kg} \cdot \text{m}^{-3}$ la masse volumique de la glace et $\rho_L = 1,02 \cdot 10^3 \text{ kg} \cdot \text{m}^{-3}$ celle de l'eau salée.



- Etablir les expressions de la poussée d'Archimède et de la force de pesanteur qui s'applique sur l'iceberg.
- Déterminer la proportion volumique de glace immergée.

15.4 Charge soulevée par une grue (★)

Une grue de chantier de hauteur h doit déplacer d'un point à un autre du chantier une charge de masse m assimilée à son centre de gravité M . Le point d'attache du câble sur le chariot de la grue est noté A .

- Le point A est à la verticale de M posé sur le sol. Déterminer la tension à appliquer au câble pour qu'il arrache très doucement le point M du sol.
- L'enrouleur de câble de la grue le remonte avec une accélération verticale a_v constante. Déterminer la tension du câble. Conclusion.

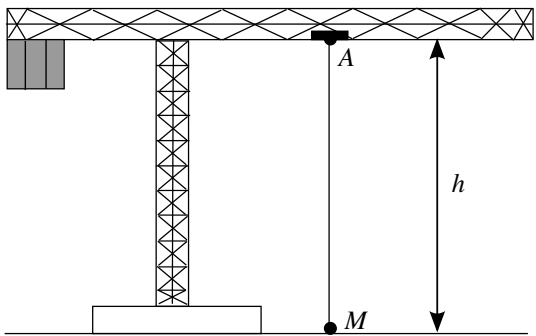


Figure 15.21

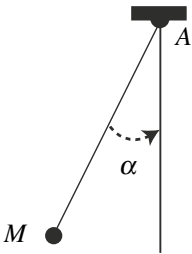


Figure 15.22

3. La montée de M est stoppée à mi-hauteur mais le chariot A se met en mouvement vers la droite (figure 15.21) avec une accélération horizontale a_h constante.
- Quelle est l'accélération de M sachant que M est alors immobile par rapport à A ?
 - Déterminer l'angle que fait le câble avec la verticale en fonction de m , g , a_h ainsi que la tension du câble.

15.5 Chaussette séchant dans un sèche-linge (★)

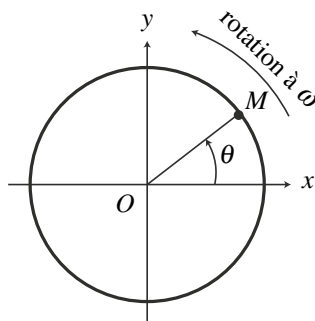
Dans le tambour d'un sèche-linge, on observe que le mouvement d'une chaussette s'effectue en une alternance de deux phases :

- dans une première phase, elle est entraînée par le tambour dans un mouvement de rotation uniforme ;
- dans une deuxième phase, elle retombe en chute libre.

L'observation montre qu'à chaque tour, elle décolle du tambour au même endroit. On cherche à déterminer ce lieu.

On modélise le tambour par un cylindre de rayon $R = 25 \text{ cm}$ tournant à $50 \text{ tour} \cdot \text{min}^{-1}$. On s'intéresse au mouvement de la chaussette que l'on assimile à un point matériel M de masse m . On étudie la première phase pendant laquelle le linge est entraîné dans un mouvement de rotation circulaire et uniforme à la même vitesse que le tambour et en restant collé aux parois du tambour. Pour les applications numériques, on considère $g = 9,8 \text{ m} \cdot \text{s}^{-2}$.

1. Déterminer l'accélération de la chaussette.
2. En déduire la réaction du tambour sur la chaussette.
3. Montrer que la réaction normale s'annule lorsque la chaussette atteint un point dont on déterminera la position angulaire.
4. Que se passe-t-il en ce point ? Quel est le mouvement ultérieur ?

**APPROFONDIR****15.6 Quelques tirs de basket-ball (★★)**

On étudie les tirs de basket-ball de manière simplifiée. On suppose que le joueur est face au panneau à une distance D de ce dernier. Le cercle du panier est situé à une hauteur $H = 3,05 \text{ m}$ au-dessus du sol et on assimilera dans un premier temps le cercle à un point situé sur le panneau. De même, le ballon sera considéré comme ponctuel. On néglige les frottements fluides de l'air. Le joueur tire d'une hauteur $h = 2,00 \text{ m}$ au-dessus du sol en imposant une vitesse initiale $\vec{v}_0$ faisant un angle α avec l'horizontale.

1. Etablir l'équation du mouvement du ballon lors du tir.
2. En déduire les équations horaires de ce mouvement.
3. Déterminer l'équation de la trajectoire du ballon.
4. On suppose que le module de la vitesse initiale est fixé. Donner l'équation à vérifier par l'angle α pour que le panier soit marqué. On la mettra sous la forme d'une équation du second degré en $\tan \alpha$.
5. Montrer que cette équation n'admet des solutions que si le module v_0 de la vitesse initiale vérifie une inéquation du second degré en v_0^2 .

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

6. En déduire l'existence d'une valeur minimale de v_0 pour que le panier soit marqué.
7. Faire l'application numérique pour un lancer franc (la distance D vaut alors 4,60 m) puis pour un panier à trois points (la distance D vaut alors 6,25 m selon les règles de la Fédération Internationale de Basket-Ball).
8. Si la condition précédente est vérifiée, donner l'expression de $\tan \alpha$ et en déduire qu'il existe deux angles possibles pour marquer le panier.
9. Donner les valeurs numériques des angles α permettant de marquer un lancer franc en supposant que $v_0 = 10,0 \text{ m.s}^{-1}$.
10. Dans la suite, on suppose que l'angle de tir est fixé. Déterminer l'expression de la vitesse initiale v_0 à imposer pour marquer le panier.
11. Faire l'application numérique pour un lancer franc et un angle de tir $\alpha = 70^\circ$.
12. On s'intéresse maintenant à l'influence de la dimension du ballon et du cercle du panier. Le rayon du ballon est de 17,8 cm et le diamètre du cercle mesure 45,0 cm. On note x la distance du centre du ballon au panneau à hauteur du cercle du panier. Donner l'intervalle des valeurs de x permettant de marquer le panier.
13. Quelle équation doit vérifier x ? On pourra la laisser sous la forme d'une équation du second degré en $(D + x)$.
14. En déduire la condition que doivent vérifier v_0 , α , g , H et h pour que cette équation ait des solutions.
15. En déduire les valeurs d'angle de tir possibles en supposant la vitesse initiale v_0 fixée. On fera l'application numérique pour $v_0 = 10,0 \text{ m.s}^{-1}$.
16. Etablir l'existence d'une valeur minimale de v_0 à angle de tir fixé. On fera l'application numérique pour $\alpha = 70^\circ$.
17. Cette condition étant vérifiée, donner l'expression de x .

15.7 Toboggan (★★)

Une personne participant à un jeu télévisé doit laisser glisser un paquet sur un toboggan, à partir du point A de manière à ce qu'il soit réceptionné, au point B par un chariot se déplaçant le long de l'axe (Ox) . On néglige les frottements de l'air. Le toboggan exerce sur le paquet non seulement une réaction normale $\vec{R}_N$ mais aussi une réaction tangentielle $\vec{R}_t$ (frottement solide) telle que $\|\vec{R}_t\| = f\|\vec{R}_N\|$ avec $f = 0,5$.

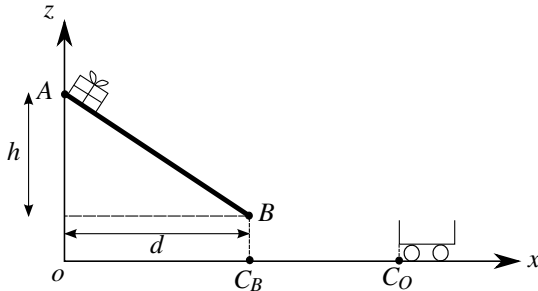


Figure 15.23

Le chariot part du point C_O à l'instant $t = 0$, vers la gauche (figure 15.23) avec une vitesse $v_c = 0,5 \text{ m}\cdot\text{s}^{-1}$. Arrivé en C_B , il s'arrête pendant un intervalle de temps $\delta t = 1 \text{ s}$.

La hauteur du toboggan est $h = 4 \text{ m}$, la distance d est égale à h et $C_O C_B = 2 \text{ m}$. La masse du paquet est $m = 10 \text{ kg}$ et $g = 9,8 \text{ m}\cdot\text{s}^{-2}$.

1. On pose $X = AP$ où P est la position du paquet à l'instant t . Déterminer $X(t)$. En déduire le temps mis par le paquet pour arriver en B .

2. A quel instant le joueur doit-il lâcher le paquet, sans vitesse initiale, pour qu'il tombe dans le chariot ?

15.8 Elastique de saut (★★)

Un fabricant d'élastiques de saut donne les caractéristiques suivantes pour un élastique de type **M** adapté à un sauteur dont le poids est compris entre 65 et 95 kg.

- Longueur à vide (notée ℓ_0) : de 6 m à 50 m pour des sites de saut de 30 m à 250 m.
- Tension appliquée sur un élastique pour un allongement de 100% : 200 kg.
- Tension appliquée sur un élastique pour un allongement de 200% : 325 kg.

On prendra $g = 9,8 \text{ m}\cdot\text{s}^{-2}$.

1. Traduire la fiche technique en langage correct en physique.

a. La tension de l'élastique obéit-elle à une loi type ressort : $\overline{T} = k(\ell - \ell_0)$ où T est la tension appliquée et ℓ la longueur de l'élastique ?

b. On supposera dans la suite que la tension est de la forme $\overline{T} = k(\ell - \ell_0)$. A partir des données, déterminer une valeur moyenne de la constante de raideur de l'élastique k pour un élastique de $\ell_0 = 6 \text{ m}$.

2. On veut savoir à quelle hauteur remonterait une masse test M de masse minimale (ici $m = 65 \text{ kg}$) si elle est lâchée sans vitesse initiale, l'élastique tendu au maximum. On prendra l'exemple d'un élastique de $\ell_0 = 6 \text{ m}$, dont le point d'attache est à la hauteur $h = 18 \text{ m}$ au-dessus du point de départ.

a. Déterminer l'expression de l'altitude $z(t)$ et de la vitesse $\dot{z}(t)$ tant que l'élastique est tendu.

b. Calculer l'instant t_1 pour lequel l'élastique n'est plus tendu. En déduire la vitesse de M en cet instant.

c. On ne tient pas compte des frottements de l'air. Déterminer la hauteur maximale atteinte par l'objet. Conclusion.

3. On réalise maintenant un saut normal, à partir du point d'attache, sans vitesse initiale avec les mêmes caractéristiques qu'à la question précédente.

a. Déterminer le point le plus bas atteint par le sauteur et l'accélération ressentie en ce point. Conclusion.

b. Après plusieurs oscillations, le sauteur s'immobilise. à quelle hauteur ?

15.9 Jeux aquatiques (★★★)

Un baigneur (masse $m = 80 \text{ kg}$) saute d'un plongoir situé à une hauteur $h = 10 \text{ m}$ au-dessus de la surface de l'eau. On considère qu'il se laisse chuter sans vitesse initiale et qu'il est

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

uniquement soumis à la force de pesanteur (on prendra $g = 10 \text{ m}\cdot\text{s}^{-2}$) durant la chute. On note (Oz) , l'axe vertical descendant, O étant le point de saut.

1. Déterminer la vitesse v_e d'entrée dans l'eau ainsi que le temps de chute t_c . Application numérique.

2. Lorsqu'il est dans l'eau, le baigneur ne fait aucun mouvement. Il subit en plus de la pesanteur :

- une force de frottement $\vec{f}_f = -k\vec{v}$ ($\vec{v}$ étant la vitesse et $k = 250 \text{ kg}\cdot\text{s}^{-1}$);
- la poussée d'Archimède $\vec{\Pi} = -\frac{m}{d_h}\vec{g}$ ($d_h = 0,9$ est la densité du corps humain).

a. Etablir l'équation différentielle à laquelle obéit la vitesse en projection sur (Oz) , notée v_z . On posera $\tau = m/k$.

b. Intégrer cette équation en prenant comme nouvelle origine des temps $t = t_c$.

c. Déterminer la vitesse limite v_L (< 0) en fonction de m, k, g et d_h . Application numérique.

d. Exprimer la vitesse v_z en fonction de $v_e, |v_L|$ et t . Déterminer à quel instant t_1 le baigneur commence à remonter.

e. En prenant la surface de l'eau comme nouvelle origine de l'axe (Oz) , exprimer $z(t)$. En déduire la profondeur maximale pouvant être atteinte.

f. En fait, il suffit que le baigneur arrive au fond de la piscine avec une vitesse de l'ordre de $1 \text{ m}\cdot\text{s}^{-1}$ pour qu'il puisse se repousser avec ses pieds sans risque : à quel instant t_2 atteint-il cette vitesse et quelle est la profondeur minimale du bassin ?

3. Le même baigneur décide maintenant d'effectuer un plongeon. On suppose qu'il entre dans l'eau avec un angle $\alpha = 60^\circ$ par rapport à l'horizontale et une vitesse $v_0 = 8 \text{ m}\cdot\text{s}^{-1}$. Les forces qui s'exercent sur lui sont les mêmes que précédemment mais le coefficient k est divisé par deux en raison d'une meilleure pénétration dans l'eau.

On repère le mouvement par les axes (Ox) (axe horizontal de même sens que $\vec{v}_0$) et (Oz) (vertical descendant comme précédemment); le point O est le point de pénétration dans l'eau.

a. Déterminer les projections des équations du mouvement sur Ox et Oz .

b. En déduire les composantes de la vitesse dans l'eau en fonction du temps. Existe-t-il une vitesse limite ?

c. Le plongeur peut-il atteindre le fond de la piscine situé à 4 m ?

15.10 Etude d'un skieur d'après ESIGETEL 2000 (★)

On étudie le mouvement d'un skieur descendant une piste selon la ligne de plus grande pente, faisant l'angle α avec l'horizontale. L'air exerce une force de frottement supposée de la forme $\vec{F} = -\lambda\vec{v}$, où λ est un coefficient constant positif et $\vec{v}$ la vitesse du skieur.

On note $\vec{T}$ et $\vec{N}$ les composantes tangentielle et normale de la force de contact exercée par la neige et f le coefficient de frottement solide tel que $\|\vec{T}\| = f\|\vec{N}\|$.

On choisit comme origine de l'axe Ox de la ligne de plus grande pente la position initiale du skieur, supposé partir à l'instant initial avec une vitesse négligeable. On note Oy la normale à la piste dirigée vers le haut.

1. Calculer $\vec{T}$ et $\vec{N}$.

2. Calculer la vitesse $\vec{v}$ et la position x du skieur à chaque instant.
3. a. Montrer que le skieur atteint une vitesse limite $\vec{v}_l$ et calculer $\vec{v}$ en fonction de $\vec{v}_l$.
b. Application numérique : calculer v_l avec $\lambda = 1 \text{ S}\cdot\text{I}$, $f = 0,9$, $g = 10 \text{ m}\cdot\text{s}^{-2}$, $m = 80 \text{ kg}$ et $\alpha = 45^\circ$.
4. Calculer littéralement et numériquement la date t_1 où le skieur a une vitesse égale à $v_l/2$.
5. A la date t_1 , le skieur tombe. On néglige alors la résistance de l'air, et on considère que le coefficient de frottement sur le sol est multiplié par 10. Calculer la distance parcourue par le skieur avant de s'arrêter.

15.11 Atterrissage en catastrophe d'un avion d'après Agrégation de Chimie 2006 (★★)

Un avion de chasse de masse $m = 9 \text{ t}$ en panne de freins atterrit à une vitesse $v_a = 241 \text{ km}\cdot\text{h}^{-1}$. Une fois au sol, il est freiné en secours par un parachute de diamètre $D = 3 \text{ m}$ déployé instantanément par le pilote au moment où les roues de l'avion touchent le sol. On néglige les forces de frottement des roues sur le sol par rapport à la force de traînée (frottement fluide) du parachute qui s'écrit $T = C_x \rho S v^2 / 2$ avec v la vitesse de l'avion, S la surface projetée du parachute sur un plan perpendiculaire à la vitesse, C_x le coefficient de traînée supposé constant et égal à 1,5, $\rho = 1,2 \text{ kg}\cdot\text{m}^{-3}$ la masse volumique de l'air.

1. On considère que le réacteur ne délivre plus aucune poussée.
 - a. Etablir l'équation différentielle à laquelle obéit la vitesse v de l'avion.
 - b. En déduire l'expression de v en fonction de la date t . On prendra comme origine des temps la date à laquelle les roues de l'avion touchent le sol.
 - c. Montrer que la force de traînée n'est pas suffisante pour stopper complètement l'avion.
2. Dans le cas où les freins fonctionnent, la distance d'atterrissage de l'avion est de l'ordre de $d = 1400 \text{ m}$. Déterminer la vitesse atteinte par l'avion après avoir été freiné uniquement par le parachute sur cette distance. Faites l'application numérique.

CORRIGÉS

15.1 Bond sur la Lune

1. Mise en équation

a. On étudie le mouvement du centre de gravité M du Capitaine Haddock dans le référentiel lié à la Lune galiléen.

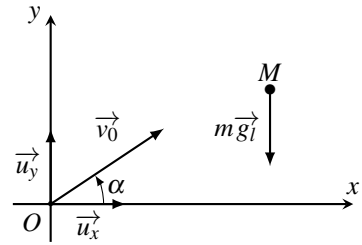
b. Après l'impulsion sur le sol qui permet d'acquérir cette vitesse initiale, la seule force est la force de pesanteur lunaire : $m\vec{g}_l = -mg_l\vec{u}_z$.

c. Le repère est (O, x, y, z) , le point O est l'origine du saut et Oz est vertical ascendant. On choisit les axes tels que :

$$\vec{v}_0 = v_0(\cos \alpha \vec{u}_x + \sin \alpha \vec{u}_z).$$

Dans ce repère cartésien, les vecteurs position, vitesse et accélération sont respectivement :

$$\overrightarrow{OM} \begin{cases} x \\ y \\ z \end{cases} \quad \vec{v} \begin{cases} \dot{x} \\ \dot{y} \\ \dot{z} \end{cases} \quad \vec{a} \begin{cases} \ddot{x} \\ \ddot{y} \\ \ddot{z} \end{cases}$$



d. La relation fondamentale de la dynamique donne : $m\vec{a} = m\vec{g}_l$ soit :

$$\ddot{x} = 0 \quad ; \quad \ddot{y} = 0 \quad ; \quad \ddot{z} = -g_l.$$

2. Résolution : par intégration successive, on trouve :

$$\begin{cases} \dot{x} = v_0 \cos \alpha \\ \dot{y} = 0 \\ \dot{z} = -g_l t + v_0 \sin \alpha \end{cases} \Rightarrow \begin{cases} x = v_0 t \cos \alpha \\ y = 0 \\ z = -\frac{1}{2} g_l t^2 + v_0 t \sin \alpha. \end{cases}$$

On obtient l'équation de la trajectoire en éliminant t :

$$z = -\frac{1}{2} g_l \frac{x^2}{(v_0 \cos \alpha)^2} + x \tan \alpha.$$

Le sauteur retombe sur le sol pour $z = 0$, soit en rejetant la solution $x = 0$ qui correspond au départ, on trouve la distance parcourue :

$$d = \frac{v_0^2}{g_l} \sin 2\alpha.$$

3. En supposant que v_0 et α soient les mêmes et puisque d est inversement proportionnel à g_l , la distance sur la Lune sera six fois plus grande soit 9 m. Il est vraisemblable que v_0 soit plus grand sur la Lune que sur Terre, donc que d soit encore plus grand sur la Lune.

15.2 Oscillations d'une masse suspendue à un ressort

1. On étudie le mouvement du petit objet de masse m assimilé à un point matériel M dans le référentiel lié au sol galiléen. Le mouvement étant vertical, on définit le repère cartésien $(Oxyz)$ tel que l'axe (Ox) coïncide avec la verticale descendante.

M est repéré par son vecteur position $\overrightarrow{OM} = x\overrightarrow{u_x}$. Sa vitesse et son accélération valent respectivement :

$$\overrightarrow{v} = \dot{x}\overrightarrow{u_x} \quad \text{et} \quad \overrightarrow{a} = \ddot{x}\overrightarrow{u_x}.$$

Dans ce référentiel, le point M est soumis à son poids $\overrightarrow{P} = m\overrightarrow{g} = mg\overrightarrow{u_x}$ et à la force rappel élastique du ressort $\overrightarrow{T} = -k(l - l_0)\overrightarrow{u_{ext}}$. Avec le choix de repère précédent, $l = x$ et $\overrightarrow{u_{ext}} = \overrightarrow{u_x}$. On applique alors le principe fondamental de la dynamique : $m\overrightarrow{a} = \overrightarrow{P} + \overrightarrow{T}$ que l'on projette sur la direction du mouvement (Ox) :

$$m\ddot{x} = mg - k(x - l_0) \quad \Leftrightarrow \quad \ddot{x} + \frac{k}{m}x = g + \frac{k}{m}l_0 = \frac{k}{m}\left(l_0 + \frac{mg}{k}\right).$$

Cette équation est bien de la forme $\ddot{x} + \omega_0^2x = \omega_0^2x_{eq}$ avec $\omega_0 = \sqrt{\frac{k}{m}}$ et $x_{eq} = l_0 + \frac{mg}{k}$.

2. Lorsque M est situé à la position M_{eq} , son accélération, et par conséquent la somme des forces auxquelles il est soumis, s'annule. Cette position est la position d'équilibre du point M . Comme on peut s'y attendre, elle correspond à un ressort étiré ($l_{eq} = l_0 + \frac{mg}{k} > l_0$) de telle sorte que la tension du ressort compense le poids.

3. Résolution :

a. Pour effectuer le changement d'origine du repère, on note $\overrightarrow{M_{eq}M} = X\overrightarrow{u_x}$ et on utilise la relation de Chasles :

$$\overrightarrow{M_{eq}M} = \overrightarrow{OM} - \overrightarrow{OM_{eq}} \quad \xrightarrow{\text{sur } \overrightarrow{u_x}} \quad X = x - x_{eq}.$$

b. x_{eq} étant une constante, on dérive l'expression précédente et on trouve :

$$\dot{X} = \dot{x} \quad \text{et} \quad \ddot{X} = \ddot{x}.$$

On remplace alors dans l'équation du mouvement pour obtenir $\ddot{X} + \omega_0^2X = 0$. Il s'agit de l'équation d'un oscillateur harmonique.

c. Les solutions sont : $X = a\cos\omega_0t + b\sin\omega_0t$ et les conditions initiales données dans l'énoncé se traduisent par $x(t=0) = l_0$ et $\dot{x}(t=0) = 0$ soit :

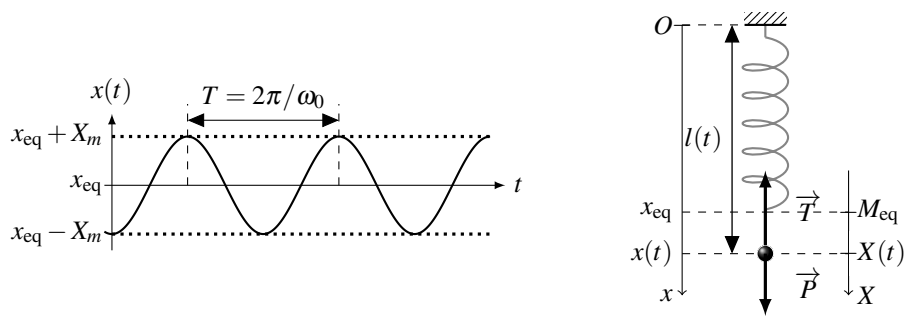
$$\begin{cases} X(t=0) = l_0 - x_{eq} = -\frac{mg}{k} = a \\ \dot{X}(t=0) = 0 = b\omega_0. \end{cases}$$

Finalement :

$$X = -\frac{mg}{k}\cos\omega_0t \quad \text{et} \quad x = x_{eq} - \frac{mg}{k}\cos\omega_0t.$$

d. Le mouvement de M présente des oscillations harmoniques centrées sur la position d'équilibre, de pulsation ω_0 , d'amplitude $X_m = \frac{mg}{k}$ et de phase à l'origine π .

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE



15.3 Partie immergée de l'iceberg

1. Dans le référentiel terrestre $\mathcal{R}_T$ galiléen, l'iceberg est soumis à son poids $\vec{P} = m\vec{g} = \rho_G V \vec{g}$ et à la poussée d'Archimède $\vec{\Pi} = -\rho_L V_I \vec{g}$.
2. On applique le principe fondamental de la dynamique à l'iceberg au repos dans $\mathcal{R}_T$:

$$\vec{0} = \vec{P} + \vec{\Pi} \Rightarrow \rho_G V = \rho_L V_I.$$

On en déduit :

$$\frac{V_I}{V} = \frac{\rho_G}{\rho_L} = \frac{0,92}{1,02} = 0,90.$$

90% du volume de l'iceberg est immergé et donc invisible.

15.4 Charge soulevée par une grue

Le référentiel lié au sol est galiléen. Le système est la masse m assimilée à son centre de gravité M repéré par ses coordonnées cartésiennes $M(x, y, z)$ où Oz est dirigé vers le haut.

1. Lorsque le point M est posé sur le sol, les forces qui s'exercent sur lui sont : le poids ($m\vec{g} = -mg\vec{u}_z$), la tension du câble ($\vec{T} = T\vec{u}_z$) et la réaction du sol ($\vec{R} = R\vec{u}_z$). La rupture de contact avec le sol implique l'annulation de la force de contact d'où $R = 0$. L'accélération de M est alors dirigée vers le haut, c'est-à-dire $\ddot{z} \geq 0$. Comme la grue arrache très doucement M du sol, on peut faire l'approximation $\ddot{z} \simeq 0$. On applique le principe fondamental de la dynamique à M de masse m à l'instant où le contact avec le sol est rompu et où R s'annule. On le projette sur l'axe Oz et on trouve :

$$T - mg = m\ddot{z} \Rightarrow T = m\ddot{z} + mg \simeq mg.$$

La tension appliquée au câble est, au minimum, égale au poids de l'objet à soulever.

2. On applique la relation fondamentale de la dynamique avec $\vec{a} = a_v\vec{u}_z$, soit $m\vec{a} = m\vec{g} + \vec{T}$, d'où $\vec{T} = m(\vec{a} - \vec{g})$, soit en projection sur $\vec{u}_z$:

$$T = m(a_v + g).$$

La tension est supérieure au poids et fonction affine de a_v . Si l'accélération devient trop forte, le câble peut se rompre.

3. On note O un point fixe dans le référentiel.

a. L'accélération de M est $\vec{a}_M = \frac{d\vec{OM}}{dt}$. On utilise la relation de Chasles $\vec{OM} = \vec{OA} + \vec{AM}$, or $\vec{AM}$ est un vecteur constant, donc :

$$\vec{a}_M = \frac{d\vec{OM}}{dt} = \frac{d\vec{OA}}{dt} + \frac{d\vec{AM}}{dt} = \frac{d\vec{OA}}{dt} = \vec{a}_h$$

b. La relation fondamentale de la dynamique appliquée à M donne $m\vec{a}_h = m\vec{g} + \vec{T}$, les différents vecteurs étant représentés sur la figure (15.24). La construction vectorielle donne alors $\tan \alpha = mg/a_h$ et $T = \sqrt{(mg)^2 + a_h^2}$.

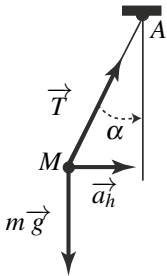


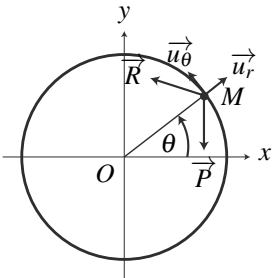
Figure 15.24

15.5 Chaussette séchant dans un sèche-linge

1. On étudie le mouvement de la chaussette assimilée à un point matériel M de masse m dans le référentiel terrestre galiléen. Lors de la première phase, M est en mouvement circulaire et uniforme de centre O et de rayon R à la vitesse angulaire ω .

Sa vitesse est tangente au cercle et de norme $v = R\omega$ et son accélération est radiale centripète de norme $\frac{v^2}{R}$ d'où :

$$\vec{a} = -\frac{v^2}{R}\vec{u}_r = -R\omega^2\vec{u}_r \quad \text{où} \quad \vec{u}_r = \frac{\vec{OM}}{\|\vec{OM}\|}.$$



2. La chaussette étant collée à la paroi du tambour, elle est soumise à son poids $\vec{P} = m\vec{g}$ et à la réaction du tambour $\vec{R}$. On lui applique le principe fondamental de la dynamique :

$$m\vec{a} = \vec{P} + \vec{R} \quad \Rightarrow \quad \vec{R} = m\vec{a} - \vec{P} = \begin{cases} \vec{R} \cdot \vec{u}_r &= -mR\omega^2 + mg \sin \theta \\ \vec{R} \cdot \vec{u}_\theta &= mg \cos \theta. \end{cases}$$

3. La composante radiale de la réaction du support s'annule lorsque $mR\omega^2 = mg \sin \theta$ soit pour $\theta = \theta_0$ tel que :

$$\sin \theta_0 = \frac{R}{g}\omega^2 = \frac{0,25}{9,8} \times \left(\frac{50 \times 2 \times 3,1415}{60} \right)^2 = 0,70 \quad \Rightarrow \quad \theta_0 = 44,4^\circ.$$

4. L'annulation de la force de contact entre le tambour et la chaussette montre qu'il y a rupture de contact. La chaussette décolle de la paroi et est alors projetée avec une vitesse $R\omega$ tangente au tambour depuis le point repéré par la coordonnée angulaire θ_0 . Le mouvement ultérieur, sous la seule action du poids, est un vol parabolique.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONNIENNE

15.6 Quelques tirs de basket

1. Le système étudié est le ballon assimilé à son centre de gravité dans le référentiel terrestre galiléen. Il est soumis à la seule force du poids. L'écriture du principe fondamental de la dynamique donne $\vec{a} = \vec{g}$.
2. La projection de la relation précédente en coordonnées cartésiennes en choisissant Oz dirigé vers le haut donne :

$$\begin{cases} \ddot{x} = 0 \\ \ddot{y} = 0 \\ \ddot{z} = -g \end{cases} \Rightarrow \begin{cases} \dot{x} = v_0 \cos \alpha \\ \dot{y} = 0 \\ \dot{z} = -gt + v_0 \sin \alpha \end{cases} \Rightarrow \begin{cases} x = v_0 t \cos \alpha - D \\ y = 0 \\ z = \frac{1}{2}gt^2 + v_0 t \sin \alpha + h \end{cases}$$

en prenant l'origine du repère au sol sous le panneau et en tenant compte des conditions initiales pour effectuer les intégrations.

3. Pour obtenir l'équation de la trajectoire, il suffit d'éliminer le temps entre les différentes équations. On en déduit $t = \frac{x+D}{v_0 \cos \alpha}$ donc :

$$z = \frac{1}{2}g \left(\frac{x+D}{v_0 \cos \alpha} \right)^2 + (x+D) \tan \alpha + h.$$

4. La résolution de l'équation précédente pour $z = H$ et $x = 0$ donne, en posant $X = \tan \alpha$ et en transformant $\frac{1}{\cos^2 \alpha} = 1 + \tan^2 \alpha$:

$$\frac{gD^2}{2v_0^2}X^2 - DX - h + \frac{gD^2}{2v_0^2} + H = 0.$$

5. Les solutions doivent être réelles pour avoir un sens physiquement. Ce sera uniquement le cas si le discriminant est positif soit $D^2 - 4 \frac{gD^2}{2v_0^2} \left(-h + \frac{gD^2}{2v_0^2} + H \right) > 0$. Cette condition peut s'écrire :

$$v_0^4 - 2g(H-h)v_0^2 - g^2D^2 \geq 0. \quad (15.8)$$

6. En posant $Y = v_0^2$, cette relation s'écrit $Y^2 - 2g(H-h)Y - g^2D^2$. Son discriminant est $\delta = g^2(H-h)^2 + g^2D^2 > 0$ et ses deux racines sont $Y_{\pm} = g(H-h) \pm g\sqrt{(H-h)^2 + D^2}$, l'une positive et l'autre négative. Ce polynôme prend des valeurs négatives entre ses racines et positives à l'extérieur. Comme $Y = v_0^2 \geq 0$, la condition (15.8) est vérifiée si v_0^2 est supérieure à la racine positive soit $v_0 \geq \sqrt{g \left(H-h + \sqrt{(H-h)^2 + D^2} \right)}$.

7. L'application numérique donne pour un lancer franc $v_0 \geq 7,59 \text{ m}\cdot\text{s}^{-1} = 27,3 \text{ km}\cdot\text{h}^{-1}$ et pour un tir à trois points $v_0 \geq 8,60 \text{ m}\cdot\text{s}^{-1} = 31,0 \text{ km}\cdot\text{h}^{-1}$.

8. La résolution à v_0 fixée de l'équation du second degré en $X = \tan \alpha$ conduit à des solutions réelles puisque les conditions précédentes sont vérifiées. Elles s'écrivent :

$$X = \tan \alpha = \frac{v_0^2 \pm \sqrt{v_0^4 - 2g(H-h)v_0^2 - g^2D^2}}{gD}.$$

On a donc deux valeurs possibles pour l'angle α .

9. L'application numérique donne $\tan \alpha = 0,52$ ou $3,93$ soit $\alpha = 27^\circ 28'$ ou $\alpha = 75^\circ 42'$.

10. La résolution à α fixé de l'équation en v_0 donne $v_0 = \sqrt{\frac{gD^2}{2\cos^2 \alpha (h-H+D\tan \alpha)}}$.

11. L'application numérique donne $v_0 = 8,75 \text{ m}\cdot\text{s}^{-1}$.

12. Le bord du ballon se situe entre $0,00 \text{ cm}$ et 45 cm donc son centre est entre $17,8 \text{ cm}$ et $27,2 \text{ cm}$.

13. D'après la question 3, en posant $Z = x + D$, on doit résoudre

$$gZ^2 - Zv_0^2 \sin 2\alpha + 2v_0^2(H-h)\cos^2 \alpha = 0$$

14. Cette équation admet au moins une solution si son discriminant est positif soit $v_0^2 \sin^2 \alpha \geq 2g(H-h)$.

15. A v_0 fixée, il faut donc que $\sin^2 \alpha \geq 0,21$ et comme on effectue le tir vers l'avant et vers le haut, l'angle α vérifie $0 \leq \alpha \leq 90^\circ$. On en déduit finalement que $26^\circ 59' \leq \alpha \leq 90^\circ$.

16. A α fixé, la condition s'écrit $v_0 \geq 4,83 \text{ m}\cdot\text{s}^{-1}$.

17. On écrit la solution $x = \frac{v_0^2 \sin 2\alpha \pm v_0 \sqrt{v_0^2 \sin^2 2\alpha - 8(H-h)\cos^2 \alpha}}{2g} - D$.

15.7 Tobogan

1. Le référentiel lié au sol est galiléen. On définit le repère A, AX, AY comme sur la figure (15.25). Les forces s'exerçant sur P sont le poids, la réaction normale et la réaction tangentielle.

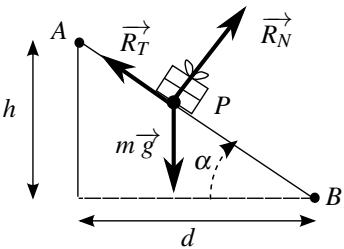


Figure 15.25

On écrit la relation fondamentale de la dynamique : $m\vec{a} = m\vec{g} + \vec{R}_T + \vec{R}_N$. La force $\vec{R}_T$ est une force de frottement donc elle s'oppose au mouvement. Les forces sont représentées sur

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

la figure (15.25). Les projections sur AX et AY donnent :

$$\begin{cases} m\ddot{X} &= mg \sin \alpha - R_T \\ m\ddot{Y} &= -mg \cos \alpha + R_N \end{cases}$$

où g , R_T et R_N sont les normes des vecteurs correspondants. Or $Y = \text{cte}$ donc $\ddot{Y} = 0$ soit $R_N = mg \cos \alpha$. On en déduit $R_T = fmg \cos \alpha$. On reporte dans l'équation sur AX et on intègre :

$$\ddot{X} = g(\sin \alpha - f \cos \alpha) \Rightarrow \dot{X} = gt(\sin \alpha - f \cos \alpha) + \text{cte}$$

La constante est nulle car la vitesse initiale est nulle puis :

$$X = \frac{gt^2}{2}(\sin \alpha - f \cos \alpha) + \text{cte}$$

La constante est nulle car $X(0) = 0$.

La position du point B correspond à $X = \frac{h}{\sin \alpha}$, donc B est atteint en un temps t_B tel que :

$$h = \frac{gt_B^2}{2}(\sin \alpha - f \cos \alpha) \sin \alpha \Rightarrow t_B = \sqrt{\frac{2h}{g \sin \alpha (\sin \alpha - f \cos \alpha)}}$$

Puisque $h = d$, l'angle α vaut $\pi/4$ soit $t_B = 1,8$ s.

2. Le chariot reste immobile $\delta t = 1$ s, il faut donc que le paquet parte au plus tard à $t_2 = t_B - \delta t = 0,8$ s avant l'arrivée du chariot et au plus tôt $t_B = 1,8$ s avant l'arrivée.

On calcule l'instant t_a d'arrivée du chariot en C_B : $t_a = C_0 C_B / v_c = 4$ s. Le joueur doit donc lâcher le paquet entre 2,2 s et 3,2 s.

15.8 Elastique de saut

1. Attention, ce n'est pas le poids du sauteur qui est donné, mais sa masse.

La tension est une force et s'exprime donc en N et non en kg. Ce qui est donné dans la fiche, correspond vraisemblablement à la masse m qui, accrochée à l'élastique donne l'allongement correspondant. Par application de la relation fondamentale à l'équilibre, on obtient $T = mg$, soit $T_1 = 1960$ N pour un allongement de 100% et $T_2 = 3185$ N pour un allongement de 200%.

2. Par définition l'allongement de l'élastique (en %) est $\frac{\ell - \ell_0}{\ell_0}$.

a. Un allongement de 100% correspond à $\ell = 2\ell_0$ et de 200% à $\ell = 3\ell_0$. Ainsi dans le deuxième cas, si la tension suivait la loi proposée, elle devrait être égale à 1,5 fois la première soit à 2940 N. Elle est en fait plus grande, ce qui signifie que lorsque l'élastique est très étirée k n'est plus une constante ; sa valeur augmente.

b. Pour l'allongement de 100%, on trouve $k_1 = T_1 / 2\ell_0 = 163 \text{ N} \cdot \text{m}^{-1}$ et pour l'allongement de 200%, on trouve $k_2 = T_2 / 3\ell_0 = 180 \text{ N} \cdot \text{m}^{-1}$ soit une valeur moyenne $k = 171,5 \text{ N} \cdot \text{m}^{-1}$ qui sera celle retenue dans la suite.

3. On étudie le mouvement de la masse test dans le référentiel terrestre galiléen. Cette masse est assimilée à son centre de gravité M de masse m soumis à son poids $m\vec{g}$ et à la tension de l'élastique $\vec{T}$. La relation fondamentale donne $m\vec{a} = m\vec{g} + \vec{T}$.

a. Le problème qui se pose avec les forces de tension d'élastique (ou de ressort) est de bien les orienter lors de la projection ; une méthode efficace est de définir le vecteur $\vec{u}_{\text{ext}}$; ce vecteur s'oriente du pont d'attache de l'élastique (ou du ressort) vers la masse en mouvement : on peut alors écrire $\vec{T} = -k(\ell - \ell_0)\vec{u}_{\text{ext}}$. La projection de la relation fondamentale sur un axe vertical Oz descendant ($\vec{u}_z = \vec{u}_{\text{ext}}$) avec O point d'attache du ressort (alors $\ell = z$) donne alors :

$$m\ddot{z} = mg - k(\ell - \ell_0) \Rightarrow \ddot{z} + \frac{k}{m}z = g + \frac{k}{m}\ell_0$$

La solution de l'équation est $z = A \cos(\omega_0 t) + B \sin(\omega_0 t) + \ell_0 - \frac{mg}{k}$, avec $\omega_0 = \sqrt{\frac{k}{m}}$. Dans la suite, on pose $\ell_e = \ell_0 + \frac{mg}{k}$.

On peut alors calculer la vitesse : $\dot{z} = \omega_0(-A \sin(\omega_0 t) + B \cos(\omega_0 t))$. On détermine A et B avec les conditions initiales, à savoir :

$$\begin{cases} z(t=0) = h \\ \dot{z}(t=0) = 0 \end{cases} \Rightarrow \begin{cases} A + \ell_e = h \\ B\omega_0 = 0 \end{cases}$$

d'où $z(t) = (h - \ell_e) \cos(\omega_0 t) + \ell_e$ et $\dot{z} = -\omega_0(h - \ell_e) \sin(\omega_0 t)$.

b. L'élastique se détend à l'instant t_1 tel que $z(t) = \ell_0$ soit :

$$t_1 = \frac{1}{\omega_0} \arccos\left(-\frac{mg}{k}\left(h - \frac{mg}{k} - \ell_0\right)^{-1}\right) = 1,25 \text{ s}$$

La vitesse $\overline{v_1}$ à l'instant t_1 est alors $\dot{z}(t_1) = -3,7 \text{ m}\cdot\text{s}^{-1}$.

c. Pour $\ell < \ell_0$, M n'est plus soumise qu'à son poids :

$$\ddot{z} = g \Rightarrow \dot{z} = g(t - t_1) + \overline{v_1} \Rightarrow z(t) = \frac{g}{2}(t - t_1)^2 + \overline{v_1}(t - t_1) + \ell_0.$$

Il faut prendre garde au fait que les conditions initiales pour cette étape du mouvement doivent être prises à l'instant t_1 et que l'axe est orienté vers le bas, d'où la projection de $\vec{g}$. La hauteur maximale atteinte, l'est lorsque la vitesse est nulle soit $t - t_1 = -\overline{v_1}/g$ et donc :

$$z_{\text{max}} = -\frac{g}{2}\left(\frac{\overline{v_1}}{g}\right)^2 - \overline{v_1}\frac{\overline{v_1}}{g} + \ell_0 = \ell_0 \frac{3\overline{v_1}^2}{2g} = 3,9 \text{ m}$$

4. Le saut se passe en deux étapes (inversées par rapport à la précédente partie) : une chute libre jusqu'à $z = \ell_0$ puis une chute soumise au poids et à la tension de l'élastique.

a. Il faut donc calculer la vitesse atteinte à la fin de la première étape en ℓ_0 pour avoir les conditions initiales pour la deuxième étape.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

Pour la première étape, le départ ayant lieu pour $z = 0$ au point d'attache de l'élastique (axe Oz vertical descendant comme précédemment) avec vitesse initiale nulle :

$$\ddot{z} = g \Rightarrow \dot{z} = gt \Rightarrow z(t) = \frac{g}{2}t^2$$

donc pour $z = \ell_0$, $t_2 = \sqrt{\frac{2\ell_0}{g}}$ et $\dot{z}(\ell_0) = \sqrt{2g\ell_0}$.

Pour la deuxième étape, la relation fondamentale donne : $m\vec{a} = m\vec{g} - k(z - \ell_0)\vec{u}_z$. Cette équation a déjà été résolue précédemment, cependant comme cette fois-ci on nous demande la hauteur maximale de chute, il est préférable de prendre la solution sous la forme : $z(t) = A' \cos(\omega_0 t + \phi) + \ell_0 + \frac{mg}{k}$ (où $\omega_0 = \sqrt{k/m}$). Avec les conditions initiales à $t = t_2$, on détermine A' (pas besoin de déterminer ϕ) :

$$\begin{cases} z(t_2) = A' \cos(\omega_0 t_2 + \phi) + \ell_0 + \frac{mg}{k} = \ell_0 \\ \dot{z}(t_2) = -A' \omega_0 \sin(\omega_0 t_2 + \phi) = \sqrt{2g\ell_0} \end{cases} \Rightarrow A' = \sqrt{\left(\frac{mg}{k}\right)^2 + \frac{2mg\ell_0}{k}}$$

Pour obtenir A' , on a élevé chaque équation au carré et utilisé la relation $\cos^2 + \sin^2 = 1$. L'application numérique donne : $A' = 7,6$ m. La hauteur maximale de chute est obtenue pour $\cos(\omega_0 t_2 + \phi) = 1$, soit $z_{\max} = A' + \ell_0 + \frac{mg}{k}$ soit $z_{\max} = 17,3$ m ce qui est pratiquement la hauteur du pont.

L'accélération est obtenue par deux dérivations de $z(t)$ soit : $\ddot{z} = -A' \omega_0^2 \cos(\omega_0 t_2 + \phi)$ soit $|\ddot{z}|_{\max} = A' \omega_0^2 = 20,1 \text{ m}\cdot\text{s}^{-2}$, soit plus de deux fois l'accélération de la pesanteur.

b. Dans l'équation du mouvement, la position d'équilibre est obtenue pour $\ddot{z} = 0$, soit $z = \ell_e = \ell_0 + \frac{mg}{k}$. C'est autour de cette position qu'ont lieu les oscillations du sauteur.

15.9 Jeux aquatiques

On considère le référentiel terrestre galiléen. Le système est le plongeur assimilé à son centre de gravité M de masse m .

1. La première partie du mouvement correspond à une chute libre. La relation fondamentale de la dynamique $m\vec{a} = m\vec{g}$ donne en projection sur l'axe (Oz) descendant $\ddot{z} = g$, soit en intégrant avec $\dot{z}(0) = 0$ et $z(0) = 0$:

$$\dot{z} = gt \quad \text{et} \quad z(t) = \frac{gt^2}{2}$$

Pour $z = h$, on trouve $t_c = \sqrt{2h/g} = 1,4$ s et $v_e = \sqrt{2gh} = 14,1 \text{ m}\cdot\text{s}^{-1}$.

2. Dans la deuxième partie du mouvement, le plongeur est dans l'eau.

a. La relation fondamentale dans l'eau devient : $m\vec{a} = m\vec{g} + \vec{\Pi} + \vec{f}_f$, ce qui donne en projection sur (Oz) et en posant $v_z = \dot{z}$:

$$m\dot{v}_z = mg \left(1 - \frac{1}{d_h}\right) - kv_z \Rightarrow \dot{v}_z + \frac{v_z}{\tau} = g \left(1 - \frac{1}{d_h}\right)$$

b. La solution générale de l'équation précédente est :

$$v_z = A \exp\left(-\frac{t}{\tau}\right) + g\tau\left(1 - \frac{1}{d_h}\right),$$

soit en changeant l'origine des temps, $v(t=0) = v_e$:

$$v_z(t) = g\tau\left(1 - \frac{1}{d_h}\right) + \left(v_e - g\tau\left(1 - \frac{1}{d_h}\right)\right) \exp\left(-\frac{t}{\tau}\right).$$

c. Lorsque $t \rightarrow \infty$, on trouve la vitesse limite, soit $v_L = g\tau\left(1 - \frac{1}{d_h}\right)$. L'application numérique donne $v_L = -0,356 \text{ m}\cdot\text{s}^{-1}$.

d. On peut alors mettre v_z sous la forme suivante :

$$v_z(t) = (v_e + |v_L|) \exp\left(-\frac{t}{\tau}\right) - |v_L|.$$

Le plongeur commence à remonter lorsque la vitesse v_z devient nulle, puisque la résultante des forces est vers le haut :

$$|v_L| = (v_e + |v_L|) \exp\left(-\frac{t_1}{\tau}\right) \Rightarrow t_1 = \tau \ln\left(1 + \frac{v_e}{|v_L|}\right) = 1,19 \text{ s}.$$

e. En intégrant la vitesse, avec $z(t=0) = 0$, on trouve :

$$z(t) = \tau(v_e + |v_L|)\left(1 - \exp\left(-\frac{t}{\tau}\right)\right) - |v_L|t.$$

La profondeur maximale est atteinte à $t = t_1 = 1,19 \text{ s}$ et vaut $z_{\max} = 4,1 \text{ m}$.

f. On cherche l'instant t_2 pour lequel $v_z = v_2 = 1 \text{ m}\cdot\text{s}^{-1}$. La résolution donne :

$$t_2 = \tau \ln\left(\frac{v_e + |v_L|}{v_2 + |v_L|}\right) = 0,76 \text{ s}, \quad \text{puis } z = 3,94 \text{ m}.$$

3. Les forces sont les mêmes que précédemment. Seules les conditions initiales changent.

a. La seule force qui a une projection selon l'horizontale est la force de frottement, ce qui donne :

$$m\ddot{x} = -k\dot{x} \quad \text{et} \quad m\ddot{z} = mg\left(1 - \frac{1}{d_h}\right) - k\dot{z}$$

b. L'intégration des équations donne :

$$\dot{x} = A \exp\left(-\frac{t}{\tau}\right) \quad \text{et} \quad \dot{z} = B \exp\left(-\frac{t}{\tau}\right) + g\tau\left(1 - \frac{1}{d_h}\right).$$

Sachant qu'à $t = 0$ (instant de pénétration dans l'eau) $\dot{x}(0) = v_0 \cos \alpha$ et $\dot{z}(0) = v_0 \sin \alpha$, on obtient :

$$\dot{x} = v_0 \cos \alpha \exp\left(-\frac{t}{\tau}\right) \quad \text{et} \quad \dot{z} = g\tau\left(1 - \frac{1}{d_h}\right) + \left(v_0 \sin \alpha - g\tau\left(1 - \frac{1}{d_h}\right)\right) \exp\left(-\frac{t}{\tau}\right).$$

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

La vitesse limite est obtenue pour $t \rightarrow \infty$. On trouve une vitesse limite verticale d'expression semblable à celle calculée précédemment. On peut mettre z sous une forme semblable à celle de la question précédente en remplaçant v_e par $v_0 \sin \alpha$.

c. Toute les expressions obtenues pour le saut, sont applicables dans le cas du plongeon en remplaçant v_e par $v_0 \sin \alpha$. Avec les nouvelles valeurs de k et donc de $\tau = 0,64$ s et de $v_L = -0,71 \text{ m}\cdot\text{s}^{-1}$, la vitesse s'annule donc à l'instant :

$$t_3 = \tau \ln \left(1 + \frac{v_0 \sin \alpha}{|v_L|} \right) = 1,52 \text{ s}$$

et à la profondeur :

$$z(t_3) = \tau(v_0 \sin \alpha + |v_L|) \left(1 - \exp \left(-\frac{t_3}{\tau} \right) \right) - |v_L|t_3 = 3,35 \text{ m.}$$

Le plongeur n'atteint donc pas le fond.

15.10 Etude d'un skieur

1. On étudie le mouvement du skieur assimilé à son centre de gravité dans le référentiel lié au sol galiléen. Les forces qui s'appliquent au skieur sont :

- le poids $m \vec{g}$;
- la réaction tangentielle $\vec{T}$;
- la réaction normale du support $\vec{N}$;
- la force de frottement $\vec{F}$.

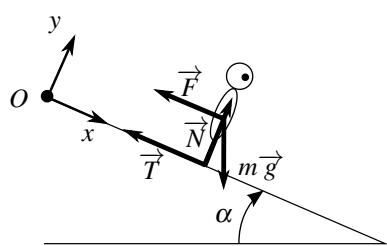


Figure 15.26

La relation fondamentale de la dynamique donne : $m \vec{a} = m \vec{g} + \vec{F} + \vec{T} + \vec{N}$, soit en projection :

$$m\ddot{x} = -T - \lambda \dot{x} + mg \sin \alpha \quad \text{et} \quad m\ddot{y} = 0 = N - mg \cos \alpha$$

On en déduit $N = mg \cos \alpha$ et donc $T = fN = fmg \cos \alpha$.

2. On reporte l'expression de T dans l'autre équation en posant $\tau = m/\lambda$ et on intègre :

$$\ddot{x} + \frac{\dot{x}}{\tau} = g(\sin \alpha - f \cos \alpha) \quad \Rightarrow \quad \dot{x} = A \exp \left(-\frac{t}{\tau} \right) + \tau g(\sin \alpha - f \cos \alpha)$$

Sachant qu'à $t = 0$ la vitesse est nulle, on obtient :

$$\dot{x} = \tau g(\sin \alpha - f \cos \alpha) \left(1 - \exp \left(-\frac{t}{\tau} \right) \right).$$

3. La présence de la force de frottement fluide dont la norme augmente avec la vitesse fait que la vitesse ne peut pas augmenter indéfiniment. Le skieur atteint une vitesse limite lorsque les frottements compensent la force motrice du mouvement. On a alors :

$$\ddot{x} = 0 \quad \text{et} \quad \frac{\dot{x}}{\tau} = g(\sin \alpha - f \cos \alpha).$$

a. On peut également trouver la vitesse limite avec l'expression de $\dot{x}$: lorsque $t \rightarrow \infty$, $\dot{x} \rightarrow \tau g(\sin \alpha - f \cos \alpha) = v_l$ qui constitue donc la vitesse limite. On obtient alors $\vec{v} = \vec{v}_l \left(1 - \exp\left(-\frac{t}{\tau}\right) \right)$ avec $\vec{v}_l = v_l \vec{u}_x$.

b. L'application numérique donne $v_l = 56 \text{ m}\cdot\text{s}^{-1}$.

4. On résout l'équation $v = v_l/2$:

$$\frac{v_l}{2} = v_l \left(1 - \exp\left(-\frac{t_1}{\tau}\right) \right) \Rightarrow t_1 = 2\tau \ln 2 = 111 \text{ s.}$$

5. Lorsque le skieur tombe, sa vitesse est $v_l/2$. L'équation du mouvement sur Oy est de la même forme que précédemment en multipliant f par 10, ce qui entraîne $T = 10fmg$. L'équation projetée sur Ox devient alors, en négligeant les frottements dus à l'air :

$$\ddot{x} = g(\sin \alpha - 10f \cos \alpha) \Rightarrow \dot{x} = g(\sin \alpha - 10f \cos \alpha)t + v_l/2$$

Le temps d'arrêt correspond à $v = 0$ soit $t_a = -v_l/(2g(\sin \alpha - 10f \cos \alpha))$. La distance d'arrêt depuis le point de chute, s'obtient en intégrant la vitesse soit :

$$d = \frac{g}{2}(\sin \alpha - 10f \cos \alpha)t_a^2 + \frac{v_l t_a}{2} \Rightarrow d = -\frac{v_l^2}{8g(\sin \alpha - 10f \cos \alpha)} = 6,2 \text{ m.}$$

15.11 Atterrissage en catastrophe d'un avion

1. On étudie le mouvement de l'avion dans le référentiel lié à la piste d'atterrissage galiléen. Ce mouvement est horizontal selon la direction de la piste repérée par l'axe (Ox).

a. On assimile le mouvement de l'avion à celui de son centre de gravité M de masse m . L'étude cinématique du mouvement de M donne :

$$\overrightarrow{OM} = x\vec{u}_x \quad ; \quad \vec{v} = \dot{x}\vec{u}_x \quad ; \quad \vec{a} = \ddot{x}\vec{u}_x.$$

Les forces s'exerçant sur l'avion sont : le poids, la réaction normale de la piste et la force de traînée T due au parachute. La seule force horizontale est donc T . La projection de la relation fondamentale de la dynamique sur (Ox) donne donc :

$$m\ddot{x} = -T = -k\dot{x}^2 \quad \text{avec} \quad k = C_x \rho S/(2m)$$

b. Pour intégrer l'équation du mouvement on sépare les variables, soit en notant $v = \dot{x}$:

$$\frac{\dot{v}}{v^2} = -k \Rightarrow -\frac{1}{v} = -kt + \text{cte}$$

or à $t = 0$, $v = v_a$: $\frac{1}{v} = kt + \frac{1}{v_a} \Rightarrow v = \frac{v_a}{v_a kt + 1}$.

c. L'avion s'arrête si la vitesse devient nulle, or d'après l'expression précédente la vitesse tend vers 0 au bout d'un temps infini. Cette force n'est pas suffisante pour stopper l'avion.

2. On obtient l'expression de la position de l'avion en intégrant l'expression précédente de la vitesse et en prenant $x(0) = 0$:

$$x = \frac{1}{k} \ln(v_a kt + 1) + \text{cte} \quad \text{soit} \quad x = \frac{1}{k} \ln(v_a kt + 1) \Leftrightarrow v_a kt + 1 = \exp(kx)$$

d'où la vitesse à la distance d : $v = v_a \exp(-kd) = 24,9 \text{ m}\cdot\text{s}^{-1}$ soit une vitesse de $89,6 \text{ km}\cdot\text{h}^{-1}$, d'où l'utilité des freins.

CHAPITRE 15 – PRINCIPES DE LA DYNAMIQUE NEWTONIENNE

Aspects énergétiques de la dynamique du point

16

L'objet de ce chapitre est de définir les notions énergétiques pour la mécanique du point. On va définir le travail et la puissance d'une force puis on établira les théorèmes de l'énergie cinétique et de l'énergie mécanique en introduisant la notion de forces conservatives et d'énergie potentielle.

1 Travail et puissance d'une force

1.1 Introduction et notations

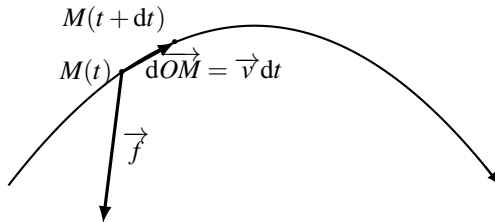


Figure 16.1 – Déplacement élémentaire d'un point M .

On étudie un point M animé d'une vitesse $\vec{v}$ dans un référentiel $\mathcal{R}$.

Le point M est soumis à une force $\vec{f}$.

À l'instant t , le point M se trouve en $M(t)$; à l'instant $t + dt$, il est situé en $M(t + dt)$.

Si dt est suffisamment petit, $M(t + dt)$ est infiniment proche de $M(t)$ et on note $d\vec{OM}$ le déplacement infinitésimal :

$$d\vec{OM} = \overrightarrow{M(t)M(t + dt)} = \vec{v} dt.$$

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

1.2 Puissance d’une force

La **puissance** de la force $\vec{f}$ appliquée au point matériel M animé de la vitesse $\vec{v}$ dans un référentiel $\mathcal{R}$ est définie par le produit scalaire :

$$\mathcal{P}(\vec{f}) = \vec{f} \cdot \vec{v}.$$

Étant donné que la vitesse dépend du référentiel d’étude et que la puissance dépend de la vitesse, on en déduit que **la puissance dépend du référentiel d’étude**. Certains auteurs adoptent la notation $\mathcal{P}_{|\mathcal{R}}$ pour insister sur ce point.

Une autre conséquence est que **la puissance est une grandeur additive**. En effet, si l’on considère deux forces $\vec{f}_1$ et $\vec{f}_2$ de résultante $\vec{f} = \vec{f}_1 + \vec{f}_2$, la puissance de la somme des forces :

$$\begin{aligned} \mathcal{P}(\vec{f}) &= \vec{f} \cdot \vec{v} = (\vec{f}_1 + \vec{f}_2) \cdot \vec{v} \\ &= \vec{f}_1 \cdot \vec{v} + \vec{f}_2 \cdot \vec{v}, \end{aligned}$$

est égale à la somme des puissances de chacune des deux forces :

$$\mathcal{P}(\vec{f}) = \mathcal{P}(\vec{f}_1) + \mathcal{P}(\vec{f}_2).$$

Pour qualifier la puissance d’une force $\vec{f}$, on distingue trois cas :

- la puissance de $\vec{f}$ est **motrice** si $\mathcal{P}(\vec{f}) > 0$. Cela correspond au cas où la projection de la force sur la trajectoire est dans le sens du mouvement,
- la puissance de $\vec{f}$ est **résistante** si $\mathcal{P}(\vec{f}) < 0$. Cela correspond au cas où la projection de la force sur la trajectoire est opposée au mouvement,
- la puissance de $\vec{f}$ est nulle si $\mathcal{P}(\vec{f}) = 0$. Cela correspond au cas où la force est perpendiculaire au mouvement (ou que le point M est immobile, ce qu’on n’envisage pas ici).

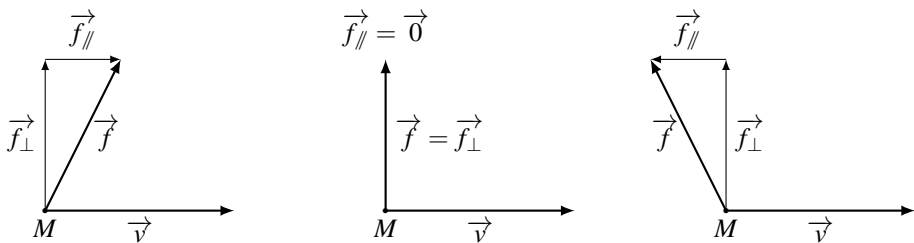


Figure 16.2 – À gauche, la projection $\vec{f}_{\parallel}$ de $\vec{f}$ sur $\vec{v}$ est dans le sens de $\vec{v}$, la puissance est motrice. À droite, $\vec{f}_{\parallel}$ est de sens opposé à $\vec{v}$, la puissance est résistante. Au centre, la force $\vec{f}$ est perpendiculaire à $\vec{v}$, la puissance est nulle.

1.3 Travail élémentaire d'une force

Le **travail élémentaire** d'une force $\vec{f}$ de puissance $\mathcal{P}(\vec{f})$ appliquée en un point M pendant un intervalle de temps dt est la quantité :

$$\delta W(\vec{f}) = \mathcal{P}(\vec{f}) dt.$$

Tout comme la puissance, le travail est une grandeur additive et dépend du référentiel d'étude. À partir de cette définition, on peut reformuler l'expression du travail élémentaire en introduisant le déplacement infinitésimal $d\vec{OM}$:

$$\delta W(\vec{f}) = \mathcal{P}(\vec{f}) dt = \vec{f} \cdot \vec{v} dt = \vec{f} \cdot d\vec{OM}.$$

L'écriture $\delta W(\vec{f}) = \mathcal{P}(\vec{f}) dt$ est utilisée avec la variable temporelle alors que l'écriture $\delta W(\vec{f}) = \vec{f} \cdot d\vec{OM}$ est utilisée avec les variables d'espace.

Le travail caractérise un échange d'énergie du système avec l'extérieur par l'intermédiaire de la force $\vec{f}$. C'est une grandeur d'échange qui n'existe ni à l'instant t ni à l'instant $t + dt$ mais seulement au cours du déplacement entre t et $t + dt$.

Remarque

Le travail infinitésimal δW est noté avec un δ et pas un d . La notation d est réservée aux différentielles définies comme des variations d'une fonction entre des points infiniment proches. Par exemple, on note $d\vec{OM} = \vec{OM}(t + dt) - \vec{OM}(t)$. Le travail n'est pas défini aux instants t et $t + dt$, on ne peut donc pas le définir comme une différentielle. Pour s'en souvenir, on note généralement δ mais certains auteurs barrent le d et notent $\bar{d}$. Dans tous les cas, il faut proscrire le d .

1.4 Travail d'une force au cours d'un déplacement

Au cours d'un déplacement le long d'une trajectoire $\widehat{AB}$ allant de A vers B , au cours duquel le mobile M quitte A à l'instant t_A et arrive en B à l'instant $t_B > t_A$, le travail de la force $\vec{f}$ correspond à la somme des travaux élémentaires calculés sur la trajectoire $\widehat{AB}$:

$$W_{A \rightarrow B}(\vec{f}) = \int_{M \in \widehat{AB}} \delta W(\vec{f}).$$

La notation $\int_{M \in \widehat{AB}}$ indique que le travail doit être calculé par une intégrale curviligne en suivant la trajectoire amenant le mobile de A vers B . Le travail dépend de la force $\vec{f}$ et de la manière dont on réalise le déplacement entre A et B . Le travail dépend du chemin suivi.

Avec la variable temporelle, $\delta W(\vec{f}) = \mathcal{P}(t) dt$ puis $W_{A \rightarrow B}(\vec{f}) = \int_{t_A}^{t_B} \mathcal{P}(t) dt$.

Avec les variables d'espace, $\delta W(\vec{f}) = \vec{f} \cdot d\vec{OM}$ puis $W_{A \rightarrow B}(\vec{f}) = \int_{M \in \widehat{AB}} \vec{f} \cdot d\vec{OM}$.

2 Premiers exemples de calculs de travaux

2.1 Travail d'une force constamment perpendiculaire au mouvement

On considère une force $\vec{f}$ qui reste en permanence perpendiculaire au mouvement. Parmi les forces vues au chapitre précédent, la force magnétique d'expression $\vec{f} = q \vec{v} \wedge \vec{B}$ rentre dans cette catégorie ; les forces de contact normales au support également lorsque le support est fixe dans le référentiel d'étude.

La force $\vec{f}$ reste toujours perpendiculaire au mouvement. Dans ces conditions, sa puissance est nulle à chaque instant : $\mathcal{P}(\vec{f}) = \vec{f} \cdot \vec{v} = 0$. Son travail l'est également :

$$\delta W(\vec{f}) = \mathcal{P} dt = 0 \implies W_{A \rightarrow B}(\vec{f}) = 0.$$

2.2 Travail d'une force constante

On considère une force $\vec{f}_0$ constante. Parmi les forces vues au chapitre précédent, le poids rentre dans cette catégorie et alors $\vec{f}_0 = m \vec{g}$; la force électrique également, à condition que le champ électrique $\vec{E}$ soit uniforme, et alors $\vec{f}_0 = q \vec{E}$.

On peut poser $f_0 = \|\vec{f}_0\|$ et définir un axe (Oz) tel que $\vec{u}_z = \frac{\vec{f}_0}{f_0}$. En notant z la coordonnée de M sur l'axe (Oz) , la puissance de $\vec{f}_0$ s'écrit : $\mathcal{P}(\vec{f}_0) = \vec{f}_0 \cdot \vec{v} = f_0 \vec{u}_z \cdot \vec{v} = f_0 \dot{z}$.

On en déduit que son travail infinitésimal vaut : $\delta W(\vec{f}_0) = \mathcal{P}(\vec{f}_0) dt = f_0 \dot{z} dt = f_0 dz$.

En notant z_A et z_B les coordonnées de A et B sur l'axe (Oz) , on obtient alors :

$$W_{A \rightarrow B}(\vec{f}_0) = \int_{M \in \widehat{AB}} f_0 dz = f_0 \left[dz \right]_{z_A}^{z_B} = f_0 (z_B - z_A).$$

Dans ce cas, le travail infinitésimal est la dérivée d'une fonction : $\delta W(\vec{f}_0) = d(f_0 z)$, et le travail sur une trajectoire partant de A et revenant en A est nul : $W_{A \rightarrow A}(\vec{f}_0) = f_0 (z_A - z_A) = 0$. Une telle force est une force conservative.

2.3 Travail d'une force de frottement de norme constante

On considère maintenant une force de frottement $\vec{f}$ dont seule la norme f est constante.

La force de frottement $\vec{f}$ étant opposée à la vitesse, sa puissance est toujours négative et s'exprime ainsi : $\mathcal{P}(\vec{f}) = \vec{f} \cdot \vec{v} = -f \|\vec{v}\| < 0$. Son travail infinitésimal s'écrit :

$$\delta W(\vec{f}) = -f \|\vec{v}\| dt = -f \| \vec{v} dt \| = -f dl$$

où dl est la longueur du déplacement infinitésimal. On en déduit :

$$W_{A \rightarrow B}(\vec{f}) = -f l_{\widehat{AB}} \quad \text{où } l_{\widehat{AB}} \text{ est la longueur du trajet } \widehat{AB}.$$

Dans ce cas, le travail infinitésimal ne s'écrit pas comme une différentielle. Le travail sur une trajectoire fermée partant de A et revenant en A est strictement négatif. Une telle force est une force dissipative.

3 Théorème de l'énergie cinétique

3.1 Définition de l'énergie cinétique

L'énergie cinétique d'un point matériel de masse m animée d'une vitesse $\vec{v}$ dans le référentiel $\mathcal{R}$ est définie par la quantité scalaire :

$$E_c = \frac{1}{2}mv^2.$$

Étant donné que la vitesse dépend du référentiel d'étude et que l'énergie cinétique dépend de la vitesse, on en déduit que **l'énergie cinétique dépend du référentiel d'étude**. Certains auteurs adoptent la notation $E_{c|\mathcal{R}}$ pour insister sur ce point.

Par ailleurs, il faut noter que l'énergie cinétique ne dépend que de la norme de la vitesse et pas de sa direction. Elle est particulièrement utile lorsque l'on veut connaître uniquement la norme de la vitesse (soit parce que la direction est connue, soit parce qu'on restreint le problème à cet aspect des choses).

3.2 Théorème de l'énergie cinétique en référentiel galiléen

a) Énoncé du théorème

La variation d'énergie cinétique d'un point matériel entre deux instants est égale au travail des forces qui s'exercent sur ce point entre les deux instants considérés.

b) Démonstration du théorème

Il s'agit d'une conséquence du principe fondamental de la dynamique qui s'écrit :

$$m \frac{d\vec{v}}{dt} = \sum_i \vec{f}_i.$$

En multipliant scalairement cette relation par $\vec{v}$, on obtient :

$$m \frac{d\vec{v}}{dt} \cdot \vec{v} = \left(\sum_i \vec{f}_i \right) \cdot \vec{v} = \sum_i (\vec{f}_i \cdot \vec{v}) = \sum_i \mathcal{P}(\vec{f}_i)$$

où $\mathcal{P}(\vec{f}_i)$ est la puissance de la force $\vec{f}_i$. Or :

$$m \frac{d\vec{v}}{dt} \cdot \vec{v} = m \frac{d}{dt} \left(\frac{\vec{v} \cdot \vec{v}}{2} \right) = \frac{d}{dt} \left(\frac{1}{2}mv^2 \right) = \frac{dE_c}{dt},$$

soit finalement :

$$\frac{dE_c}{dt} = \sum_i \mathcal{P}(\vec{f}_i).$$

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

On intègre alors cette relation par rapport au temps entre l’instant t_A où le mobile quitte A avec une vitesse v_A et l’instant t_B où il atteint B avec la vitesse v_B :

$$\int_{t_A}^{t_B} \frac{dE_c}{dt} dt = \int_{t_A}^{t_B} \left(\sum_i \mathcal{P}(\vec{f}_i) \right) dt.$$

Or, d’une part :

$$\int_{t_A}^{t_B} \left(\sum_i \mathcal{P}(\vec{f}_i) \right) dt = \sum_i \int_{t_A}^{t_B} \mathcal{P}(\vec{f}_i) dt = \sum_i W_{A \rightarrow B}(\vec{f}_i)$$

est égal à la somme des travaux des forces agissant sur M entre A et B et d’autre part :

$$\int_{t_A}^{t_B} \frac{dE_c}{dt} dt = \left[E_c(t) \right]_{t_A}^{t_B} = E_c(t_B) - E_c(t_A) = \frac{1}{2}mv_B^2 - \frac{1}{2}mv_A^2$$

est la variation d’énergie cinétique de M entre A et B que l’on note souvent ΔE_c .

Le théorème de l’énergie cinétique peut se formuler de deux manières :

- Loi de l’énergie cinétique :

$$\Delta E_c = E_c(t_B) - E_c(t_A) = \frac{1}{2}mv_B^2 - \frac{1}{2}mv_A^2 = \sum_i W_{A \rightarrow B}(\vec{f}_i)$$

que l’on peut également écrire au niveau infinitésimal : $dE_c = \sum_i \delta W(\vec{f}_i)$.

- Loi de la puissance cinétique :

$$\frac{dE_c}{dt} = \sum_i \mathcal{P}(\vec{f}_i).$$

c) Discussion suivant le signe de $W_{A \rightarrow B}(\vec{f})$

- Si $\sum_i W_{A \rightarrow B}(\vec{f}_i) > 0$, le travail est moteur, l’énergie cinétique augmente donc la norme de la vitesse augmente. Le mouvement est accéléré.
- Si $\sum_i W_{A \rightarrow B}(\vec{f}_i) < 0$, le travail est résistant, l’énergie cinétique diminue donc la norme de la vitesse diminue. Le mouvement est décéléré.
- Si $\sum_i W_{A \rightarrow B}(\vec{f}_i) = 0$, le travail est nul, l’énergie cinétique est conservée donc la norme de la vitesse est constante. Le mouvement est uniforme.



Pour montrer qu’un mouvement est uniforme, il faut avoir le réflexe d’utiliser le théorème de l’énergie cinétique. Plus généralement, c’est souvent le théorème à utiliser pour déterminer si un mouvement est accéléré, décéléré ou uniforme.

3.3 Utilisation du théorème de l'énergie cinétique

L'utilisation du théorème de l'énergie cinétique est de deux sortes :

- on déduit le travail de la somme des forces appliquées à partir de la connaissance de la variation d'énergie cinétique,
- on déduit la variation de l'énergie cinétique connaissant les travaux des forces.

Cette deuxième utilisation est souvent délicate car le travail des forces s'exprime rarement de façon simple. En revanche, il est particulièrement efficace dans le cas où une expression du travail est connue.

3.4 Intérêt d'une approche énergétique

En mécanique du point, le théorème de l'énergie cinétique est une conséquence du principe fondamental de la dynamique et n'introduit pas de nouveau postulat. Une approche énergétique revient implicitement à utiliser le principe fondamental de la dynamique.

En revanche, l'expression vectorielle du principe fondamental de la dynamique conduit par projection à trois équations scalaires tandis que le théorème de l'énergie cinétique ne fournit qu'une équation scalaire. En passant du principe fondamental de la dynamique au théorème de l'énergie cinétique, on a perdu deux équations scalaires. Dans la démonstration, cette perte a lieu lorsque l'on multiplie le principe fondamental de la dynamique par $\vec{v}$, ce qui conduit implicitement à projeter sur la trajectoire.

Lorsque l'on traite un problème dont la résolution ne concerne qu'une variable, une seule équation suffit et une méthode énergétique est appropriée. Les problèmes de ce type sont appelés problèmes à un degré de liberté. Si le problème a plus d'un degré de liberté, la seule utilisation du théorème de l'énergie cinétique ne permet pas la résolution complète du problème et on doit alors recourir au principe fondamental de la dynamique.

On retiendra qu'une méthode énergétique est adaptée aux cas où un seul paramètre de position suffit à décrire l'évolution du système.

3.5 Étude d'un problème à l'aide du théorème de l'énergie cinétique

a) Énoncé du problème

On considère un pratiquant de saut à ski qui s'élance sur un tremplin (figure 16.3). Il part sans élan du sommet A de la piste et parcourt la distance $e = 100$ m jusqu'à la tête de tremplin B d'où il saute. Au cours de la prise d'élan, il descend une dénivellation $h = 50$ m. Pour des raisons de sécurité, la vitesse du sauteur en B ne doit pas dépasser $v_{max} = 30 \text{ m}\cdot\text{s}^{-1}$. Dans un premier temps, on veut savoir si les conditions de sécurité sont respectées lorsque les conditions de neige et de vent sont tellement bonnes que l'on peut négliger tout frottement. Ensuite, on va évaluer le travail des forces de frottement dans des conditions de glisse normales.

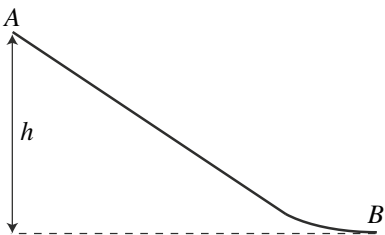


Figure 16.3 – Profil d'un tremplin de saut à ski.

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

b) Mise en équation et résolution du problème en l'absence de frottement

- Système : on modélise le skieur par un point matériel M de masse m .
- Référentiel : on étudie le mouvement de M dans le référentiel terrestre supposé galiléen.
- Bilan des forces : le skieur est soumis à :
 - son poids $\vec{P} = m\vec{g}$;
 - la force de contact normale de la neige sur ses skis $\vec{R}$ qui assure son mouvement le long de la piste ;
 - des forces de frottements solides (sous les skis) et fluides (résistance de l'air) de résultante $\vec{f}$ opposée au mouvement que l'on néglige car les conditions de glisse sont idéales.
- Méthode de résolution : on utilise le théorème de l'énergie cinétique.

On applique le théorème de l'énergie cinétique entre le point A et le point B :

$$E_c(B) - E_c(A) = W_{A \rightarrow B}(\vec{R}) + W_{A \rightarrow B}(\vec{P}) \quad (16.1)$$

et on détermine les deux membres de cette égalité.

Le sauteur part de A avec une vitesse $\vec{v}_A = \vec{0}$. Il atteint le point B avec une vitesse $\vec{v}_B$ de norme v_B . La variation d'énergie cinétique entre A et B vaut donc :

$$\Delta E_c = E_c(B) - E_c(A) = \frac{1}{2}mv_B^2 - \frac{1}{2}mv_A^2 = \frac{1}{2}mv_B^2.$$

Le travail de la force de contact est nul puisqu'elle est en permanence perpendiculaire au mouvement ce qui implique que :

$$\mathcal{P}(\vec{R}) = \vec{R} \cdot \vec{v} = 0 \quad \text{puis} \quad W_{A \rightarrow B}(\vec{R}) = 0.$$

Il ne reste qu'à calculer le travail du poids entre le point A et le point B . Pour cela, on définit l'axe (Oz) vertical ascendant de sorte que $\vec{P} = -mg\vec{u}_z$ et on trouve alors :

$$\delta W(\vec{P}) = -mg\vec{u}_z \cdot d\vec{OM} = -mgdz,$$

puis :

$$W_{A \rightarrow B}(\vec{P}) = \int_{z_A}^{z_B} -mgdz = -mg \left[z \right]_{z_A}^{z_B} = -mg(z_B - z_A) = mg(z_A - z_B) = mgh.$$

Le travail du poids est positif : c'est le moteur du mouvement du skieur vers le bas. En injectant les relations précédentes dans l'équation (16.1), on obtient alors :

$$\frac{1}{2}mv_B^2 = mgh \quad \text{puis} \quad v_B = \sqrt{2gh}.$$

Numériquement, $v_B = \sqrt{2 \times 10 \times 50} = 32 \text{ m} \cdot \text{s}^{-1}$ qui est supérieure à la vitesse limite de sécurité de $30 \text{ m} \cdot \text{s}^{-1}$. Le tremplin n'est pas sûr par conditions idéales de glisse.

c) Calcul des forces de frottement lors d'une prise d'élan chronométrée

Lors d'un essai pour lequel les conditions de glisse ne sont pas idéales, les responsables de la sécurité du tremplin mesurent la vitesse en tête du tremplin : $v_B = 25 \text{ m}\cdot\text{s}^{-1}$. On veut évaluer le travail des forces de frottement au cours de cette descente pour un sauteur de masse $m = 75 \text{ kg}$ équipé.

Comme précédemment, on applique le théorème de l'énergie cinétique :

$$E_c(B) - E_c(A) = W_{A \rightarrow B}(\vec{R}) + W_{A \rightarrow B}(\vec{P}) + W_{A \rightarrow B}(\vec{f}),$$

d'où on tire :

$$W_{A \rightarrow B}(\vec{f}) = E_c(B) - E_c(A) - W_{A \rightarrow B}(\vec{R}) - W_{A \rightarrow B}(\vec{P}),$$

et, en injectant les expressions des travaux trouvées précédemment, on obtient :

$$W_{A \rightarrow B}(\vec{f}) = \frac{1}{2}mv_B^2 - mgh.$$

Numériquement, on trouve alors $W_{A \rightarrow B}(\vec{f}) = 0,5 \times 75 \times 25^2 - 75 \times 10 \times 50 = -14 \text{ kJ}$.

Le travail des forces de frottement est négatif, ce qui correspond bien à une force résistante. Il est en général très difficile de distinguer la part dissipée par les frottements des skis sur la neige de la part dissipée par la résistance de l'air.

4 Énergie potentielle et forces conservatives

4.1 Définitions

On dit qu'une force $\vec{f}$ **dérive d'un potentiel** ou encore qu'elle est **conservative** si son travail $W_{A \rightarrow B}(\vec{f})$ entre deux points A et B ne dépend pas du chemin suivi mais uniquement des points A et B . On peut alors écrire le travail $W_{A \rightarrow B}(\vec{f})$ comme la différence $E_p(A) - E_p(B)$ où E_p est une fonction de la variable de position. Au niveau élémentaire, la relation devient :

$$\delta W(\vec{f}) = -dE_p.$$

On appelle alors **énergie potentielle de la force** cette fonction E_p .

Une force $\vec{f}$ est dite conservative si on peut trouver une fonction énergie potentielle E_p telle que

$$\vec{f} \cdot d\vec{OM} = \delta W(\vec{f}) = -dE_p.$$

L'énergie potentielle est définie à partir de sa variation, elle n'est donc déterminée qu'à une constante près : si E_p est une énergie potentielle possible, $E_p + K$ où K est une constante est également une énergie potentielle possible. On peut donc fixer une position arbitraire O pour laquelle $E_p(O) = 0$. Le point O est alors le point de **référence de l'énergie potentielle**.

Remarque

Lors de la définition du travail, on a insisté sur le fait que le travail n'est en général pas une différentielle. C'est pour cela qu'on le note δW avec un δ plutôt qu'un d . Les forces conservatives sont les forces singulières pour lesquelles le travail est une différentielle.

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

4.2 Exemples de forces conservatives

a) Poids d'un corps

Le poids d'un point matériel de masse m dans le champ de pesanteur $\vec{g}$ vaut : $m\vec{g}$, soit en projection dans un système de coordonnées cartésiennes dont l'axe (Oz) est vertical et dirigé vers le haut : $-mg\vec{u}_z$. Le travail de cette force s'écrit donc :

$$\delta W = m\vec{g} \cdot d\vec{OM} = (-mg\vec{u}_z) \cdot d\vec{OM} = -mg(\vec{u}_z \cdot d\vec{OM}).$$

Or $\vec{u}_z \cdot d\vec{OM}$ est la composante du déplacement infinitésimal selon $\vec{u}_z$ donc $\vec{u}_z \cdot d\vec{OM} = dz$. On en déduit :

$$\delta W = -mgdz = -d(mgz) = -dE_p \quad \text{puis} \quad E_p = mgz + C$$

où C est une constante d'intégration que l'on peut choisir de manière arbitraire.

Dans cette expression, l'axe (Oz) est dirigé selon la verticale ascendante. Il s'agit donc d'un axe des altitudes, et l'énergie potentielle augmente lorsque l'altitude augmente. On peut alors retenir la formule ci-dessus sous la forme :

$$E_p = mg \cdot \text{altitude.}$$

Le choix arbitraire de l'altitude 0, c'est-à-dire de la référence des altitudes, correspond au choix arbitraire de la constante C .

b) Force gravitationnelle exercée par un astre ponctuel

On considère un point matériel M de masse m soumis à la force gravitationnelle exercée par un astre de centre P et de masse m_p . M est soumis à la force gravitationnelle :

$$\vec{F}_{P \rightarrow M} = -\mathcal{G} \frac{mm_p}{r^2} \vec{u}_r,$$

où $r = \|\vec{PM}\|$ et $\vec{u}_r = \frac{\vec{PM}}{r}$. Le travail élémentaire de cette force s'écrit :

$$\delta W = \vec{F}_{P \rightarrow M} \cdot d\vec{PM} = -\mathcal{G} \frac{mm_p}{r^2} \vec{u}_r \cdot d\vec{PM} = -\mathcal{G} \frac{mm_p}{r^2} \left(\frac{\vec{PM}}{r} \cdot d\vec{PM} \right).$$

Or $\vec{PM} \cdot d\vec{PM} = d\left(\frac{\vec{PM}^2}{2}\right) = d\left(\frac{r^2}{2}\right) = r dr$ donc

$$\delta W = -\mathcal{G} \frac{mm_p}{r^2} dr = \mathcal{G} mm_p d\left(-\frac{1}{r}\right) = \mathcal{G} mm_p d\left(\frac{1}{r}\right) = d\left(\frac{\mathcal{G} mm_p}{r}\right) = -dE_p.$$

On en déduit l'énergie potentielle associée : $E_p = -\frac{\mathcal{G} mm_p}{r} + C$, où C est une constante. On prend généralement comme référence des énergies potentielles $E_p \rightarrow 0$ lorsque $r \rightarrow \infty$, ce qui revient à choisir $C = 0$ et on obtient :

$$E_p = -\frac{\mathcal{G}mm_p}{r}.$$

Remarque

Pour établir cette relation, on peut également remarquer que l'expression de la force $\vec{F}_{P \rightarrow M} = -\mathcal{G} \frac{mm_p}{r^2} \vec{u}_r$ suppose de travailler en coordonnées sphériques d'origine P . Or, le déplacement infinitésimal radial en sphérique vaut dr d'où $\vec{u}_r \cdot d\vec{PM} = dr$.

c) Force de rappel d'un élastique d'un ressort

Un ressort de raideur k ayant subi un allongement $\Delta l = l - l_0$ exerce une force de rappel élastique dans la direction de l'allongement :

$$\vec{f} = -k\Delta l \vec{u}_{ext}.$$

où le vecteur unitaire $\vec{u}_{ext}$ est dirigé du ressort vers le système. Lors d'un déplacement infinitésimal $d\vec{OM}$ du point M , le travail élémentaire de la force de rappel élastique vaut :

$$\delta W = \vec{f} \cdot d\vec{OM} = -k(l - l_0) \vec{u}_{ext} \cdot d\vec{OM}.$$

Le produit scalaire $\vec{u}_{ext} \cdot d\vec{OM}$ est la projection du déplacement infinitésimal dans la direction du ressort et conduit à un variation de la longueur du ressort de dl . On en déduit le travail élémentaire :

$$\delta W = -k(l - l_0) dl = -dE_p$$

puis

$$E_p = \frac{1}{2}k(l - l_0)^2 + C \quad \text{où } C \text{ est une constante.}$$

On prend généralement comme référence des énergies potentielles $E_p = 0$ lorsque le ressort est à sa longueur à vide donc que $l = l_0$. Cela revient à choisir $C = 0$ et on obtient alors l'expression l'énergie potentielle élastique dont dérive la force de rappel du ressort :

$$E_p = \frac{1}{2}k(l - l_0)^2.$$

d) Force électrique

Force coulombienne exercée par une charge ponctuelle On considère un point matériel M de charge q soumis à la force coulombienne exercée par une charge q_p placée en P . M est soumis à la force coulombienne :

$$\vec{F}_{P \rightarrow M} = \frac{1}{4\pi\epsilon_0} \frac{qq_p}{r^2} \vec{u}_r,$$

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

où $r = \|\vec{PM}\|$ et $\vec{u}_r = \frac{\vec{PM}}{r}$. Le travail élémentaire de cette force s'écrit :

$$\delta W = \vec{F}_{P \rightarrow M} \cdot d\vec{PM} = \frac{1}{4\pi\epsilon_0} \frac{qq_p}{r^2} \vec{u}_r \cdot d\vec{PM} = \frac{1}{4\pi\epsilon_0} \frac{qq_p}{r^2} \left(\frac{\vec{PM}}{r} \cdot d\vec{PM} \right).$$

Comme pour la force gravitationnelle, $\vec{PM} \cdot d\vec{PM} = r dr$ donc :

$$\delta W = \frac{1}{4\pi\epsilon_0} \frac{qq_p}{r^2} dr = -\frac{qq_p}{4\pi\epsilon_0} \left(-\frac{dr}{r^2} \right) = -\frac{qq_p}{4\pi\epsilon_0} d\left(\frac{1}{r} \right) = -d\left(\frac{qq_p}{4\pi\epsilon_0 r} \right) = -dE_p.$$

On déduit alors l'énergie potentielle correspondante :

$$E_p = \frac{qq_p}{4\pi\epsilon_0 r} + C \quad \text{où } C \text{ est une constante.}$$

On prend généralement comme référence des énergies potentielles $E_p \rightarrow 0$ lorsque $r \rightarrow \infty$, ce qui revient à choisir $C = 0$ et on obtient :

$$E_p = \frac{qq_p}{4\pi\epsilon_0 r}.$$

Force exercée par un champ électrique uniforme On étudie un point matériel de charge électrique q placé dans un champ électrique uniforme $\vec{E}$. Il est soumis à la force

$$\vec{F} = q\vec{E},$$

dont le travail élémentaire s'écrit :

$$\delta W = \vec{F} \cdot d\vec{OM} = q\vec{E} \cdot d\vec{OM}.$$

On note $E = \|\vec{E}\|$ et on choisit l'axe (Oz) de telle sorte que $\vec{E} = E\vec{u}_z$.

Le travail élémentaire de cette force s'écrit alors :

$$\delta W = (qE\vec{u}_z) \cdot d\vec{OM} = qE \left(\vec{u}_z \cdot d\vec{OM} \right).$$

Or $\vec{u}_z \cdot d\vec{OM} = dz$ d'où :

$$\delta W = qEdz = d(qEz) = -dE_p.$$

On déduit l'énergie potentielle correspondante :

$$E_p = -qEz + C,$$

où C est une constante.

4.3 Exemples de forces non conservatives

On considère une force de frottement fluide proportionnelle à la vitesse $\vec{f} = -\lambda \vec{v}$. Le travail élémentaire de cette force s'écrit :

$$\delta W = \vec{f} \cdot d\vec{OM} = -\lambda \vec{v} \cdot \vec{v} dt = -\lambda v^2 dt < 0.$$

Ce travail dépend non seulement du déplacement $d\vec{OM}$, mais aussi de la vitesse à laquelle il est fait. On ne peut pas l'écrire sous la forme d'une différentielle : il ne s'agit donc pas d'une force conservative. Les forces de frottement dissipant de l'énergie, on parle de **forces dissipatives**.

Remarque

Dans le cas d'une force conservative, on a : $W_{A \rightarrow B} = E_p(A) - E_p(B) = -W_{B \rightarrow A}$.

Ceci ne peut pas être réalisé pour les forces de frottement car le travail est toujours négatif : $W_{A \rightarrow B} < 0$ et $W_{B \rightarrow A} < 0$. Les frottements ne sont pas conservatifs.

5 Énergie mécanique

5.1 Définition de l'énergie mécanique

La quantité E_m qui est la somme de l'énergie cinétique et des énergies potentielles est appelée **énergie mécanique** du point matériel :

$$E_m = E_c + E_p.$$

Dans cette définition, E_p est l'énergie potentielle du système, c'est-à-dire la somme des énergies potentielles des différentes forces conservatives.

5.2 Conservation de l'énergie mécanique

On considère un système soumis à un ensemble de forces $\vec{f}_i$. Dans le cas où ces forces sont conservatives ou ne travaillent pas, on peut écrire pour chacune des forces :

$$W(\vec{f}_i) = -\Delta E_{pi}.$$

On définit alors l'énergie potentielle du système comme la somme des énergies potentielles dont dérive chacune des forces :

$$E_p = \sum_i E_{pi}.$$

Le théorème de l'énergie cinétique implique alors que :

$$\Delta E_c = \sum_i W(\vec{f}_i) = \sum_i -\Delta E_{pi} = -\Delta E_p \quad \text{soit} \quad \Delta(E_c + E_p) = 0.$$

On en déduit :

$$E_c + E_p = E_m \quad \text{où } E_m \text{ est une constante.}$$

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

L'énergie mécanique d'un point matériel soumis uniquement à des forces conservatives ou qui ne travaillent pas est une constante du mouvement.

Lorsque toutes les forces sont conservatives, le mouvement est qualifié de mouvement conservatif car l'énergie mécanique est conservée.

Dans le cas d'un mouvement conservatif, l'énergie mécanique est une quantité qui se conserve au cours du mouvement et qui n'est fonction que de la position et de ses dérivées premières par rapport au temps. L'énergie mécanique est alors une **intégrale première du mouvement**.

L'énergie mécanique est répartie sous deux formes : énergie cinétique et énergie potentielle. L'énergie potentielle peut être convertie en énergie cinétique et réciproquement, mais, pour un mouvement conservatif, la somme des deux formes d'énergie reste constante. La conservation de l'énergie mécanique et l'échange permanent entre les deux formes d'énergies ont été mis en évidence lors de l'étude de l'oscillateur harmonique.

5.3 Cas général : non conservation de l'énergie mécanique

Dans le cas général, il est nécessaire de distinguer les forces conservatives notées $\vec{f}_C$ de celles qui ne le sont pas que l'on notera $\vec{f}_{NC}$. En notant leurs travaux respectifs $W(\vec{f}_C)$ et $W(\vec{f}_{NC})$, on sait, d'après ce qui précède, que $W(\vec{f}_C)$ peut se mettre sous la forme : $W(\vec{f}_C) = -\Delta E_p$ où E_p est l'énergie potentielle du système. En revanche, il n'est pas aisé d'avoir une expression simple de $W(\vec{f}_{NC})$. L'expression du travail total vaut alors :

$$W_{tot} = W(\vec{f}_C) + W(\vec{f}_{NC}) = -\Delta E_p + W(\vec{f}_{NC}).$$

En appliquant le théorème de l'énergie cinétique à cette situation, on a donc :

$$\Delta E_c = W_{tot} = -\Delta E_p + W(\vec{f}_{NC})$$

puis

$$\Delta E_m = \Delta(E_c + E_p) = W(\vec{f}_{NC})$$

ou sous forme élémentaire :

$$dE_m = d(E_c + E_p) = \delta W(\vec{f}_{NC}).$$

Parfois ce résultat est appelé « théorème de l'énergie mécanique ».

La variation d'énergie mécanique au cours du mouvement est égale au travail des forces qui ne dérivent pas d'un potentiel, autrement dit des forces non conservatives.

C'est une autre formulation du théorème de l'énergie cinétique.

Il est à noter que le problème de la conservation de l'énergie fait partie d'un cadre beaucoup plus général que celui de la mécanique. On reviendra sur ce point dans le cours de thermodynamique et plus précisément lors de l'étude du premier principe qui postule la conservation de l'énergie totale. Le résultat qui vient d'être établi sera alors généralisé.

6 Étude qualitative des mouvements et des équilibres

6.1 Exemple introductif

On considère un jeu dans lequel une perle est enfilée sur une tige rigide formant des puits et des bosses (figure 16.4). Le profil d'altitude de la tige est noté $h(x)$ et la pesanteur dérive de l'énergie potentielle $E_p = mg \cdot \text{altitude} = mgh(x)$. La courbe d'altitude $h(x)$ est, à un facteur d'échelle près, identique à la courbe d'énergie potentielle de pesanteur. Tous ceux qui ont joué à ce jeu savent que :

- si on lâche la perle sans vitesse initiale, elle glisse spontanément vers le fond du puits le plus proche ;
- les positions où la perle peut rester en équilibre sont les sommets des bosses (positions instables) et les fonds des puits (positions stables) ;
- pour pouvoir sortir la perle d'un puits, il faut lui procurer une vitesse (une énergie cinétique) de norme suffisante pour pouvoir gravir la bosse d'à côté sur son élan.

Dans ce paragraphe, on va modéliser, expliquer et généraliser ces observations.

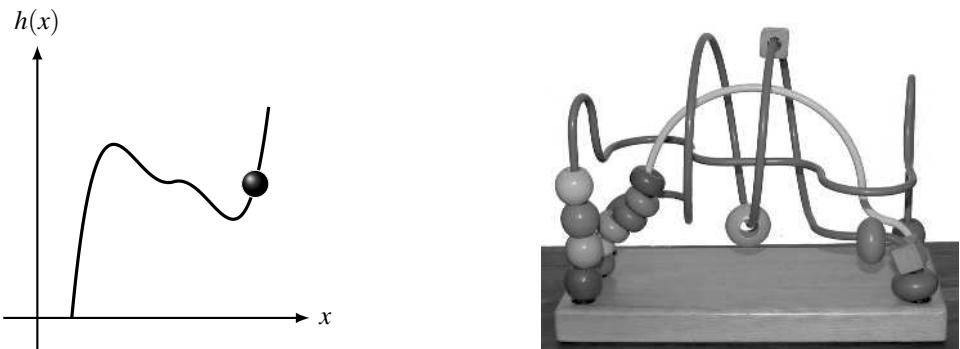


Figure 16.4 – Jeu constitué d'une perle enfilée sur une tige rigide. Le profil d'altitude $h(x)$ coïncide, à un facteur d'échelle près, au profil d'énergie potentielle.

6.2 Position du problème

Pour simplifier l'analyse, on considère que le point matériel M de masse m n'est soumis qu'à une force conservative notée $\vec{F}$ qui dérive de l'énergie potentielle E_p . On suppose que le problème est paramétré par une variable d'espace cartésienne notée x . La force $\vec{F}$ est alors de la forme $\vec{F} = F_x(x)\vec{u}_x$ et l'énergie potentielle une fonction de la variable x : $E_p(x)$.

D'après le théorème de l'énergie cinétique, l'énergie mécanique se conserve soit :

$$E_c + E_p = E_m = \text{constante.}$$

Par ailleurs, l'énergie cinétique est une quantité positive : $E_c = \frac{1}{2}mv^2 = E_m - E_p \geq 0$, ce qui entraîne que l'énergie mécanique est supérieure ou égale à l'énergie potentielle :

$$E_m \geq E_p.$$

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

On peut généraliser cette situation au cas où toutes les forces qui travaillent sont conservatives. $\vec{F}$ est alors la résultante des forces et E_p l'énergie potentielle totale. Par ailleurs, les résultats que l'on démontre ici sont transposables au cas où la variable est une variable angulaire qui sera traitée en exercice.

6.3 Analyse du mouvement à l'aide d'un graphe énergétique

La connaissance de l'énergie potentielle $E_p(x)$ en fonction de la variable de position x permet de distinguer les différents mouvements possibles. Pour cela, on trace le graphe de l'énergie potentielle $E_p(x)$ et la droite horizontale d'ordonnée E_m .

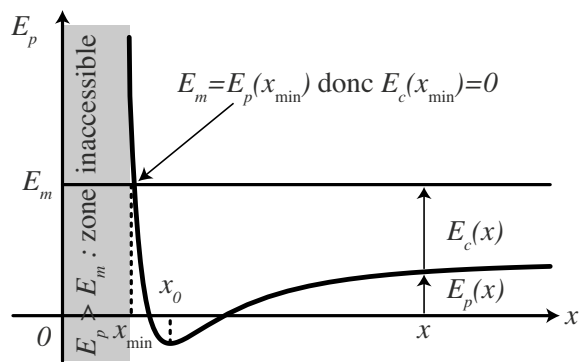


Figure 16.5 – Graphe d'énergie potentielle. Avec l'énergie mécanique E_m , le mobile peut atteindre l'ensemble des positions telles que $x > x_{\min}$. En $x = x_{\min}$, sa vitesse s'annule, en $x = x_0$, elle est maximale.

a) Positions accessibles

Graphiquement, l'inégalité $E_m \geq E_p(x)$ entraîne que le mobile ne peut accéder qu'aux positions pour lesquelles la courbe d'énergie potentielle est en dessous de la droite horizontale d'ordonnée E_m .

b) Vitesse en un point

Graphiquement, l'égalité $E_m = E_c + E_p$ permet de déterminer l'énergie cinétique E_c en un point d'abscisse x : c'est l'écart entre la droite d'ordonnée E_m et la courbe $E_p(x)$. On en déduit la norme de la vitesse $v(x)$ par la relation :

$$E_c = \frac{1}{2}mv^2 \quad \Longleftrightarrow \quad v(x) = \sqrt{\frac{2E_c}{m}} = \sqrt{\frac{2(E_m - E_p(x))}{m}}.$$

En un point où $E_p = E_m$, l'énergie cinétique est nulle. Cette position représente un point de vitesse nulle. C'est le cas en $x = x_{\min}$ sur la figure 16.5.

En un point où l'énergie potentielle est minimale, l'énergie cinétique du mobile est maximale. En ce point, le mobile atteint sa vitesse maximale. C'est le cas en $x = x_0$ sur la figure 16.5.

Exemple

Considérons l'énergie potentielle décrite par la figure 16.6 ci-dessous : $E_p(x)$ est une fonction décroissante puis croissante de x qui admet un minimum en x_0 et tend vers E_∞ lorsque $x \rightarrow \infty$. On veut établir graphiquement les différents types de mouvements que l'on peut observer suivant la valeur de l'énergie mécanique du mobile.

- Si $E_m < E_0$, alors $E_m < E_p$ quelle que soit la position. Aucun mouvement n'est possible avec une énergie mécanique aussi basse.
- Si $E_m = E_0$, le point matériel ne peut être qu'en $x = x_0$ et son énergie cinétique (donc sa vitesse) est nulle. Le point matériel est en équilibre en $x = x_0$.
- Si $E_0 < E_m < E_\infty$ ($E_m = E_1$ dans l'exemple considéré) alors le point matériel ne peut évoluer qu'entre les positions x_1 et x'_1 . Il reste dans une zone bornée de l'espace et ne peut pas s'en aller à l'infini. Il est dans un **état lié**. Lorsqu'il passe en x_0 , son énergie potentielle est minimale, sa vitesse est alors maximale. En x_1 et x'_1 , la vitesse du mobile s'annule, il fait demi-tour et repart dans l'autre sens.
- Si $E_m \geq E_\infty$ ($E_m = E_2$ dans l'exemple considéré) alors le point matériel peut évoluer aux positions d'abscisses x vérifiant $x \geq x_3$. Le point matériel peut s'échapper à l'infini. Il est dans un **état libre** ou **état de diffusion**. x_3 est une position de vitesse nulle et x_0 une position de vitesse maximale.

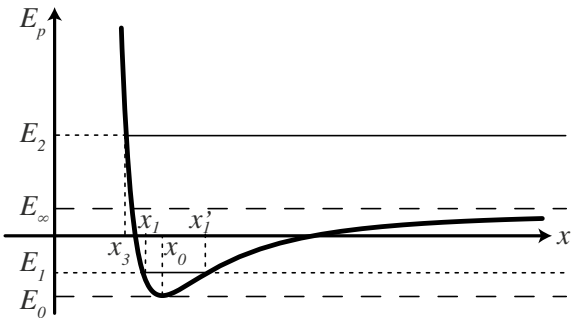


Figure 16.6 – Analyse graphique des positions accessibles au mobile en fonction de son énergie mécanique.

6.4 Analyse des équilibres à l'aide d'un graphe énergétique

a) Définitions

Un point matériel est en équilibre lorsque sa vitesse est nulle à tout instant. On en déduit que son accélération est nulle à tout instant puis, d'après le principe fondamental de la dynamique, que la somme des forces qui s'appliquent sur ce point est nulle.

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

En une position d'équilibre, la somme des forces qui s'appliquent au mobile est nulle :

$$\sum_i \vec{f}_i = \vec{0}.$$

Le mobile peut rester indéfiniment en une position d'équilibre.

Quand on écarte le mobile de sa position d'équilibre,

- s'il revient vers sa position d'équilibre initiale, la position d'équilibre est **stable** ;
- s'il s'éloigne définitivement de sa position d'équilibre initiale, la position d'équilibre est **instable**.

Pour qu'une position d'équilibre soit stable, il est donc nécessaire que, au voisinage de cette position d'équilibre, la résultante des forces soit dirigée vers cette dernière.

b) Lien entre la force et l'énergie potentielle

Lorsqu'on considère uniquement la force $\vec{F} = F_x(x)\vec{u}_x$ qui dérive de l'énergie potentielle $E_p(x)$, le travail s'écrit :

$$\delta W = \vec{F} \cdot d\vec{OM} = F_x(x)\vec{u}_x \cdot d\vec{OM} = F_x(x)dx = -dE_p$$

puis la force :

$$F_x(x) = -\frac{dE_p}{dx}.$$

Cette relation explique pourquoi l'on dit que la force $\vec{F}$ dérive d'une énergie potentielle. Cette relation montre également que :

- aux endroits où E_p est croissante, sa dérivée par rapport à x est positive et la force est dirigée vers les x décroissantes et donc vers les énergies potentielles décroissantes,
- aux endroits où E_p est décroissante, sa dérivée par rapport à x est négative et la force est dirigée vers les x croissantes et donc vers les énergies potentielles décroissantes,
- aux endroits où E_p admet une tangente horizontale, sa dérivée par rapport à x est nulle et la force également.

Ceci est illustré sur la figure 16.7.

Si la force $\vec{F} = F_x(x)\vec{u}_x$ dérive de l'énergie potentielle $E_p(x)$ alors :

$$\vec{F} = -\frac{dE_p}{dx}\vec{u}_x.$$

En un point d'abscisse x , la force $\vec{F}$ est dirigée dans le sens des énergies potentielles décroissantes donc vers le minimum d'énergie potentielle le plus proche.

c) Détermination énergétique des positions d'équilibre

Pour déterminer les positions pour lesquelles le point matériel est à l'équilibre, il faut chercher les positions où la force $\vec{F}$ s'annule. Cela implique que la dérivée de la fonction $E_p(x)$ s'annule et donc que la courbe représentative de $E_p(x)$ admet une tangente horizontale.

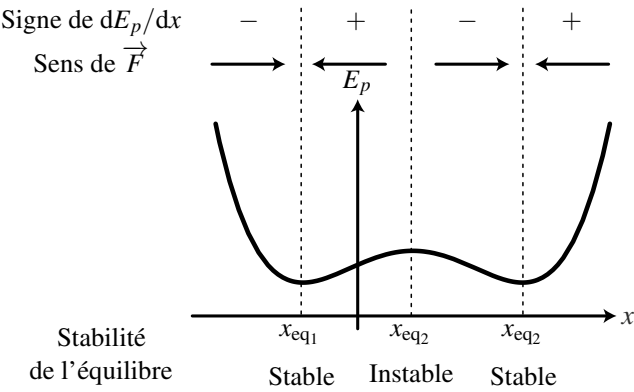


Figure 16.7 – Sens de la force $\vec{F}$ en fonction des variations de l'énergie potentielle E_p .

Une position d'équilibre correspond à un extremum d'énergie potentielle. Son abscisse x_{eq} vérifie l'équation :

$$\left(\frac{dE_p}{dx}\right)_{x=x_{eq}} = 0.$$

d) Étude énergétique de la stabilité des équilibres

Pour analyser la stabilité de la position d'équilibre d'abscisse x_{eq} , il faut déterminer si, au voisinage de x_{eq} , la résultante des forces est dirigée vers x_{eq} ou non.

On se place alors au voisinage de x_{eq} en un point d'abscisse $x = x_{eq} + dx$. La force $\vec{F}$ est dirigée vers les énergies potentielles décroissantes donc :

- si x_{eq} est un minimum d'énergie potentielle, $\vec{F}$ tend à ramener le mobile vers x_{eq} qui est donc une position d'équilibre stable (positions x_{eq1} et x_{eq3} de la figure 16.7),
- si x_{eq} est un maximum d'énergie potentielle, $\vec{F}$ tend à éloigner le mobile de x_{eq} qui est donc une position d'équilibre instable (positions x_{eq2} de la figure 16.7).

Lorsque la fonction E_p est deux fois dérivable et que sa dérivée seconde ne s'annule pas en x_{eq} , on peut exprimer la force $\vec{F}$ au voisinage de x_{eq} . Pour cela, on utilise un développement de Taylor (voir appendice mathématique) de la fonction $F_x(x)$:

$$F_x(x_{eq} + dx) = F_x(x_{eq}) + \left(\frac{dF_x}{dx}\right)_{x=x_{eq}} dx = \left(\frac{dF_x}{dx}\right)_{x=x_{eq}} dx = dF_x,$$

puisque x_{eq} est une position d'équilibre donc $F_x(x_{eq}) = 0$.

Par ailleurs, $F_x(x) = -\frac{dE_p}{dx}$ donc $\left(\frac{dF_x}{dx}\right)_{x=x_{eq}} = -\left(\frac{d^2E_p}{dx^2}\right)_{x=x_{eq}}$. Au voisinage de x_{eq} , le

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

mobile est donc soumis à une force élémentaire dF_x telle que

$$dF_x = - \left(\frac{d^2 E_p}{dx^2} \right)_{x=x_{eq}} dx.$$

Si on exerce un déplacement $dx > 0$, il faut que $dF_x < 0$ pour ramener le point vers sa position initiale. Cela impose donc :

$$\left(\frac{d^2 E_p}{dx^2} \right)_{x=x_{eq}} > 0.$$

De même, si on exerce un déplacement $dx < 0$, il faut que $dF_x > 0$ pour ramener le point vers sa position initiale et on retrouve le même résultat.

La position d'équilibre x_{eq} est **stable** si l'énergie potentielle est minimale en x_{eq} . Cela se traduit généralement par :

$$\left(\frac{d^2 E_p}{dx^2} \right)_{x=x_{eq}} > 0.$$

La position d'équilibre x_{eq} est **instable** dans les autres cas. Généralement, l'énergie potentielle est alors maximale en x_{eq} et $\left(\frac{d^2 E_p}{dx^2} \right)_{x=x_{eq}} < 0$.

7 Portraits de phase et lien avec le profil d'énergie potentielle

Dans ce paragraphe, on définit une représentation graphique du mouvement appelée **portrait de phase**. On se limite à un **système à un degré de liberté** mécanique, c'est-à-dire à un système dont l'évolution est décrite par un seul paramètre de position. Pour un tel mouvement, la notion de trajectoire est souvent de peu d'intérêt car les trajectoires sont des portions de droites, de cercles, ou de courbes connues à l'avance. Le portrait de phase permet de représenter sur un même graphe la position de M ainsi que sa vitesse et permet une représentation plus riche. Comme dans le paragraphe précédent, on considère que le mouvement du point matériel M est paramétré par une variable d'espace cartésienne notée x pour simplifier l'analyse, mais tout est transposable au cas où la coordonnée de M n'est pas cartésienne.

7.1 Définitions

Le **plan de phase** d'un système à un degré de liberté est le plan $(x, \dot{x})$ où x est la variable de position de M et $\dot{x}$ sa dérivée. À chaque instant, le mobile M est repéré par sa position x et on peut déterminer sa vitesse $\dot{x}$. On peut le représenter par un point P de coordonnées $(x, \dot{x})$ du plan de phase. Au cours du mouvement du mobile, la succession des points P trace une courbe dans le plan de phase appelée **trajectoire de phase**. On a déjà établi et tracé une trajectoire de phase pour l'oscillateur harmonique, on va généraliser cette notion pour d'autres mouvements. On peut noter que les conditions initiales du mouvement ($x(t=0), \dot{x}(t=0)$) sont représentées par un point P_0 du plan de phase. Le déterminisme implique que chaque condition initiale donne une unique trajectoire de phase. Le **portrait de phase** d'un système est un ensemble de trajectoires de phases associées à différentes conditions initiales.

7.2 Exemple introductif

a) Position du problème

On considère un mobile M de masse m soumis à la seule force conservative $\vec{F} = F_x(x)\vec{u}_x$ qui dérive de l'énergie potentielle $E_p(x)$ tracée sur la figure 16.8. Comme la seule force qui travaille est conservative, le théorème de l'énergie cinétique implique que l'énergie mécanique E_m du système est conservée : $E_m = E_c + E_p = \frac{1}{2}mv^2 + E_p(x)$.

En coordonnées cartésiennes $\vec{v} = \dot{x}\vec{u}_x$ et l'énergie cinétique vaut alors :

$$E_c = \frac{1}{2}mv^2 = \frac{1}{2}m\dot{x}^2.$$

Il s'en suit que l'on peut déterminer $\dot{x}$ en fonction x à partir de l'expression de l'énergie potentielle $E_p(x)$ et de la valeur de l'énergie mécanique E_m grâce à la relation :

$$\dot{x}^2 = \frac{2(E_m - E_p(x))}{m}$$

qui admet deux solutions de signes opposés :

$$\dot{x} = \sqrt{\frac{2(E_m - E_p(x))}{m}} > 0 \quad \text{et} \quad \dot{x} = -\sqrt{\frac{2(E_m - E_p(x))}{m}} < 0.$$

On peut alors tracer une trajectoire de phase point par point en fixant une valeur pour E_m et en relevant la valeur de $E_p(x)$ sur le profil d'énergie potentielle. La fonction « racine carré » étant une fonction croissante, la norme de la vitesse $|\dot{x}|$ est une fonction croissante de l'écart $E_m - E_p(x)$.

b) Observation du portrait de phase

On a tracé sur la figure 16.8 un exemple de profil d'énergie potentielle et le portrait de phase associé. L'énergie potentielle $E_p(x)$ est une fonction décroissante puis croissante de x qui admet un minimum en O . Sa courbe représentative dessine une « cuvette » dont les bords ont une hauteur E_∞ . Elle forme un **puits de potentiel** de profondeur E_∞ . L'origine de l'axe (Ox), ainsi que la référence d'énergie potentielle sont choisies arbitrairement au minimum d'énergie potentielle. Sur cette courbe d'énergie potentielle, on a fait apparaître trois niveaux d'énergie mécanique E_1 , E_2 et E_3 et tracé sur le graphe du dessous les trajectoires de phase associées.

c) Analyse des différents mouvements possibles

Dans tous les cas, on observe que la norme de la vitesse augmente quand le mobile s'approche du fond du puits puis diminue quand il s'en éloigne. Elle est maximale au fond du puits.

Pour l'énergie mécanique $E_1 \ll E_\infty$ ou $E_2 < E_\infty$, le mobile est piégé dans le puits de potentiel. Il est dans un état lié et suit un mouvement périodique. Sa trajectoire de phase est une courbe fermée décrite dans le sens horaire. Lorsque le mobile passe par le minimum d'énergie potentielle, sa vitesse est maximale. La trajectoire de phase admet alors une tangente horizontale.

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

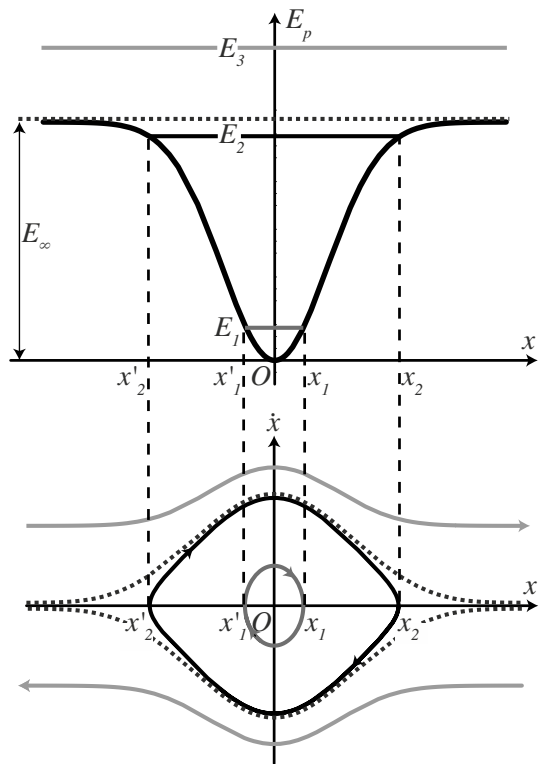


Figure 16.8 – Energie potentielle et portrait de phase. Exemple d'un puits de potentiel de profondeur E_0 .

Pour $E_m = E_1$, la vitesse du mobile s'annule en x_1 et x'_1 . La trajectoire de phase coupe l'axe des abscisses en x_1 et x'_1 . Le mobile réalise des oscillations de faible amplitude autour de sa position d'équilibre. On démontrera qu'on peut assimiler ces oscillations à des oscillations harmoniques et la trajectoire de phase à une ellipse.

Pour $E_m = E_2$, la vitesse du mobile s'annule en x_2 et x'_2 . La trajectoire de phase coupe l'axe des abscisses en x_2 et x'_2 . Le mobile réalise des oscillations de grande amplitude qui ne sont pas harmoniques. Sa trajectoire de phase change de forme. On montrera que ce changement de forme est une preuve de la non-linéarité du problème.

Pour l'énergie mécanique $E_3 > E_\infty$, le mobile est libre de sortir du puits de potentiel. Il est dans un état de diffusion. Deux trajectoires de phase ouvertes sont possibles selon le signe de la vitesse : une trajectoire parcourue de gauche à droite située dans le demi-plan supérieur qui correspond à $\dot{x} > 0$ et une trajectoire parcourue de droite à gauche située dans le demi-plan inférieur qui correspond à $\dot{x} < 0$.

7.3 Caractéristiques principales des portraits de phase

Dans ce paragraphe, on cherche les caractéristiques principales des portraits de phase à partir de l'exemple ci-dessus et on les justifie brièvement.

1. **Les trajectoires de phases sont parcourues de gauche à droite dans le demi-plan supérieur et de droite à gauche dans le demi-plan inférieur.**

Dans le demi-plan supérieur $\dot{x} > 0$ donc x est croissante et évolue de gauche à droite.
Dans le demi-plan inférieur, $\dot{x} < 0$ donc x est décroissante et évolue de droite à gauche.

2. **Les mouvements périodiques correspondent à des trajectoires de phase fermées décrites dans le sens horaire.**

Si le mouvement du mobile est périodique, il retrouve la même position et la même vitesse après une période donc le même point du plan de phase. D'après le point 1, elles sont décrites dans le sens horaire.

3. **Les trajectoires de phases ne se croisent pas.**

Ceci se démontre par l'absurde en supposant que deux trajectoires se croisent en P_0 . P_0 représente une condition initiale pour l'évolution du système. Si deux trajectoires de phases passent par P_0 , c'est qu'une même condition initiale peut conduire à deux mouvements différents. C'est incompatible avec le déterminisme classique.

4. **Une trajectoire de phase coupe généralement l'axe (Ox) selon la verticale.**

En effet, la tangente à la trajectoire de phase en un point $(x, \dot{x})$ peut être définie comme la droite passant par ce point et faisant un angle α avec l'axe Ox tel que :

$$\tan \alpha = \frac{d\dot{x}}{dx} = \frac{\frac{d\dot{x}}{dt}}{\frac{dx}{dt}} = \frac{\ddot{x}}{\dot{x}}.$$

Sur l'axe (Ox) , la vitesse $\dot{x}$ est nulle. Dans le cas général, l'accélération $\ddot{x}$ est non nulle et $\tan \alpha \rightarrow \infty$ d'où $|\alpha| \rightarrow \frac{\pi}{2}$.

On peut noter que les positions d'équilibres pour lesquelles $\ddot{x} = 0$ sont des exceptions à cette règle. En pratique, seules les positions d'équilibre instables posent problème car, si le mobile est sans vitesse en une position d'équilibre stable, il ne bouge plus et sa trajectoire de phase est restreinte à un point.

5. **Les positions d'équilibres stables sont entourées de trajectoires de phase elliptiques.**

On a vu que les trajectoires de phases d'un oscillateur harmonique sont des ellipses. On démontrera dans un prochain chapitre que les mouvements de faible amplitude au voisinage d'une position d'équilibre stable sont des mouvements harmoniques. Ce point en sera la conséquence directe.

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

SYNTHÈSE

SAVOIRS

- puissance et travail d'une force
- savoir que la puissance dépend du référentiel
- énergie cinétique d'un système
- savoir qu'elle dépend du référentiel
- loi de l'énergie cinétique et de la puissance cinétique dans un référentiel galiléen
- définition d'une force conservative
- expressions des énergies potentielles de pesanteur (champ uniforme), gravitationnelle (champ créé par un astre ponctuel), élastique, électrique (champ uniforme et champ créé par une charge ponctuelle)
- définition d'une position d'équilibre
- savoir que les positions d'équilibre stable correspondent à des minima d'énergie potentielle et que les positions d'équilibre instable correspondent à des maxima
- définitions d'un état lié et d'un état de diffusion

SAVOIR-FAIRE

- reconnaître le caractère moteur ou résistant d'une force
- utiliser la loi de l'énergie cinétique ou la loi de la puissance cinétique à bon escient
- établir les expressions des énergies potentielles à connaître
- reconnaître les cas de conservation de l'énergie mécanique
- déterminer l'énergie mécanique initiale en exploitant les conditions initiales
- tracer un graphe d'énergie potentielle
- repérer les positions d'équilibre et déduire leur stabilité à partir d'un graphe d'énergie potentielle
- déduire d'un graphe d'énergie potentielle le comportement qualitatif d'un mobile : trajectoire bornée ou non, mouvement périodique, position de vitesse nulle
- expliquer qualitativement le lien entre le profil d'énergie potentielle et le portrait de phase

MOTS-CLÉS

- | | | |
|----------------------|--------------------------|---------------------|
| • puissance | • énergie potentielle | • état lié |
| • travail | • énergie mécanique | • état de diffusion |
| • énergie cinétique | • mouvement conservatif | |
| • force conservative | • équilibre et stabilité | |

S'ENTRAÎNER

16.1 Distance de freinage (★)

Une voiture de masse $m = 1,5 \cdot 10^3$ kg roule à la vitesse de $50 \text{ km} \cdot \text{h}^{-1}$ sur une route horizontale. Devant un imprévu, le conducteur écrase la pédale de frein et s'arrête sur une distance $d = 15$ m. On modélise la force de freinage par une force constante opposée à la vitesse.

1. Calculer le travail de la force de freinage.
2. En déduire la norme de cette force.
3. Quelle distance faut-il pour s'arrêter si la vitesse initiale est de $70 \text{ km} \cdot \text{h}^{-1}$?
4. Commenter cette phrase relevée dans un livret d'apprentissage de la conduite : « La distance de freinage est proportionnelle au carré de la vitesse du mobile ».

16.2 Le Marsupilami (★)

Le Marsupilami est un animal de bande dessinée créé par Franquin aux capacités physiques remarquables, en particulier grâce à sa queue qui possède une force importante. Pour se déplacer, le Marsupilami enroule sa queue comme un ressort entre lui et le sol et s'en sert pour se propulser vers le haut.

On note $\ell_0 = 2$ m la longueur à vide du ressort équivalent. Lorsqu'il est complètement comprimé, la longueur du ressort est $\ell_m = 50$ cm. La masse m de l'animal est 50 kg et la queue quitte le sol lorsque le ressort mesure ℓ_0 . On prendra $g = 10 \text{ m} \cdot \text{s}^{-2}$.

1. Quelle est la constante de raideur du ressort équivalent si la hauteur maximale d'un saut est $h = 10$ m ? Quelle est sa vitesse lorsque la queue quitte le sol ?

16.3 Interaction entre particules chargées (★)

On considère deux particules A (fixe) et B (mobile), de même masse m et de charge respective q_A et q_B . On considère la force de Coulomb entre ces deux particules comme étant la seule force en jeu dans ce problème.

1. Rappeler l'expression de la force de Coulomb notée $\vec{f}$.
2. Déterminer l'énergie potentielle dont dérive la force $\vec{f}$.
3. On suppose $q_A = q_B = q$. On lance B vers A avec la vitesse v_0 . à quelle distance minimale B s'approche-t-elle de A ? On pourra s'aider d'un graphe d'énergie potentielle.
4. On suppose $q_A = -q_B = q$. Quelle vitesse minimale faut-il donner à B pour qu'elle puisse s'échapper à l'infini ? On pourra s'aider d'un graphe d'énergie potentielle.

16.4 Toboggan (★★)

Un adulte ($m = 70$ kg) descend un toboggan d'une hauteur $h = 5$ m faisant un angle $\alpha = 45^\circ$ avec le sol. La norme de la force de frottement $\vec{T}$ est donnée par $\|\vec{T}\| = f \|\vec{R}\|$, où $f = 0,4$ est le coefficient de frottement et $\vec{R}$ la réaction normale. On prendra $g = 9,8 \text{ m} \cdot \text{s}^{-2}$

1. Calculer la variation d'énergie mécanique due au frottement entre le haut et la bas du toboggan.

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

2. Déterminer la vitesse de la personne en bas du toboggan. Comparer la à celle qu'il aurait s'il n'y avait pas de frottement.

APPROFONDIR

16.5 Cycliste au Tour de France d'après Centrale 2006 (★★)

Un cycliste assimilé à un point matériel se déplace en ligne droite. Il fournit une puissance mécanique constante P , les forces de frottement de l'air sont proportionnelles au carré de la vitesse v du cycliste ($\vec{F}_f = -kv\vec{v}$) où k est une constante positive. On néglige les forces de frottement du sol sur la roue et on choisit un axe horizontal (Ox) orienté dans la direction du mouvement du cycliste.

1. En appliquant le théorème de la puissance cinétique, établir une équation différentielle en v et montrer qu'on peut la mettre sous la forme :

$$mv^2 \frac{dv}{dx} = k(v_l^3 - v^3),$$

où v_l est une constante homogène à une vitesse dont on cherchera la signification physique.

2. On pose $f(x) = k(v_l^3 - v^3)$.

a. Dédire des résultats précédent, l'équation différentielle vérifiée par f .

b. Déterminer l'expression de la vitesse en fonction de x , s'il aborde la ligne droite avec une vitesse v_0 .

c. Application numérique : lors d'un sprint, la puissance développée vaut $P = 2 \text{ kW}$ et la vitesse limite v_l vaut $20 \text{ m}\cdot\text{s}^{-1}$. Déterminer la valeur de k et en déduire la distance caractéristique pour qu'un coureur de masse $m = 85 \text{ kg}$ avec son vélo atteigne cette vitesse. Conclure.

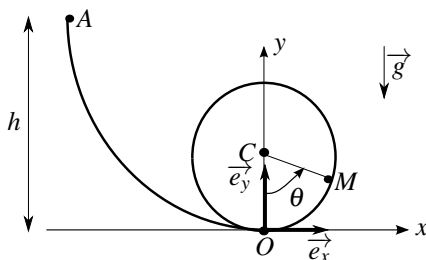
16.6 Bille dans une gouttière ICNA 2006 (★★)

Une bille, assimilée à un point matériel M de masse m , est lâchée sans vitesse initiale depuis le point A d'une gouttière situé à une hauteur h du point le plus bas O de la gouttière. Cette dernière est terminée en O par un guide circulaire de rayon a , disposé verticalement. La bille, dont on suppose que le mouvement a lieu sans frottement, peut éventuellement quitter la gouttière à l'intérieur du cercle. On note $\vec{g} = -g\vec{u}_y$ l'accélération de la pesanteur.

1. Calculer la norme v_0 de la vitesse en O puis en un point M quelconque du cercle repéré par l'angle θ .

2. Déterminer la réaction de la gouttière en un point du guide circulaire.

3. Déterminer la hauteur minimale de A pour que la bille ait un mouvement révolatif dans le guide.



16.7 Tige avec ressort (★)

On considère une tige fixe, dans un plan vertical xOz , faisant l'angle α avec l'axe Oz . Un anneau M de masse m est enfilé sur la tige et astreint à se déplacer sans frottement le long de celle-ci. Cet anneau est également relié à un ressort de raideur k et de longueur à vide ℓ_v dont l'autre extrémité est fixée en O . On repère la position de M par $OM = X$.

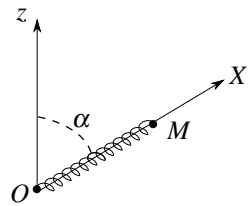


Figure 16.9

1. Quelles sont les forces conservatives appliquées à M ? Déterminer l'expression de leur énergie potentielle E_p en fonction de X et α .
2. Etablir l'équation différentielle du mouvement à l'aide du théorème de l'énergie mécanique.
3. On souhaite étudier graphiquement les différents mouvements possibles.
 - a. Etudier la fonction $E_p(X)$ dans le cas où $mg \cos(\alpha) < k\ell_v$. Tracer son allure.
 - b. Discuter sur le graphique les mouvements possibles en prenant à $t = 0$ les conditions initiales suivantes : $X = \ell_v$ et $\frac{dX}{dt} = V_0$. Préciser la valeur maximale de V_0 pour que le mouvement se fasse entre deux positions extrêmes $X_1 > 0$ et $X_2 > 0$.
 - c. Déterminer les fonctions $V(X)$ et $X(t)$ dans les conditions de la question précédente.

16.8 Pendule simple modifié d'après ENSTIM 2005 (★★)

On considère un pendule simple modifié (figure 16.10). Un mobile ponctuel M de masse m , est accroché à l'extrémité d'un fil inextensible de longueur L et de masse négligeable, dont l'autre extrémité est fixe en O . On néglige tout frottement. Lorsque $\theta > 0$, le système se comporte comme un pendule simple de centre O et de longueur de fil L . A la verticale et en dessous de O , un clou est planté en O' avec $OO' = L/3$, qui bloquera la partie haute du fil vers la gauche.

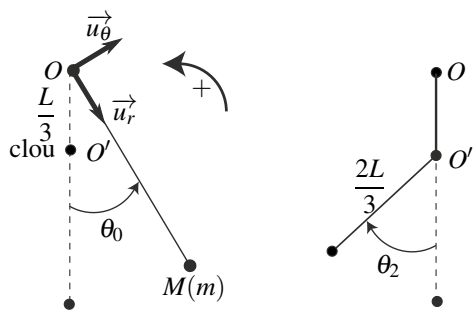


Figure 16.10

Quand $\theta < 0$, le système se comporte donc comme un pendule simple de centre O' et de longueur de fil $2L/3$. On repère la position du pendule par l'angle θ qu'il fait avec la verticale.

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

à la date $t = 0$, on abandonne sans vitesse initiale le mobile M en donnant au fil une inclinaison initiale $\theta(0) = \theta_0 > 0$. On note t_1 la date de la première rencontre du fil avec le clou, t_2 la date de première annulation de la vitesse du mobile pour $\theta < 0$. L'intervalle de date $[0, t_1[$ est nommé première phase du mouvement, l'intervalle $]t_1, t_2]$ est nommé deuxième phase. À la date t_1^- immédiatement inférieure à t_1 , le fil n'a pas encore touché le clou et à la date t_1^+ , immédiatement supérieure, le fil vient de toucher le clou.

1. Par le théorème de la puissance mécanique, établir l'équation différentielle vérifiée par θ pour la première phase du mouvement.
2. Dans l'hypothèse des petites oscillations, on suppose $\sin\theta \simeq \theta$. Déterminer la durée δt_I de la première phase du mouvement sans résoudre l'équation.
3. En utilisant le théorème de l'énergie mécanique, déterminer la vitesse v_1^- de M à la date t_1^- . En déduire la vitesse angulaire ω_1^- à cette date.
4. Le blocage de la partie supérieure du fil par le clou ne s'accompagne d'aucun transfert énergétique. déterminer la vitesse v_1^+ de M à la date t_1^+ . En déduire la vitesse angulaire ω_1^+ à cette date.
5. En utilisant les résultats des questions (1 et 2), donner sans calcul la durée δt_{II} de la deuxième phase.
6. Déterminer l'expression de l'angle θ_2 à la date t_2 .
7. Décrire brièvement la suite du mouvement de ce système et donner l'expression de sa période T .
8. Dresser l'allure du portrait de phase, dans le système d'axes $\left(\theta, \frac{d\theta}{dt}\right)$. Conclure.

CORRIGÉS

16.1 Distance de freinage

On étudie le mouvement du véhicule, assimilé à son centre de gravité M de masse m dans le référentiel terrestre galiléen. Ce véhicule est soumis à son poids, à la réaction normale de la route et à la force de freinage $\vec{f}$. La réaction normale de la route et le poids sont dirigés selon la verticale alors que le déplacement du véhicule est horizontal. Leur puissance est nulle. Seule la force $\vec{f}$ fournit un travail.

1. On applique le théorème de l'énergie cinétique entre l'instant initial où le véhicule à la vitesse $v_i = 50 \text{ km}\cdot\text{h}^{-1}$ et l'instant final où il s'arrête (vitesse $v_f = 0 \text{ km}\cdot\text{h}^{-1}$) :

$$E_{cf} - E_{ci} = \frac{1}{2}mv_f^2 - \frac{1}{2}mv_i^2 = W(\vec{f}),$$

soit :

$$W(\vec{f}) = -\frac{1}{2}mv_i^2 = -0,5 \times 1,5 \cdot 10^3 \times 50/3,6 = -10 \text{ kJ}.$$

Le travail de $\vec{f}$ est négatif car la force est résistante.

2. Par ailleurs, en désignant par I le point initial atteint à l'instant t_i et F le point final atteint à t_f , on peut déterminer le travail de cette force de freinage de norme constante :

$$W(\vec{f}) = \int_{t_i}^{t_f} \vec{f} \cdot \vec{v} dt = \int_{t_i}^{t_f} -fv dt = \int_{IF} -f d\ell = -fd$$

d'où :

$$f = -\frac{W(\vec{f})}{d} = \frac{1}{2}mv_i^2/d = 0,69 \text{ kN}.$$

3. La norme de la force étant constante, la distance d'arrêt est proportionnelle au travail de $\vec{f}$ donc à l'énergie cinétique initiale à dissiper. La distance d'arrêt est donc proportionnelle au carré de la vitesse initiale. Pour une vitesse initiale $v'_i = 70 \text{ km}\cdot\text{h}^{-1}$, la distance d'arrêt vaut donc :

$$d' = d \left(\frac{v'_i}{v_i} \right)^2 = 15 \times \left(\frac{70}{50} \right)^2 = 29 \text{ m}.$$

Elle a quasiment doublée.

4. La phrase trouvée dans le livret d'apprentissage de conduite est parfaitement justifiée par ce modèle physique simple.

16.2 Le Marsupilami

Le référentiel lié au sol est considéré galiléen. Le système est le corps du Marsupilami assimilé à son centre de gravité M de masse m auquel s'applique trois forces lorsqu'il est sur le sol :

- le poids (force conservative);

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

- la tension de la queue-ressort (force conservative) ;
- la réaction du sol (force de travail nul).

Lorsqu'il est en l'air et que la queue ne touche plus le sol, seul le poids est à prendre en compte.

1. Le système n'échange de l'énergie que par l'intermédiaire de forces conservatives. On en déduit que l'énergie mécanique se conserve au cours du mouvement. On définit l'axe (Oz) vertical ascendant dont l'origine est située au niveau du sol sorte que l'altitude de M correspond à son abscisse z et que la longueur de la queue de l'animal est également égale à z lorsque la queue touche le sol.

L'énergie potentielle de pesanteur vaut $E_{p1} = mg \text{ altitude} = mgz$.

L'énergie potentielle élastique vaut $E_{p2} = \frac{1}{2}k(\ell - \ell_0)^2 = \frac{1}{2}k(z - \ell_0)^2$.

L'expression de l'énergie mécanique est alors :

$$E_m = \frac{1}{2}mv^2 + mgz + \frac{1}{2}k(z - \ell_0)^2 \quad \text{pendant que la queue touche le sol}$$

et :

$$E_m = \frac{1}{2}mv^2 + mgz \quad \text{lorsque l'animal est en l'air.}$$

On écrit l'égalité de l'énergie mécanique en trois points :

1. Départ : $z = \ell_m, v = 0$.
2. Décollage : $z = \ell_0, v_D$.
3. Point le plus haut : $z = h, v = 0$.

Soit :

$$mg\ell_m + \frac{1}{2}k(\ell_m - \ell_0)^2 = \frac{1}{2}mv_D^2 + mg\ell_0 = mgh$$

On en déduit :

$$k = \frac{2mg(h - \ell_m)}{(\ell_m - \ell_0)^2} = 4,1 \cdot 10^3 \text{ N} \cdot \text{m}^{-1}$$

ce qui est assez important.

La vitesse est au départ vaut alors $v_D = \sqrt{2g(h - \ell_0)}$ soit $12,5 \text{ m} \cdot \text{s}^{-1}$, ce qui est aussi une valeur importante.

16.3 Interactions entre particules chargées

1. Le point A étant fixe on le choisit comme origine du repère et on note $r\vec{u}_r$ le vecteur $\overrightarrow{AB}$ où $r = AB$ et $\vec{u}_r$ est le vecteur unitaire dirigé de A vers B . La force $\vec{f}$ s'écrit alors : $\vec{f} = \frac{q_A q_B}{4\pi\epsilon_0 r^2} \vec{u}_r$.

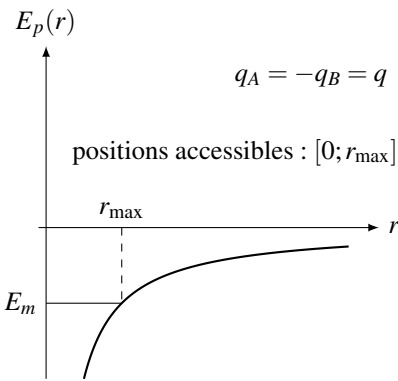
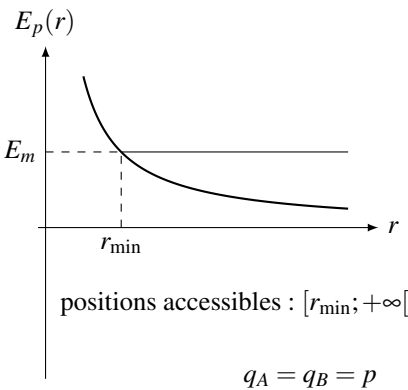
2. Pour déterminer l'énergie potentielle, on utilise la relation entre le travail et la variation d'énergie potentielle valable lorsque les forces sont conservatives : $\delta W = -dE_p$, soit :

$$\delta W = \vec{f} \cdot d\vec{r} = \frac{q_A q_B}{4\pi\epsilon_0} \frac{\vec{u}_r}{r^2} \cdot (dr\vec{u}_r + r d\vec{u}_r)$$

or $\vec{u}_r \cdot d\vec{u}_r = d(\vec{u}_r \cdot \vec{u}_r) = d(1) = 0$ car $\vec{u}_r$ est un vecteur unitaire. On peut ainsi déterminer E_p :

$$\delta W = -dE_p = \frac{q_A q_B}{4\pi\epsilon_0 r^2} dr \Rightarrow E_p = \frac{q_A q_B}{4\pi\epsilon_0 r} + \text{constante}.$$

On choisit la référence d'énergie potentielle nulle à l'infini, ce qui fixe la constante à zéro.



3. Puisque $\vec{f}$ est la seule force à prendre en compte et qu'elle est conservative, l'énergie mécanique se conserve. En considérant A immobile et $q_A = q_B = q$ on peut écrire :

$$E_m = E_c + E_p = \frac{1}{2}mv_B^2 + \frac{q^2}{4\pi\epsilon_0 r}.$$

Le point B s'approche de A, comme la force est répulsive, sa vitesse diminue puis s'annule, puis B s'éloigne de A. La distance minimale entre les deux points est obtenue pour $v_B = 0$. On écrit l'égalité de l'énergie mécanique entre le point de départ ($AB = r = a$) et le point où $v_B = 0$:

$$\frac{1}{2}mv_0^2 + \frac{q^2}{4\pi\epsilon_0 a} = \frac{q^2}{4\pi\epsilon_0 d_{\min}} \Rightarrow d_{\min} = \left(\frac{1}{a} + \frac{2\pi\epsilon_0 mv_0^2}{q^2} \right)^{-1}.$$

4. Avec $q_A = -q_B = q$, l'énergie mécanique s'écrit :

$$E_m = E_c + E_p = \frac{1}{2}mv_B^2 - \frac{q^2}{4\pi\epsilon_0 r}.$$

La force est maintenant attractive. Il faut lancer B à l'opposé de A pour qu'elle puisse s'éloigner. La vitesse de B diminue, et la charge ne peut pas accéder aux points plus éloignés que $r_{\max}$ (voir figure ci-dessus). Pour que B puisse atteindre l'infini, il faut que $r_{\max}$ soit repoussé à l'infini. Pour cela, l'énergie mécanique doit être suffisante pour que la trajectoire ne soit pas bornée et on voit sur le graphe d'énergie potentielle qu'il est nécessaire que l'énergie mécanique soit positive ou nulle. La vitesse minimale à donner à B correspond à l'énergie cinétique minimale et donc à l'énergie mécanique minimale $E_m = 0$. La conservation de l'énergie mécanique donne alors :

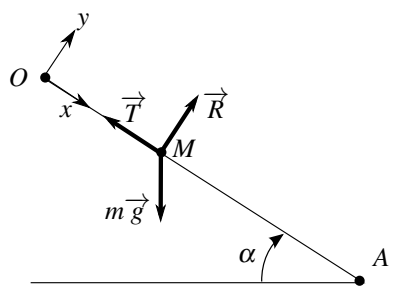
$$\frac{1}{2}mv_0^2 - \frac{q^2}{4\pi\epsilon_0 a} = E_m(r \rightarrow \infty) = 0 \Rightarrow v_0 = \sqrt{\frac{q^2}{2\pi\epsilon_0 am}}.$$

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

Cette vitesse porte le nom de vitesse de libération car la charge B échappe alors à l'attraction de la charge A . On peut remarquer que, dans ces conditions, la vitesse s'annule en l'infini.

16.4 Toboggan

Le référentiel terrestre est supposé galiléen. Le système M subit trois forces : le poids, la force de frottement et la réaction normale. Le poids est une force conservative, la réaction normale ne travaille pas et la force de frottement est une force non conservative.



1. La variation d'énergie mécanique étant égale au travail des forces non conservatives, on obtient entre le point de départ O et le point d'arrivée A :

$$E_m(A) - E_m(O) = W_{O \rightarrow A}(\vec{T}).$$

Il faut donc calculer la force $\vec{T}$. La relation fondamentale de la dynamique en projection sur Oy donne $m\ddot{y} = R - mg \sin \alpha$, or $\ddot{y} = 0$ donc $R = mg \sin \alpha$. On en déduit $\vec{T} = -fmg \sin \alpha \vec{u}_x$. Ainsi :

$$E_m(A) - E_m(O) = \int_0^{x_A} \vec{T} \cdot (dx \vec{u}_x) = -fmgL \sin \alpha.$$

Sachant que L (distance entre O et A) vaut $h / \cos \alpha$, on trouve finalement :

$$E_m(A) - E_m(O) = -fmgh \tan \alpha = -fmgh.$$

2. En explicitant l'expression de l'énergie mécanique, on obtient :

$$\left(\frac{1}{2}mv_A^2 + mgZ_A \right) - \left(\frac{1}{2}mv_O^2 + mgZ_O \right) = -fmgh \Rightarrow v_A = \sqrt{2gh(1-f)}$$

en notant Z l'altitude des points.

Application numérique : $v_A = 7,7 \text{ m}\cdot\text{s}^{-1}$. Ne pas tenir compte des frottements revient à prendre $f = 0$ dans les expressions précédentes. On aurait alors une vitesse d'arrivée égale à $9,9 \text{ m}\cdot\text{s}^{-1}$.

16.5 Cycliste au tour de france

Le sol est lié au référentiel galiléen terrestre. Le système constitué du cycliste et de son vélo est assimilé à son centre de gravité et soumis à son poids (vertical), à la réaction du sol $\vec{R}$ (verticale), à la force de frottement de l'air (horizontale) et à une force motrice de puissance P (horizontale).

1. On lui applique le théorème de la puissance cinétique donne :

$$\frac{dE_c}{dt} = \mathcal{P}(m\vec{g}) + \mathcal{P}(\vec{F}_f) + \mathcal{P}(\vec{R}) + P$$

or le déplacement est horizontal donc le poids et la réaction ont une puissance nulle. La puissance de la force de frottement est : $\mathcal{P}(\vec{F}_f) = -kv\vec{v} \cdot \vec{v} = -kv^3$. En dérivant l'énergie cinétique on trouve :

$$mv\frac{dv}{dt} = P - kv^3.$$

D'autre part, on peut écrire : $\frac{dv}{dt} = \frac{dv}{dx} \frac{dx}{dt} = \frac{dv}{dx}v$. En remplaçant dans la précédente équation, on trouve l'équation demandée :

$$mv^2\frac{dv}{dx} = P - kv^3 = k(v_l^3 - v^3),$$

où $kv_l^3 = P \Rightarrow v_l = \left(\frac{P}{k}\right)^{\frac{1}{3}}$ est la vitesse limite pour laquelle la puissance motrice P est compensée par la puissance de freinage des frottements fluides. On peut remarquer que, pour $v = v_l$, $\frac{dv}{dt} = 0$ et le mouvement est uniforme.

2. On introduit $f(x) = k(v_l^3 - v^3)$.

a. La dérivée de $f(x)$ par rapport à x donne $f'(x) = -3kv^2\frac{dv}{dx}$. L'équation différentielle précédente devient alors :

$$-\frac{m}{3k}f'(x) = f(x) \Leftrightarrow f'(x) + \frac{3k}{m}f(x) = 0,$$

équation différentielle linéaire du premier ordre à coefficient constant.

b. On pose $L = \frac{m}{3k}$, ce qui donne alors : $f(x) = A\exp(-x/L)$. Or à $t = 0$, $x = 0$ et $v = v_0$ soit $f(0) = A = k(v_l^3 - v_0^3)$. On en déduit finalement la vitesse :

$$v(x) = v_l \left(1 - \left(1 - \left(\frac{v_0}{v_l} \right)^3 \right) \exp\left(-\frac{x}{L}\right) \right)^{1/3}.$$

Si $v_l > v_0$, le cycliste accélère jusqu'à atteindre la vitesse v_l . Si $v_l < v_0$, le cycliste ralentit jusqu'à atteindre la vitesse v_l . Cela confirme bien que v_l est la vitesse limite du cycliste. L est la distance caractéristique nécessaire pour que le cycliste atteigne sa vitesse limite.

c. Avec $P = 2 \text{ kW}$ et $v_l = 20 \text{ m}\cdot\text{s}^{-1}$, on trouve $k = \frac{P}{v_l^3} = 0,25 \text{ kg}\cdot\text{m}^{-1}$. On en déduit $L = \frac{m}{3k} = 113 \text{ m}$. Il faut quelques L pour atteindre cette vitesse limite de $70 \text{ km}\cdot\text{h}^{-1}$ (de l'ordre de 3 à 5). C'est une distance importante sur laquelle le cycliste ne pourra pas maintenir une telle puissance. C'est pour cela que des équipiers lui « lance » le sprint en le protégeant de l'air ce qui diminue le coefficient k .

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

16.6 Bille dans une gouttière

Le référentiel lié au sol est galiléen. Les forces s'appliquant au système M sont le poids (force conservative) et la réaction du support $\vec{R}$ (force de travail nul).

1. Le poids est conservatif et la réaction normale ne travaille pas. On est donc dans un cas de conservation de l'énergie mécanique :

$$E_m(A) = E_m(O) \Rightarrow \frac{1}{2}mv_A^2 + mgy_A = \frac{1}{2}mv_O^2 + mgy_O,$$

soit puisque $v_A = 0$:

$$\frac{1}{2}mv_O^2 = mg(y_A - y_O) = mgh \Rightarrow v_O = \sqrt{2gh}.$$

Pour calculer la vitesse en un point M quelconque du cercle. On repère M par sa coordonnée angulaire θ et on reprend le même raisonnement entre O et M . D'après la figure (16.11), l'altitude de M est $z = a(1 - \cos \theta)$. On obtient alors :

$$v_m = \sqrt{2(h + a(\cos \theta - 1))}.$$

2. Pour déterminer la réaction, on fait l'hypothèse que la bille reste en contact avec la gouttière. Dans ce cas, son mouvement est circulaire et on utilise les relations cinématiques :

$$\vec{CM} = a\vec{u}_r \quad ; \quad \vec{v} = a\dot{\theta}\vec{u}_\theta \quad ; \quad \vec{a} = -\frac{v^2}{a}\vec{u}_r + a\ddot{\theta}\vec{u}_\theta.$$

On applique alors la relation fondamentale de la dynamique en projection sur $\vec{u}_r$ soit :

$$-m\frac{v^2}{a} = -R + mg \cos \theta.$$

En utilisant la vitesse v_M calculée précédemment, on obtient la réaction en M :

$$R = mg \left(3 \cos \theta + 2\frac{h}{a} - 2 \right).$$

3. Pour que la bille ait un mouvement révolutif, il faut que la réaction ne s'annule en aucun point du cercle. Son expression montre qu'elle est minimale en $\theta = \pi$ (au sommet S du cercle). On souhaite donc que $R(\theta = \pi) > 0$ soit :

$$mg \left(-3 + 2\frac{h}{a} - 2 \right) > 0 \Rightarrow h > \frac{5}{2}a.$$

Il faut donc lâcher la bille d'une hauteur 2,5 fois supérieure à celle de S pour que la bille puisse faire un tour en étant plaquée sur la piste.

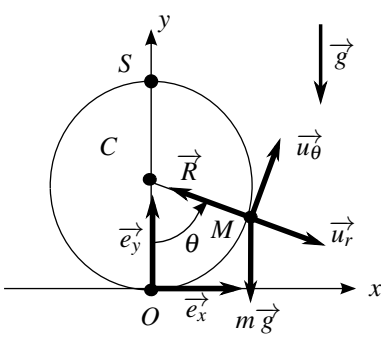


Figure 16.11

16.7 Tige avec ressort

1. Le référentiel du laboratoire est considéré galiléen. Le système M est soumis à son poids, à la tension du ressort et à la réaction de la tige.

La réaction ne travaille pas et les deux autres forces sont conservatives donc l'énergie mécanique est constante.

L'énergie potentielle associée au poids est $mgz + cte$ et celle dont dérive la tension du ressort $\frac{1}{2}k(\ell - \ell_v)^2 + cte$. Dans le cas présent, la longueur du ressort $\ell = X$ et l'altitude $z = X \cos \alpha$. d'où l'expression de l'énergie potentielle totale :

$$E_p = mgX \cos \alpha + \frac{1}{2}k(X - \ell_v)^2 + cte.$$

On prend par exemple $E_p = 0$ pour $X = \ell_v$ et on obtient l'expression de la constante : $cte = -mg\ell_v \cos \alpha$.

2. On exprime l'énergie mécanique :

$$E_m = \frac{1}{2}m\dot{X}^2 + mgX \cos \alpha + \frac{1}{2}k(X - \ell_v)^2 - mg\ell_v \cos \alpha.$$

Comme elle est constante, sa dérivée est nulle :

$$\frac{dE_m}{dt} = \dot{X}(m\ddot{X} + mg \cos \alpha + k(X - \ell_v)) = 0.$$

On exclut $\dot{X} = 0 \forall t$ qui correspond à l'immobilité, d'où l'équation du mouvement :

$$\ddot{X} + \frac{k}{m}X = \frac{k}{m}\ell_v - mg \cos \alpha.$$

3. Etude graphique.

a. Pour étudier la fonction E_p , on commence par calculer sa dérivée :

$$\frac{dE_p}{dX} = mg \cos \alpha + k(X - \ell_v).$$

Cette dérivée est nulle pour $X = \ell_v - \frac{mg \cos \alpha}{k}$ qui est bien positif d'après l'hypothèse de l'énoncé. Cette position correspond à un minimum $E_{p,\min} = -\frac{(mg \cos \alpha)^2}{2k}$. On trace alors $E_p(x)$:

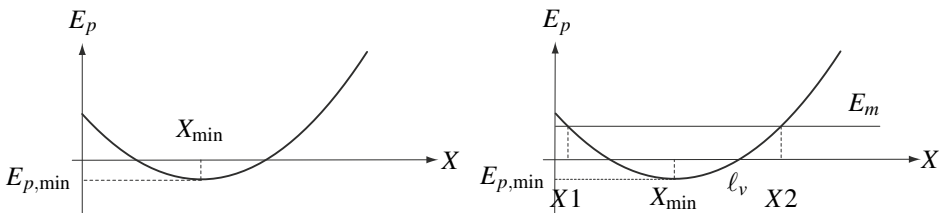


Figure 16.12

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

b. Comme l'énergie cinétique est positive ou nulle, le mouvement a lieu dans les zones telles que $E_m \geq E_p$, l'égalité correspondant aux positions extrêmes. Sur la figure (16.12), on a tracé l'énergie mécanique égale à $\frac{1}{2}V_0^2$, car à $t = 0$ $E_p = 0$. Elle est supérieure à E_p en ℓ_v car l'énergie cinétique initiale n'est pas nulle.

Les points extrêmes atteints correspondent à une énergie cinétique nulle, donc :

$$E_m = E_p \Rightarrow \frac{1}{2}mV_0^2 = mgX \cos \alpha + \frac{1}{2}k(X - \ell_v)^2 - mg\ell_v \cos \alpha.$$

D'après la figure $X_1 > 0$ si $E_m < E_p(0)$ soit :

$$\frac{1}{2}mV_0^2 < \frac{1}{2}k\ell_v^2 - mg\ell_v \cos \alpha \Rightarrow V_0 < \pm \sqrt{\frac{\ell_v^2 - 2mg\ell_v \cos \alpha}{m}}$$

En pratique le point O ne pourra pas être atteint à cause de la longueur minimale du ressort lorsque toutes les spires se touchent.

c. On obtient l'expression de la vitesse en écrivant que l'énergie mécanique à tout instant est égale sa valeur à $t = 0$:

$$\frac{1}{2}mV^2 + mgX \cos \alpha + \frac{1}{2}k(X - \ell_v)^2 - mg\ell_v \cos \alpha = \frac{1}{2}mV_0^2$$

soit

$$V = \pm \sqrt{V_0^2 - 2gX \cos \alpha - \frac{k(X - \ell_v)^2}{m} + 2g\ell_v \cos \alpha}.$$

Pour déterminer $X(t)$, on utilise l'équation du mouvement établie à la question 2. On pose $X_{\min} = \ell_v - \frac{mg \cos \alpha}{k}$ et $\omega_0 = \sqrt{k/m}$, d'où :

$$\ddot{X} + \omega_0^2 X = \omega_0^2 X_{\min} \Rightarrow X = X_{\min} + A \cos \omega_0 t + B \sin \omega_0 t.$$

On détermine les constantes à $t = 0$ soit, $X(0) = \ell_v = X_{\min} + A$ et $\dot{X}(0) = V_0 = \omega_0 B$ d'où :

$$X(t) = X_{\min} + (\ell_v - X_{\min}) \cos \omega_0 t + \frac{V_0}{\omega_0} \sin \omega_0 t.$$

16.8 Pendule simple modifié

Le référentiel du laboratoire est considéré galiléen. Les forces s'exerçant sur le système sont : le poids $m\vec{g}$ et la tension du fil $\vec{T}$.

1. Le théorème de la puissance mécanique appliqué à la première phase donne :

$$\frac{dE_c}{dt} = \mathcal{P} \Rightarrow m\vec{v} \frac{d\vec{v}}{dt} = m\vec{g} \cdot \vec{v} + \vec{T} \cdot \vec{v}.$$

Sachant que pendant la première phase, le mouvement est circulaire de rayon L , on a $\vec{v} = L\dot{\theta}\vec{u}_\theta$. Ainsi puisque $\vec{T}$ est perpendiculaire à $\vec{v}$, on trouve :

$$L^2\dot{\theta}\ddot{\theta} = -mgL\dot{\theta} \sin \theta.$$

En excluant l'immobilité ($\dot{\theta} = 0 \forall t$), on trouve l'équation classique du pendule simple :

$$\ddot{\theta} + \frac{g}{L} \sin \theta = 0.$$

2. Pour $\sin \theta \simeq \theta$, on trouve des oscillations sinusoïdales de pulsation $\omega_0 = \sqrt{g/L}$ et de période $T_0 = 2\pi/\omega_0$. Entre l'instant initial et t_1^- il y a un quart de période, soit $T_0/4$, d'où

$$\delta t_I = \frac{\pi}{2} \sqrt{\frac{L}{g}}.$$

3. La tension ne travaille pas et le poids est conservatif donc l'énergie mécanique E_m se conserve. L'énergie potentielle du poids a pour expression $mgz = -mgL \cos \theta$ si on la prend nulle en O . On écrit l'égalité des expressions de E_m à $t = 0$ et à t_1^- :

$$-mgL \cos \theta_0 = \frac{1}{2} m v_1^{-2} - mgL \rightarrow v_1^- = -\sqrt{2gL(1 - \cos \theta_0)}.$$

La vitesse angulaire est donc $\omega_1^- = v_1^-/L$ soit $\omega_1^- = -\sqrt{\frac{2g(1 - \cos \theta_0)}{L}}$.

4. D'après l'énoncé $E_m(t_1^-) = E_m(t_1^+)$, donc puisque la position est pratiquement la même, l'énergie potentielle est la même et il y a égalité des énergies cinétique d'où $v_1^- = v_1^+$. En revanche $\omega_1^+ = v_1^+/(2L/3)$ d'où :

$$\omega_1^+ = -\frac{3}{2} \sqrt{\frac{2g(1 - \cos \theta_0)}{L}}.$$

Du fait du raccourcissement brutal de la longueur de fil, la vitesse angulaire du pendule varie brutalement.

5. La deuxième phase correspond à une demi période d'un pendule de longueur $2L/3$, donc

la durée est $\delta t_{II} = \pi \sqrt{\frac{2L}{3g}}$.

6. On écrit de nouveau l'égalité de l'énergie mécanique entre t_1^+ et t_2 sachant qu'à t_2 l'altitude z vaut $-\frac{L}{3} - \frac{2L}{3} \cos \theta_2$:

$$\frac{1}{2} m v_1^{+2} - mgL = -mg \frac{L}{3} (1 + 2 \cos \theta_2) \Rightarrow \cos \theta_2 = \frac{3 \cos \theta_0 - 1}{2}.$$

7. Après la phase 2, on effectue une phase 2' inversée par rapport à la phase 2 puis une phase 1' inversée par rapport à la phase 1 qui ramène le pendule dans sa position initiale. Le pendule repart alors pour un mouvement identique (mouvement périodique). Dans la réalité l'oscillation va peu à peu s'amortir en raison des frottements.

La période T est la somme d'une demi période d'un pendule de longueur L et d'une demi période d'un pendule de longueur $2L/3$ soit $T = \pi \sqrt{\frac{L}{g}} \left(1 + \sqrt{\frac{2}{3}} \right)$.

CHAPITRE 16 – ASPECTS ÉNERGÉTIQUES DE LA DYNAMIQUE DU POINT

8. Le portrait de phase d'un oscillateur sinusoïdal est une ellipse, donc ici on aura deux demi ellipses :

phase 1 :
$$\begin{cases} \theta &= \theta_0 \cos \omega_0 t \\ \dot{\theta} &= -\omega_0 \theta_0 \sin \omega_0 t \end{cases}$$

phase 2 :
$$\begin{cases} \theta &= \theta_2 \cos \left(\sqrt{\frac{3}{2}} \omega_0 t + \phi \right) \\ \dot{\theta} &= -\sqrt{\frac{3}{2}} \omega_0 \theta_2 \sin \left(\sqrt{\frac{3}{2}} \omega_0 t + \phi \right) \end{cases}$$

La discontinuité de la vitesse angulaire est due à la discontinuité de la longueur du fil.

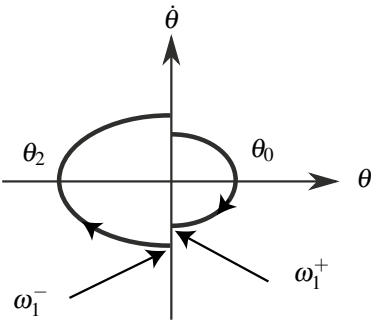


Figure 16.13

Mouvement dans un puits de potentiel

17

L'objet de ce chapitre est l'étude des mouvements dans un puits de potentiel, que l'on va définir. On y applique les résultats établis dans les chapitres précédents.

1 Mouvement conservatif dans un puits de potentiel

On définit un **puits de potentiel** comme une zone de l'espace qui entoure une position d'équilibre stable. On rappelle qu'en une telle position, l'énergie potentielle admet un minimum. Un puits de potentiel est donc une zone entourant un minimum d'énergie potentielle. Dans toute cette partie, on considère des mouvements pour lesquels les seules forces qui travaillent sont conservatives. On néglige ainsi systématiquement les forces de frottement.

1.1 Mouvement dans un puits de potentiel harmonique

On considère un objet modélisé par un point matériel M de masse m se déplaçant sur un axe (Ox) dans un référentiel $\mathcal{R}$ supposé galiléen. Cet objet se déplace dans un puits de potentiel harmonique de la forme :

$$E_p(x) = \frac{1}{2}kx^2.$$

On néglige toute force non conservative et notamment tout frottement. Les conditions initiales sont les suivantes : à l'instant $t = 0$, le mobile passe par le point d'abscisse $x = x_0$ avec une vitesse $\vec{v} = v_0\vec{u}_x$.

Remarque

Cette situation est un cas d'école auquel on peut très fréquemment se ramener.

Le point M n'est soumis qu'à des forces conservatives dont la résultante dérive de $E_p(x)$. Le théorème de l'énergie mécanique appliqué à M stipule donc que l'énergie mécanique du système est une intégrale première du mouvement :

$$E_m = \text{constante}.$$

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

Le mouvement de M a lieu sur l'axe (Ox) . Le point M est repéré par son vecteur position $\overrightarrow{OM} = x\vec{u}_x$ et sa vitesse vaut $\vec{v} = \dot{x}\vec{u}_x$. L'équation du mouvement donnée par la conservation de l'énergie mécanique s'écrit donc :

$$E_m = \frac{1}{2}mv^2 + E_p(x) = \frac{1}{2}m\dot{x}^2 + \frac{1}{2}kx^2 = \text{constante},$$

et on peut déterminer cette constante à l'instant initial : $E_m(t = 0) = \frac{1}{2}mv_0^2 + \frac{1}{2}kx_0^2 = E_{m_0}$.

a) Graphe énergétique

On peut alors tracer le graphe énergétique de ce mouvement en faisant apparaître l'énergie mécanique E_{m_0} . La courbe d'énergie potentielle est une parabole et la position $x = 0$ est un minimum d'énergie potentielle. Il s'agit donc d'une position d'équilibre stable. L'ensemble des positions accessibles vérifie $E_{m_0} > E_p(x)$. C'est un segment $[-X_m, X_m]$ symétrique par rapport à la position d'équilibre. La trajectoire est une trajectoire liée et le mobile M oscille périodiquement avec une amplitude X_m autour de sa position d'équilibre.

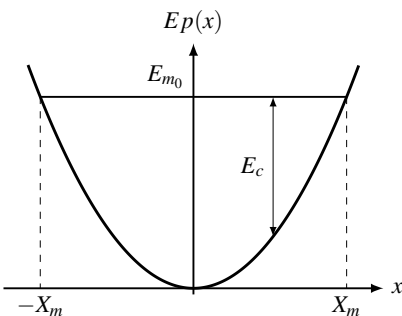


Figure 17.1 – Potentiel harmonique.

Pour le mouvement d'énergie mécanique E_{m_0} , le point $x = X_m$ est un point où l'énergie cinétique E_c s'annule. C'est donc un point de vitesse nulle et X_m vérifie :

$$E_p(x = X_m) = \frac{1}{2}kX_m^2 = E_{m_0} = \frac{1}{2}mv_0^2 + \frac{1}{2}kx_0^2 \quad \implies \quad X_m = \sqrt{x_0^2 + \frac{m}{k}v_0^2}.$$

b) Trajectoire de phase

On peut alors réécrire l'équation du mouvement en introduisant l'amplitude X_m :

$$\frac{1}{2}kx^2 + \frac{1}{2}m\dot{x}^2 = \frac{1}{2}kX_m^2$$

puis divisant cette équation par $\frac{1}{2}k$ et en posant $\omega_0 = \sqrt{\frac{k}{m}}$:

$$\omega_0^2 x^2 + \dot{x}^2 = \omega_0^2 X_m^2. \quad \iff \quad \frac{x^2}{X_m^2} + \frac{\dot{x}^2}{(\omega_0 X_m)^2} = 1.$$

On reconnaît l'équation cartésienne d'une ellipse du plan de phase $(x, \dot{x})$ dont la forme générale est décrite dans l'appendice mathématique :

$$\frac{x^2}{a^2} + \frac{\dot{x}^2}{b^2} = 1,$$

avec ici $a = X_m$ et $b = \omega_0 X_m$. La trajectoire de phase d'un mouvement harmonique est une ellipse.

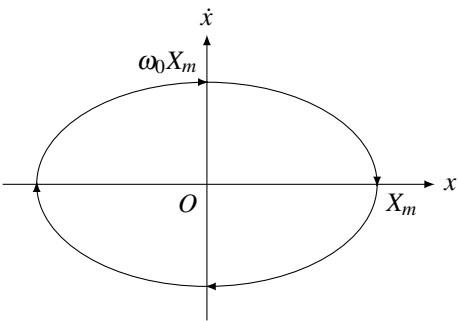


Figure 17.2 – Trajectoire de phase du mouvement de M.

La trajectoire de phase est une courbe fermée. Le mouvement est donc périodique. Par ailleurs, l'amplitude est nécessairement égale à X_m .

c) Équation horaire du mouvement

L'équation énergétique du mouvement $E_m = E_c + E_p = \frac{1}{2}m\dot{x}^2 + \frac{1}{2}kx^2$ a déjà été étudiée dans le chapitre sur l'oscillateur harmonique. On obtient l'équation du mouvement en la dérivant par rapport au temps. On pose $\omega_0 = \sqrt{\frac{k}{m}}$ et elle se met sous la forme :

$$\ddot{x} + \omega_0^2 x = 0.$$

C'est l'équation du mouvement d'un oscillateur harmonique de pulsation propre $\omega_0 = \sqrt{\frac{k}{m}}$. Les solutions de cette équation ont été étudiées en détail au chapitre sur l'oscillateur harmonique. Elles présentent des oscillations sinusoïdales de pulsation ω_0 que l'on écrit :

$$x(t) = X_m \cos(\omega_0 t + \varphi_0).$$

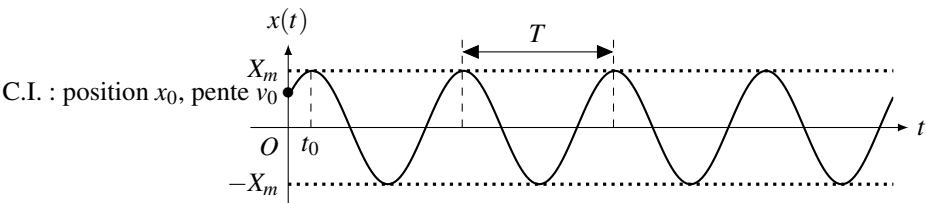


Figure 17.3 – Évolution temporelle de la position x du mobile.

À l'aide des conditions initiales, on trouve les expressions de X_m et de φ_0 :

$$X_m = \sqrt{x_0^2 + \left(\frac{v_0}{\omega_0}\right)^2} \quad \text{et} \quad \varphi_0 \text{ vérifie } \begin{cases} \cos \varphi_0 = \frac{x_0}{X_m} \\ \sin \varphi_0 = -\frac{v_0}{\omega_0 X_m}. \end{cases}$$

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

Le mouvement de M présente des oscillations sinusoïdales de période $T = \frac{2\pi}{\omega_0} = 2\pi\sqrt{\frac{m}{k}}$, d'amplitude X_m et de phase à l'origine φ_0 . La figure 17.3 représente l'évolution de son abscisse $x(t)$ au cours du temps.

Dans un puits de potentiel harmonique $E_p(x) = \frac{1}{2}kx^2$, un mobile de masse m est animé d'oscillations sinusoïdales dont la période $T = 2\pi\sqrt{\frac{m}{k}}$ ne dépend pas de l'amplitude. On parle d'**isochronisme** des oscillations.

1.2 Mouvement dans un puits de potentiel quelconque

a) Mouvement de faible amplitude au voisinage d'une position d'équilibre stable

On considère une particule ponctuelle de masse m soumise à la seule force $\vec{F}$ conservative dirigée selon le vecteur $\vec{u}_x$ et ne dépendant que de x . Cette force dérive de l'énergie potentielle E_p ne dépendant que d'un paramètre de position x . On a établi dans le chapitre précédent :

$$\vec{F} = F_x(x)\vec{u}_x \quad \text{avec} \quad F_x(x) = -\frac{dE_p}{dx}.$$

On a également établi qu'en un point d'équilibre stable d'abscisse x_{eq} , la force $\vec{F}$ s'annule et l'énergie potentielle admet un minimum. Dans le cas où $E_p(x)$ est deux fois dérivable et où sa dérivée seconde n'est pas nulle en $x = x_{eq}$, on en a déduit que :

$$\left(\frac{dE_p}{dx}\right)_{x=x_{eq}} = 0 \quad \text{et} \quad \left(\frac{d^2E_p}{dx^2}\right)_{x=x_{eq}} > 0.$$

Même si la force $\vec{F}$ n'est pas une force de rappel élastique, on va montrer que, sous réserve de ne pas trop s'éloigner de la position d'équilibre stable d'abscisse x_{eq} , on peut assimiler le mouvement du mobile au mouvement obtenu dans un puits de potentiel harmonique.

Graphiquement, cette approximation revient à assimiler l'énergie potentielle tracée en trait plein sur la figure 17.4 au potentiel harmonique tangent de forme parabolique tracé en pointillés sur la même figure. Il est visible que cette approximation n'est valable qu'au voisinage du minimum d'énergie potentielle.

Analytiquement, on peut effectuer un développement limité à l'ordre 2 de l'énergie potentielle au voisinage d'une position d'équilibre stable x_{eq} :

$$E_p(x) = E_p(x_{eq}) + (x - x_{eq}) \left(\frac{dE_p}{dx}\right)_{x=x_{eq}} + \frac{(x - x_{eq})^2}{2} \left(\frac{d^2E_p}{dx^2}\right)_{x=x_{eq}}. \quad (17.1)$$

Le fait d'effectuer un développement limité restreint la validité de l'étude qui en découle à de faibles écarts à la position d'équilibre x_{eq} .

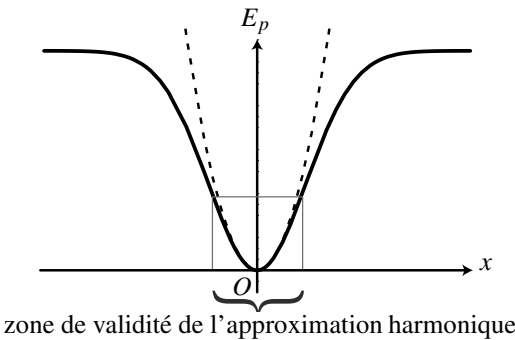


Figure 17.4 – Puits de potentiel (trait plein) et son approximation harmonique valable au voisinage du minimum d'énergie potentielle (en pointillé).

En utilisant le fait que $\left(\frac{dE_p}{dx}\right)_{x=x_{eq}} = 0$, en notant $k = \left(\frac{d^2E_p}{dx^2}\right)_{x=x_{eq}} > 0$ et en restant au voisinage de x_{eq} , l'équation (17.1) permet d'assimiler la fonction $E_p(x)$ à une parabole d'équation :

$$E_p(x) = E_p(x_{eq}) + k \frac{(x - x_{eq})^2}{2}.$$

L'énergie potentielle étant définie à une constante près, il est toujours possible de choisir la référence des énergies potentielles en $x = x_{eq}$ pour fixer $E_p(x_{eq}) = 0$. En posant $X = x - x_{eq}$, ce qui revient à décaler l'origine de l'axe des x à la position d'équilibre, on peut ensuite réécrire l'énergie potentielle sous la forme :

$$E_p(x) = \frac{1}{2}kX^2.$$

On peut assimiler le potentiel au voisinage de la position x_{eq} à un potentiel harmonique de constante de rappel élastique $k = \left(\frac{d^2E_p}{dx^2}\right)_{x=x_{eq}}$. On est alors ramené au cas d'école du paragraphe précédent.

Les mouvements de faibles amplitudes au voisinage d'une position d'équilibre stable peuvent être assimilés à des mouvements d'oscillations harmoniques centrées sur la position d'équilibre x_{eq} et de pulsation $\omega_0 = \sqrt{\frac{k}{m}}$ où la constante de rappel élastique k est donnée par la relation :

$$k = \left(\frac{d^2E_p}{dx^2}\right)_{x=x_{eq}}.$$

b) Mouvements de grande amplitude dans un puits de potentiel quelconque

On a vu dans le paragraphe précédent que les mouvements de faible amplitude au voisinage d'une position d'équilibre stable peuvent être modélisés par des mouvements dans un puits de

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

potentiel harmonique. On se pose ici la question de savoir ce qui change lorsque l'amplitude des oscillations est trop grande pour pouvoir réaliser cette approximation.

On rappelle les caractéristiques essentielles du mouvement harmonique : il s'agit d'oscillations périodiques, de forme sinusoïdale et de période indépendante de l'amplitude. On va montrer que, lorsque la forme du potentiel n'est plus harmonique, le mouvement est toujours constitué d'oscillations, mais que ces oscillations ne sont ni sinusoïdales, ni isochrones.

Pour visualiser ces changements, on trace les solutions, obtenues par résolution numérique, des mouvements de grande amplitude dans le puits de potentiel de la figure 17.4. Ce puits est de la forme :

$$E_p(x) = E_0 \left(1 - \exp \left(- \left(\frac{x}{L} \right)^2 \right) \right),$$

ce qui fournit une échelle d'énergie E_0 qui caractérise la profondeur du puits de potentiel et une échelle de longueur L qui caractérise sa largeur. Pour un mobile de masse m , on trouve

alors une échelle de vitesse $v_0 = \sqrt{\frac{E_0}{m}}$ puisqu'une énergie est égale au produit d'une masse par une vitesse au carré. On en déduit une échelle de temps $t_0 = \frac{L}{v_0} = L \sqrt{\frac{m}{E_0}}$.

On peut alors écrire l'équation du mouvement en appliquant le principe fondamental de la dynamique projeté sur l'axe (Ox) :

$$\frac{d(m\dot{x})}{dt} = -\frac{dE_p}{dx} = -E_0 \frac{2x}{L^2} \exp \left(- \left(\frac{x}{L} \right)^2 \right).$$

On introduit les grandeurs adimensionnées $x^* = \frac{x}{L}$ et $t^* = \frac{t}{t_0}$ et on obtient l'équation différentielle du mouvement adimensionnée :

$$\frac{mL}{t_0^2} \frac{d^2(x^*)}{dt^{*2}} = -\frac{2E_0}{L} x^* \exp(-x^{*2}) \quad \Longleftrightarrow \quad \frac{d^2 x^*}{dt^{*2}} = -2x^* \exp(-x^{*2})$$

que l'on peut résoudre numériquement. Les courbes de la figure 17.5 sont établies avec comme condition initiale : $x(t=0) = 0$ et $\dot{x}(t=0) > 0$.

Les courbes en trait plein noir correspondent à l'énergie mécanique E_{m1} faible. Le mouvement est alors de faible amplitude et quasi-sinusoïdal de période $T^* \simeq 4,5$. Les courbes en trait pointillé correspondent à une énergie mécanique E_{m2} plus importante. L'amplitude et la période du mouvement augmentent toutes les deux. La période atteint $T^* \simeq 5,8$. Les courbes en trait plein gris correspondent à l'énergie mécanique E_{m3} encore plus grande. L'amplitude et la période augmentent encore. La période atteint $T^* \simeq 10,9$. La trajectoire de phase n'est clairement plus elliptique.

On peut noter que, dans l'approximation harmonique, les mouvements de faible amplitude sont sinusoïdaux de pulsation $\omega_0 = \sqrt{\frac{k}{m}}$ où $k = \left(\frac{d^2 E_p}{dx^2} \right)_{x=0} = 2 \frac{E_0}{L^2} = 2 \frac{m}{t_0^2}$ soit $\omega_0 = \frac{\sqrt{2}}{t_0}$.

La période $T = \frac{2\pi}{\omega_0} = \sqrt{2}\pi t_0$ correspond alors à $T^* = \sqrt{2}\pi \simeq 4,5$.

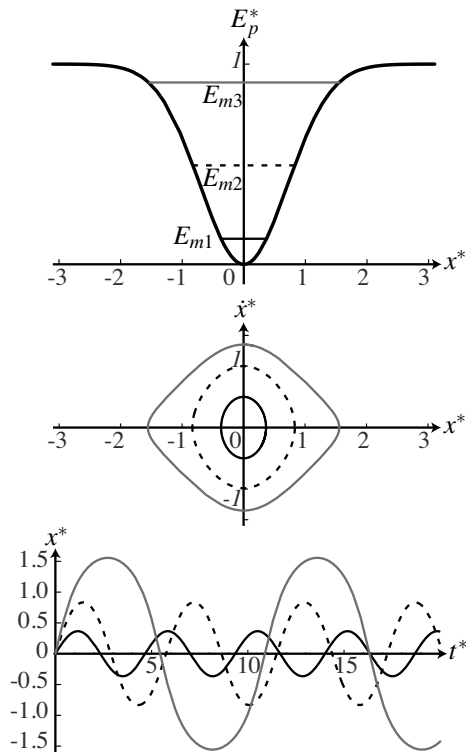


Figure 17.5 – Mouvement dans un potentiel non harmonique pour trois énergies mécaniques différentes.

On remarque l’évolution de la période des oscillations avec l’amplitude, ce qui prouve que l’isochronisme des oscillations n’est pas vérifié par les mouvements de grande amplitude. Par ailleurs, la forme des oscillations évolue également, ce qui est particulièrement visible sur les trajectoires de phase. Les oscillations de grande amplitude ne sont plus sinusoïdales.

La présence de mouvements d’oscillations sinusoïdales isochrones est une signature des mouvements dans des potentiels rigoureusement ou quasiment harmoniques.

c) Condition de sortie d’un puits de potentiel

On se pose maintenant la question de savoir quelle énergie il est nécessaire de procurer au mobile pour qu’il puisse sortir d’un puits de potentiel. La réponse est directement issue du théorème de l’énergie mécanique et du signe positif de l’énergie cinétique :

$$E_m = E_p + E_c \geq E_p.$$

Pour que le mobile puisse accéder à tous les points, il faut donc que son énergie mécanique dépasse la plus haute valeur de l’énergie potentielle. Pour franchir une barrière de potentiel de profondeur E_0 , il doit donc avoir une énergie mécanique supérieure à E_0 (voir figure 17.6).

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

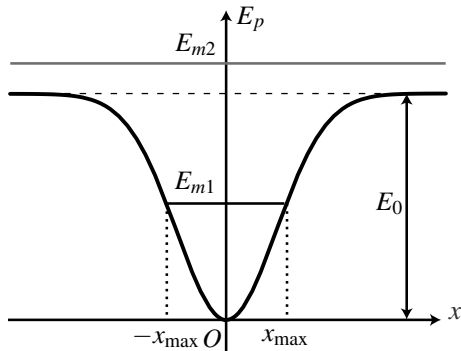


Figure 17.6 – Puits de potentiel. Avec l'énergie mécanique E_{m1} , le mobile est confiné entre les abscisses $x_{\max}$ et $-x_{\max}$. Avec l'énergie mécanique E_{m2} , il peut sortir du puits et accéder à l'ensemble des positions.

2 Mouvements dans un puits de potentiel : influence des frottements

Dans la partie précédente, on n'a pas tenu compte des forces de frottement. Or les frottements existent toujours mais se font plus ou moins ressentir. On s'intéresse ici au cas où on peut les modéliser par une force proportionnelle à la vitesse :

$$\vec{f} = -\lambda \vec{v}.$$

On restreint l'étude aux mouvements dans un puits de potentiel harmonique pour lequel la résolution analytique est possible. On peut envisager d'autres formes de potentiels mais il faudra alors chercher des solutions numériques.

2.1 Équation différentielle du mouvement

On reprend le puits de potentiel harmonique étudié au paragraphe 1. L'énergie mécanique du système vaut alors $E_m = E_c + E_p = \frac{1}{2} m \dot{x}^2 + \frac{1}{2} k x^2$. Le mouvement est à un seul degré de liberté et on le résout à l'aide du théorème de l'énergie cinétique :

$$\frac{d(E_c + E_p)}{dt} = \mathcal{P}(\vec{f}_{NC}), \tag{17.2}$$

où $\mathcal{P}(\vec{f}_{NC})$ est la puissance des forces non conservatives, c'est-à-dire ici des frottements fluides. On calcule alors :

$$\frac{d(E_c + E_p)}{dt} = m \dot{x} \ddot{x} + k x \dot{x},$$

et

$$\mathcal{P}(\vec{f}_{NC}) = \mathcal{P}(\vec{f}) = \vec{f} \cdot \vec{v} = -\lambda \vec{v} \cdot \vec{v} = -\lambda \dot{x}^2,$$

que l'on remplace dans l'équation (17.2) pour aboutir à $m \dot{x} \ddot{x} + k x \dot{x} = -\lambda \dot{x}^2$. On simplifie ensuite par $\dot{x}$ qui correspond à un point M immobile à tout instant et l'équation différentielle

du mouvement devient :

$$\ddot{x} + \frac{\lambda}{m} \dot{x} + \frac{k}{m} x = 0 \Leftrightarrow \frac{d^2x}{dt^2} + \frac{\lambda}{m} \frac{dx}{dt} + \frac{k}{m} x = 0.$$

Il s'agit de l'équation différentielle d'un oscillateur amorti.

On a déjà rencontré ce type d'équation différentielle lors de l'étude du régime libre du circuit *RLC* série dans le chapitre *Circuit linéaire du second ordre*. La tension u_C aux bornes du condensateur vérifiait alors l'équation différentielle :

$$\frac{d^2u_C}{dt^2} + \frac{R}{L} \frac{du_C}{dt} + \frac{1}{LC} u_C = 0,$$

et on l'avait écrite sous forme canonique :

$$\frac{d^2u_C}{dt^2} + \frac{\omega_0}{Q} \frac{du_C}{dt} + \omega_0^2 u_C = 0.$$

On va procéder de même ici.

En pratique, les phénomènes oscillatoires sont présents dans divers domaines de la physique et il arrive souvent que l'on puisse les modéliser par une équation différentielle du type oscillateur amorti. Lorsque c'est le cas, le modèle étant identique, ses prévisions seront analogues à celles que l'on a établi dans le chapitre *Circuit linéaire du second ordre*. Tous les phénomènes étudiés en électricité (existence de trois régimes transitoires dépendant de l'importance de la dissipation, phénomène de résonance, etc...) peuvent alors exister de manière analogue. On étudie ici l'influence des frottements sur la nature des régimes libres observés.

2.2 Équation caractéristique et observation

a) Équation caractéristique

L'équation du mouvement :

$$\ddot{x} + \frac{\lambda}{m} \dot{x} + \frac{k}{m} x = 0,$$

peut être écrite sous une forme dite canonique :

$$\ddot{x} + \frac{\omega_0}{Q} \dot{x} + \omega_0^2 x = 0, \tag{17.3}$$

ce qui revient à introduire la pulsation propre ω_0 et le facteur de qualité Q tels que :

$$\begin{cases} \omega_0^2 = \frac{k}{m} \\ \frac{\omega_0}{Q} = \frac{\lambda}{m} \end{cases} \quad \text{soit} \quad \begin{cases} \omega_0 = \sqrt{\frac{k}{m}} \\ Q = \frac{\sqrt{km}}{\lambda} \end{cases}.$$

On utilise cette notation dans toute la suite mais d'autres écritures peuvent être adoptées. On peut notamment utiliser la forme canonique suivante :

$$\frac{\ddot{x}}{\omega_0^2} + 2\xi \frac{\dot{x}}{\omega_0} + x = 0,$$

ce qui revient à utiliser ω_0 et un facteur d'amortissement ξ relié à Q par la relation : $\xi = \frac{1}{2Q}$.

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

b) Observation de régimes différents selon la valeur de Q

La figure 17.7 illustre les différents comportements adoptés par le système en fonction de la valeur du facteur de qualité Q . Pour ces graphiques, on a choisi $x(t = 0) = x_0$ et $\dot{x}(t = 0) = 0$ comme conditions initiales. On étudie donc le **régime libre** de l'oscillateur amorti qui, éloigné d'une distance x_0 de sa position d'équilibre, y retourne librement. On se sert de x_0 pour définir une longueur sans dimension $x^* = \frac{x}{x_0}$ et de la période propre $T_0 = \frac{2\pi}{\omega_0}$ pour définir un temps sans dimension $t^* = \frac{t}{T_0}$. La vitesse sans dimension est alors $\dot{x}^* = \frac{dx^*}{dt^*} = \dot{x} \frac{T_0}{x_0}$. Enfin, on définit l'énergie mécanique initiale E_0 comme échelle d'énergie et on trace $E_c^* = \frac{E_c}{E_0}$ et $E_p^* = \frac{E_p}{E_0}$. On distingue :

- un cas où l'on observe des oscillations dont l'amplitude diminue au cours du temps appelé **régime pseudo-périodique** ;
- un cas où l'on n'observe pas d'oscillations appelé **régime aperiodique**.

Un cas limite appelé **régime critique** sépare les deux cas précédents. Il n'est pas représenté mais ressemble fortement au régime aperiodique.

2.3 Résolution : les trois régimes

La méthode de résolution de l'équation différentielle (17.3) est détaillée dans l'annexe mathématique et a été pratiquée en électricité. On introduit l'équation caractéristique associée :

$$r^2 + \frac{\omega_0}{Q}r + \omega_0^2 = 0, \quad \text{de discriminant } \Delta = \left(\frac{\omega_0}{Q}\right)^2 - 4\omega_0^2 = (2\omega_0)^2 \left(\frac{1}{4Q^2} - 1\right).$$

On distingue plusieurs types de solutions suivant le signe du discriminant.

a) Régime pseudo-périodique

Il s'agit du cas où le discriminant est négatif :

$$\Delta < 0 \Leftrightarrow Q > \frac{1}{2} \Leftrightarrow \lambda < \frac{1}{2\sqrt{km}}.$$

Cela correspond donc au cas où le frottement défini grâce au coefficient λ est faible. On peut noter que l'oscillateur harmonique correspond à la limite $\lambda = 0$. Les deux racines de l'équation caractéristique sont complexes conjuguées :

$$r_{\pm} = -\frac{\omega_0}{2Q} \pm i\omega_0 \sqrt{1 - \frac{1}{4Q^2}} = -\frac{\omega_0}{2Q} \pm i\omega,$$

et la solution de l'équation différentielle s'écrit :

$$x(t) = X_m \exp\left(-\frac{\omega_0}{2Q}t\right) \cos(\omega t + \varphi_0),$$

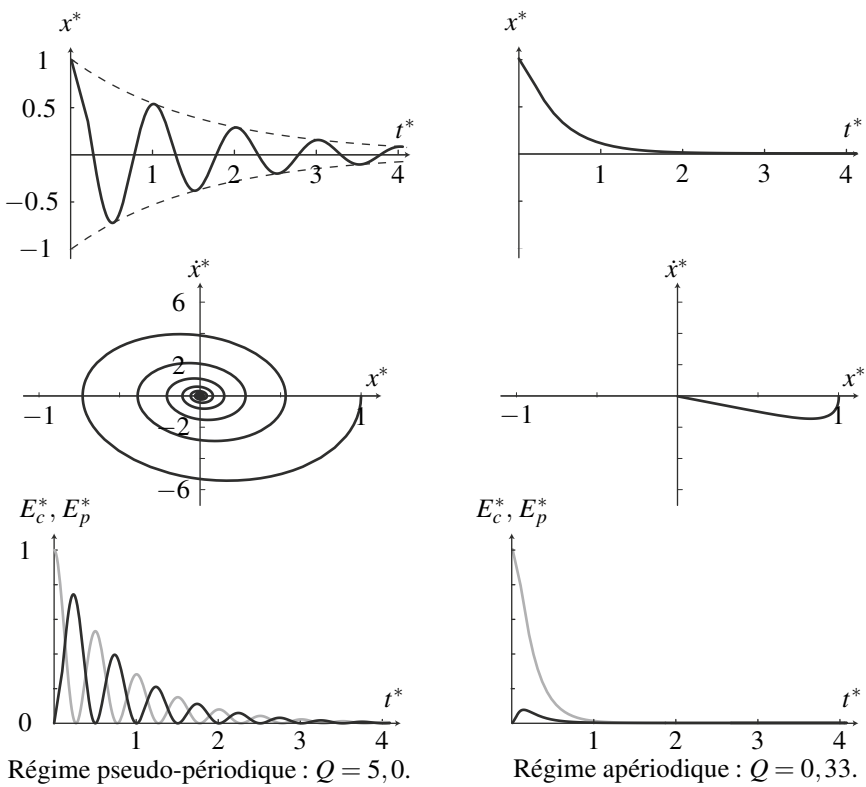


Figure 17.7 – Différents comportements du système. De haut en bas : évolution temporelle de x^* , trajectoire de phase et évolution temporelle des énergies cinétiques (courbe noire) et potentielle (courbe grise).

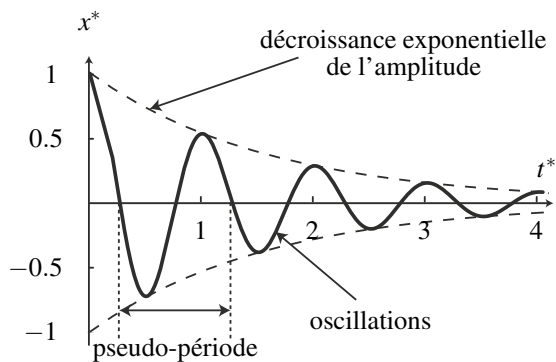


Figure 17.8 – Régime pseudo-périodique. $Q = 5,0$; $x^*(t = 0) = 1$ et $\dot{x}^*(t = 0) = 0$.

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

où (X_m, φ_0) sont des constantes à déterminer à partir des conditions initiales.

La solution que l'on a représenté sur la figure 17.8 est le produit d'une fonction d'amplitude : $A(t) = X_m \exp\left(-\frac{\omega_0}{2Q}t\right)$ qui diminue au cours du temps, par un terme d'oscillations sinusoïdales : $\cos(\omega t + \varphi_0)$. Elle correspond à des oscillations dont l'amplitude diminue au cours du temps. Cela justifie l'appellation de **régime pseudo-périodique** donnée à ce type de solutions. L'amplitude $A(t)$ des oscillations décroît de façon exponentielle avec un temps caractéristique $\tau = \frac{2Q}{\omega_0}$.

Pour caractériser les oscillations d'amplitude décroissante, on définit deux grandeurs : la **pseudo-période** et le **décroissement logarithmique**.

La pseudo-période Elle caractérise la durée d'une oscillation et mesure la période du terme oscillant $\cos(\omega t + \varphi_0)$ qui vaut :

$$T = \frac{2\pi}{\omega} = \frac{2\pi}{\omega_0 \sqrt{1 - \frac{1}{4Q^2}}}.$$

$\sqrt{1 - \frac{1}{4Q^2}} < 1$ donc $T > \frac{2\pi}{\omega_0} = T_0$ où T_0 est la période de l'oscillateur harmonique associé ou période propre. La pseudo-période est un peu plus longue que la période propre.

Le décroissement logarithmique Il caractérise l'amortissement des oscillations et mesure la diminution de leur amplitude pendant une pseudo-période :

$$\delta = \ln\left(\frac{A(t)}{A(t+T)}\right) = \ln\left(\frac{A(t)}{A(t) \exp\left(-\frac{\omega_0}{2Q}T\right)}\right) = \frac{\omega_0}{2Q}T = \frac{\pi}{Q \sqrt{1 - \frac{1}{4Q^2}}}.$$

Cas d'un amortissement faible Dans le cas d'un amortissement faible, caractérisé par la présence de nombreuses oscillations avant l'arrêt du mobile, $Q \gg 1$. On obtient alors :

$$T \simeq \frac{2\pi}{\omega_0} = T_0 \quad \text{et} \quad \delta \simeq \frac{\pi}{Q}.$$

La mesure de la pseudo-période T et du décroissement logarithmique δ donne directement accès à la pulsation propre ω_0 et au facteur de qualité Q . Les paramètres de l'équation différentielle sont donc immédiatement accessibles à la mesure.

À la limite où il n'y a pas de frottement, la période est rigoureusement égale à la période propre ($T = T_0$) et l'amplitude des oscillations est constante ($\delta = 0$), ce qui est cohérent.

Évaluation rapide du facteur de qualité L'amplitude des oscillations décroît de façon exponentielle avec un temps caractéristique $\tau = \frac{2Q}{\omega_0}$. Lorsque le facteur de qualité Q est suffisam-

ment grand (en pratique supérieur à 3), le nombre d'oscillations observables correspond approximativement à $Q \pm 1$. En effet, pour Q grand, la durée d'une oscillation vaut $T \simeq T_0 = \frac{2\pi}{\omega_0}$.

Après Q oscillations, il s'est donc écoulé le temps $\Delta t = Q \frac{2\pi}{\omega_0} = \pi \frac{2Q}{\omega_0}$ et l'amplitude $A(t)$ est réduite à $\exp(-\pi) \simeq 4.10^{-2}$ fois sa valeur initiale. Si l'amplitude initiale représente 5 cm, l'amplitude vaut 2 mm après Q oscillations. L'amplitude des oscillations suivantes est encore plus petite ce qui les rend difficiles à voir.

Remarque

Un oscillateur de bonne qualité présente un grand nombre d'oscillations avant de s'arrêter. Son facteur de qualité est grand.

b) Régime apériodique

Il s'agit du cas où le discriminant est positif :

$$\Delta > 0 \Leftrightarrow Q < \frac{1}{2} \Leftrightarrow \lambda > \frac{1}{2\sqrt{km}}.$$

Cela correspond donc au cas où le frottement défini grâce au coefficient λ est fort. L'équation caractéristique admet alors deux racines réelles :

$$r_{\pm} = -\frac{\omega_0}{2Q} \pm \omega_0 \sqrt{\frac{1}{4Q^2} - 1} = -\frac{\omega_0}{2Q} \pm \frac{\omega_0}{2Q} \sqrt{1 - 4Q^2}.$$

Comme $\frac{\omega_0}{2Q} > \frac{\omega_0}{2Q} \sqrt{1 - 4Q^2}$, les deux racines sont négatives et la solution de l'équation différentielle :

$$x(t) = X_+ \exp(r_+ t) + X_- \exp(r_- t)$$

est la somme de deux exponentielles décroissantes. L'amplitude diminue au cours du temps sans que n'apparaissent d'oscillations. Cela correspond bien à ce que l'on observe sur la figure 17.9 et justifie l'appellation **régime apériodique** donnée à ce type de solutions.

c) Régime critique

Ce régime est un cas limite qui sépare les deux cas précédents. Il correspond à $\Delta = 0$ soit $Q = \frac{1}{2}$. L'équation caractéristique admet la racine double $r = -\omega_0$ et la solution s'écrit :

$$x(t) = (\lambda t + \mu) \exp(-\omega_0 t),$$

où λ et μ sont des constantes à déterminer à partir des conditions initiales. L'allure est la même que pour le régime apériodique et on n'observe pas d'oscillations.

2.4 Évaluation rapide de la durée des différents régimes

La valeur d’une exponentielle décroissante de la forme $\exp\left(-\frac{t}{\tau}\right)$ de temps caractéristique τ est réduite à 1% de sa valeur initiale après une durée $\Delta t = -\tau \ln\left(\frac{1}{100}\right) \simeq 4,6\tau$.

- pour le régime pseudo-périodique, les oscillations voient leur amplitude décroître de façon exponentielle avec un temps caractéristique $\tau_1 = \frac{2Q}{\omega_0}$. L’amplitude est réduite à 1% de sa valeur initiale au bout d’une durée $\Delta t_1 = 4,6\frac{2Q}{\omega_0}$;
- pour le régime apériodique, les deux termes décroissent de façon exponentielle avec des temps caractéristiques $-\frac{1}{r_+}$ et $-\frac{1}{r_-}$. La décroissance est limitée par le terme qui possède la plus grande constante de temps, à savoir $\tau_2 = -\frac{1}{r_+} = \frac{2Q}{\omega_0} \frac{1}{1 - \sqrt{1 - 4Q^2}}$. Ce terme est réduit à 1% de sa valeur initiale au bout d’une durée $\Delta t_2 = 4,6\tau_2$;
- pour le régime critique, la solution comporte un terme exponentiel de temps caractéristique $\tau_3 = \frac{1}{\omega_0}$. Il est réduit à 1% de sa valeur initiale au bout d’une durée $\Delta t_3 = 4,6\tau_3$;

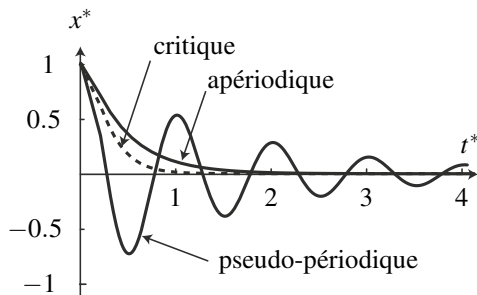


Figure 17.9 – Comparaison des durées des trois régimes.

On peut voir sur la figure 17.9 que, pour une période propre donnée, le régime critique est celui pour lequel le retour en $x = 0$ sans oscillations est le plus rapide. On cherche donc à se rapprocher de ce cas lorsque l’on veut que le système revienne rapidement à l’équilibre. Cette propriété est expliquée par les inégalités suivantes :

$$\tau_3 < \tau_1 \text{ pour } Q > \frac{1}{2} \quad \text{et} \quad \tau_3 < \tau_2 \text{ pour } Q < \frac{1}{2}.$$

Il faut noter qu’on a raisonné ici uniquement sur l’amplitude de la solution pseudo-périodique. Quand on étudie plus finement ce cas, on montre qu’il permet de revenir plus rapidement en $x = 0$, à condition que le système puisse osciller sans dommage. Pour une analyse plus complète, on peut se reporter au chapitre *Circuit linéaire du second ordre* où on définit précisément le temps de réponse à 5% qui permet de comparer la durée de chacun des régimes transitoires.

Exemple

Les amortisseurs sont destinés à ramener une voiture rapidement à l’équilibre sans oscillations excessives. On les règle pour qu’ils fonctionnent quasiment au régime critique.

2.5 Aspects énergétiques

a) Variations d'énergie

À partir des expressions de $x(t)$ obtenues précédemment, on peut tracer les différentes énergies en fonction du temps, c'est-à-dire :

- l'énergie cinétique : $E_c = \frac{1}{2}m\dot{x}^2$;
- l'énergie potentielle : $E_p = \frac{1}{2}kx^2 = \frac{1}{2}m\omega_0^2 x^2$ car $k = m\omega_0^2$;
- l'énergie mécanique : $E_m = E_c + E_p$.

La figure 17.10 représente l'évolution de ces différentes énergies dans le cas d'un régime pseudo-périodique.

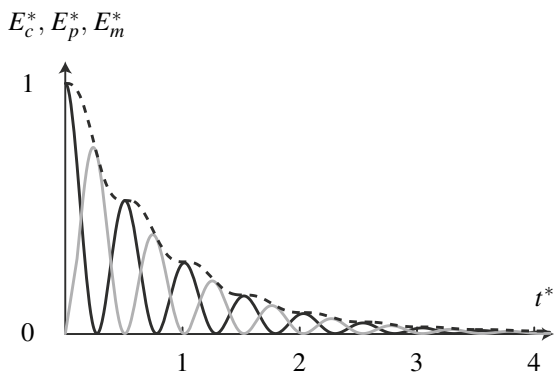


Figure 17.10 – Évolution des énergies cinétiques (courbe noire), potentielle (courbe grise) et mécanique (tirets) pour $Q = 5, 0$.

Dans ce cas, on connaît $x(t) = A(t) \cos(\omega t + \varphi_0)$ où $A(t) = X_m \exp\left(-\frac{t}{\tau}\right)$ avec $\tau = \frac{2Q}{\omega_0}$. On en déduit :

$$E_p(t) = \frac{1}{2}m\omega_0^2 A^2(t) \cos^2(\omega t + \varphi_0) = \frac{1}{2}m\omega_0^2 A^2(t) \frac{1}{2} (1 + \cos(2\omega t + 2\varphi_0)) .$$

L'énergie potentielle présente des oscillations dont l'amplitude diminue au cours du temps. La pseudo-période des oscillations vaut $\frac{T}{2}$ et l'amplitude décroît de façon exponentielle avec un temps caractéristique $\tau' = \frac{\tau}{2}$ car $A^2(t) = X_m^2 \exp\left(-2\frac{t}{\tau}\right)$.

L'énergie cinétique présente les mêmes caractéristiques mais est en opposition de phase. Au final, l'énergie mécanique n'oscille quasiment plus et présente quasiment une décroissance exponentielle de temps caractéristique τ' .

b) Origine de la variation d'énergie mécanique

En appliquant le théorème de l'énergie cinétique, on obtient :

$$\frac{d(E_c + E_p)}{dt} = \mathcal{P}(\overrightarrow{f_{NC}}) = -\lambda \dot{x}^2 < 0 ,$$

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

où $\mathcal{P}(\vec{f}_{NC})$ est la puissance des forces non conservatives, soit ici la puissance des forces de frottement. L'énergie mécanique du point matériel diminue sous l'action des frottements. Elle est dissipée sous forme de chaleur.

SYNTHÈSE

SAVOIRS

- puits de potentiel, barrière de potentiel
- puits de potentiel harmonique
- savoir que les mouvements harmoniques sont isochrones et présentent des trajectoires de phase elliptiques
- savoir que seules les forces linéaires dérivent d'un potentiel harmonique
- savoir qu'un puits de potentiel quelconque peut être assimilé à un puits harmonique au voisinage de son minimum
- savoir que les mouvements de faibles amplitudes au voisinage d'une position d'équilibre stable peuvent être assimilés à des oscillations harmoniques
- définition d'un oscillateur amorti
- forme canonique de l'équation différentielle d'un oscillateur amorti en terme de pulsation propre et de facteur de qualité
- différents régimes libres possibles
- nature du régime libre en fonction de la valeur du facteur de qualité

SAVOIR-FAIRE

- identifier les mouvements au voisinage d'une position d'équilibre stable aux mouvements d'un oscillateur harmonique
- mettre en évidence des forces non linéaires par la perte d'isochronisme et les déformations des trajectoires de phase
- évaluer l'énergie minimale nécessaire pour franchir une barrière de potentiel
- écrire sous forme canonique l'équation différentielle d'un oscillateur amorti
- identifier la pulsation propre et le facteur de qualité
- déterminer le polynôme caractéristique associé et en trouver les racines
- déterminer la réponse détaillée dans le cas d'un régime libre
- déterminer un ordre de grandeur de la durée du régime transitoire selon la valeur du facteur de qualité

MOTS-CLÉS

- | | | |
|--------------------------|------------------------|----------------------|
| • puits de potentiel | • effets non-linéaires | • facteur de qualité |
| • barrière de potentiel | • oscillateur amorti | • régime libre |
| • oscillateur harmonique | • pulsation propre | • régime transitoire |

S'ENTRAÎNER

17.1 Oscillations d'un pendule simple - partie 1 (★)

Un objet ponctuel A de masse m est suspendu à l'extrémité d'un fil de masse négligeable et de longueur L dont l'autre extrémité O est fixe. On ne considère pas les mouvements en dehors d'un plan vertical Oxy perpendiculaire à un axe Oz horizontal. On repère A par l'angle θ entre le fil et la verticale. On suppose que le référentiel terrestre est galiléen. On néglige tous les frottements. On désigne par $\vec{g} = g\vec{u}_x$ l'accélération du champ de pesanteur.

1. Déterminer l'équation différentielle du mouvement de A en appliquant le théorème de l'énergie cinétique.
2. Retrouver le résultat par le principe fondamental de la dynamique.
3. Tracer la courbe d'énergie potentielle et déterminer les positions d'équilibre du système et leur stabilité.
4. Initialement on abandonne A sans vitesse initiale alors que le fil est écarté d'un angle θ_0 . On se place dans le cas où θ_0 est petit. Montrer que le système constitue un oscillateur harmonique dont on exprimera la pulsation ω_0 et la période T_0 en fonction de g et L .
5. Compte tenu des conditions initiales, déterminer l'évolution de l'angle θ en fonction du temps t .
6. Quelle est la valeur maximale $v_{\max}$ de la vitesse de l'objet A au cours de son mouvement ? On l'exprimera en fonction de θ_0 , L et g .
7. Tracer la courbe donnant θ en fonction du temps.

17.2 Oscillations d'un pendule simple - partie 2 (★★)

On reprend le problème précédent mais on tient compte d'une force de frottement fluide proportionnelle à la vitesse $\vec{v}$. On note h le coefficient de proportionnalité.

1. Déterminer l'équation différentielle du mouvement de A en appliquant le théorème de l'énergie cinétique.
2. Les frottements sont suffisamment faibles pour que le régime soit pseudo-périodique. Donner la condition sur h pour qu'il en soit ainsi.
3. On lâche le pendule sans vitesse initiale d'une position faisant un angle θ_0 faible. Linéariser l'équation différentielle du mouvement puis déterminer l'expression de θ en fonction du temps t .
4. Donner l'allure de la courbe représentant θ en fonction du temps.

17.3 Mouvement d'un anneau sur une piste circulaire d'après CCP 2004 (★★)

On considère le dispositif de la figure (17.11) où un objet assimilable à un point matériel M de masse m se déplace solidairement à une piste formée de deux parties circulaires de rayon R_1 et R_2 , de centre C_1 et C_2 dans un plan vertical.

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

On repère la position de M par l'angle θ . Pour la partie (1), $\theta \in \left[-\frac{\pi}{2}, \pi\right]$ et pour la partie (2), $\theta \in \left[\pi, \frac{5\pi}{2}\right]$. Il n'y a pas de frottement et on note g l'accélération de la pesanteur.

1. Exprimer l'énergie potentielle de pesanteur E_p du point M en supposant $E_p = 0$ au point B ($\theta = \pi$). On distinguera les cas (1) et (2).
2. Tracer l'allure de $E_p(\theta)$.
3. Déterminer les positions angulaires d'équilibre et leur stabilité.

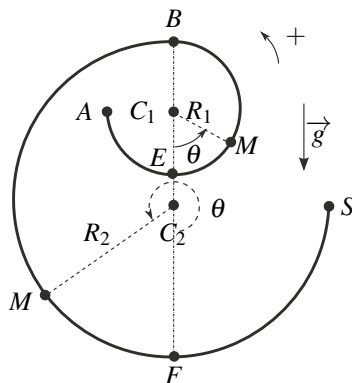


Figure 17.11

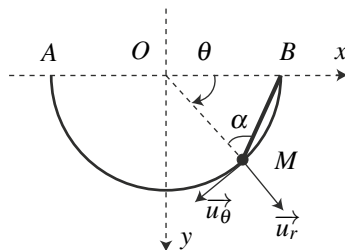
4. L'anneau est initialement en A ($\theta = -\pi/2$), il est lancé à une vitesse V_0 dans le sens trigonométrique.

- a. A quelle condition sur la vitesse V_0 l'anneau peut-il atteindre le point F en $\theta = 2\pi$?
- b. Cette condition étant remplie, donner l'expression de sa vitesse V_F en fonction des données du problèmes.
- c. A quelle condition sur V_0 , l'anneau sort-il de la piste en S en $\theta = 5\pi/2$.

17.4 Mouvement d'une bille reliée à un ressort sur un cercle (★★)

On considère le mouvement d'une bille M de masse m pouvant coulisser sans frottement sur un cerceau de centre O et de rayon R disposé dans un plan vertical. On note AB le diamètre horizontal du cerceau, Ox l'axe horizontal, Oy l'axe vertical descendant et θ l'angle entre Ox et OM . La bille est attachée à un ressort de longueur à vide nulle et de raideur k dont la seconde extrémité est fixée en B . Elle ne peut se déplacer que sur le demi-cercle inférieur.

1. Déterminer l'angle α entre MO et MB en fonction de θ .
2. Etablir l'expression de la longueur de ressort en fonction de R et θ .
3. En déduire l'énergie potentielle totale du système. Représenter la courbe d'énergie potentielle et en déduire les positions d'équilibres éventuelles et leur stabilité.



4. Si l'on écarte faiblement la bille de sa position d'équilibre stable θ_e et qu'on la lâche sans vitesse initiale, à quel type de mouvement peut-on s'attendre ?
5. Etablir l'équation différentielle du mouvement.
6. On note ε l'écart $\theta - \theta_e$. Initialement, on écarte la bille d'un angle $\varepsilon_0 \ll \frac{\pi}{2}$ à partir de sa position d'équilibre et on la lâche sans vitesse initiale. Linéariser l'équation du mouvement et conclure.

APPROFONDIR

17.5 Etude d'un oscillateur à l'aide de son portrait de phase (★)

On fait l'étude d'un oscillateur M de masse $m = 0,2 \text{ kg}$ astreint à se déplacer suivant l'axe Ox de vecteur unitaire $\vec{u}_x$. Il est soumis uniquement aux forces suivantes :

- La force de rappel d'un ressort de caractéristiques (k, ℓ_0) .
- Une force de frottement visqueux : $\vec{f}_v = -\lambda \dot{x} \vec{u}_x$.
- Une force constante $\vec{F}_c = F_c \vec{u}_x$

1. Equation du mouvement.

a. Etablir l'équation différentielle du mouvement de M et la mettre sous la forme :

$$\ddot{x} + \frac{\omega_0}{Q} \dot{x} + \omega_0^2 x = \omega_0^2 X_0$$

où x est l'allongement du ressort (par rapport à ℓ_0). Les grandeurs ω_0 , Q et X_0 sont à exprimer en fonction des données.

b. Dans le cas d'une solution pseudo-périodique, exprimer $x(t)$: on définira le temps caractéristique de décroissance des oscillations τ et la pseudo-pulsation Ω que l'on exprimera en fonction de ω_0 et Q .

2. Le portrait de phase ($v(t) = \dot{x}(t), x(t)$) de l'oscillateur étudié est donné sur la figure (17.12). On souhaite pouvoir en tirer les valeurs des différents paramètres de l'oscillateur.

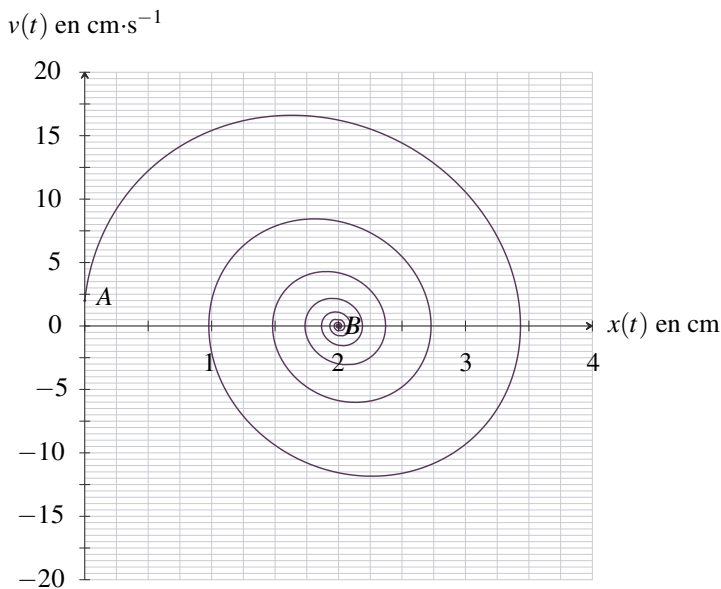


Figure 17.12

a. Quel est le type de mouvement ?

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

- b. Déterminer la vitesse et l'élongation au début et à la fin du mouvement.
- c. Déterminer la vitesse maximale atteinte ainsi que l'élongation maximale.
- d. On donne les différentes dates correspondant aux croisements de la trajectoire de phase avec l'axe des abscisses :

$t \text{ (s)}$	0,31	0,65	0,97	1,3	1,62
-----------------	------	------	------	-----	------

En déduire le pseudo-période T et la pseudo-pulsation Ω .

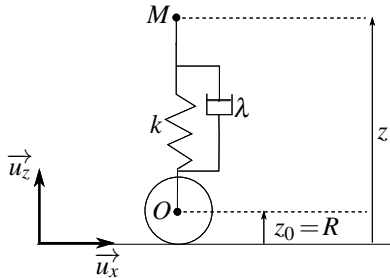
- e. On définit le décrément logarithmique par $\delta = \frac{1}{n} \ln \left(\frac{x(t) - x_B}{x(t+nT) - x_B} \right)$ où $x(t)$ et $x(t+nT)$ sont les élongations aux instants t et $t+nT$ (n entier naturel) et x_B l'élongation finale de M . Exprimer δ en fonction de T et τ .

En choisissant une valeur de n la plus grande possible pour les données dont on dispose, déterminer δ puis τ .

- f. Déduire des résultats précédents le facteur de qualité Q et la pulsation propre ω_0 .
- g. Déterminer la raideur du ressort k , le coefficient de frottement λ et la force F_c sachant que $\ell_0 = 1 \text{ cm}$.

17.6 Modélisation d'un amortisseur (★★)

On considère l'amortisseur d'un véhicule. Chaque roue supporte un quart de la masse de la voiture assimilé à un point M de masse $m = 500 \text{ kg}$, et est reliée à un amortisseur dont le ressort a une constante de raideur $k = 2,5 \cdot 10^4 \text{ N} \cdot \text{m}^{-1}$. Le point M subit aussi un frottement visqueux $\vec{f}_v = -\lambda \vec{v}$ où $\vec{v}$ est la vitesse verticale de M et $\lambda = 5,0 \cdot 10^3 \text{ kg} \cdot \text{s}^{-1}$.



1. Le véhicule franchit à vitesse constante un défaut de la chaussée de hauteur $h = 5,0 \text{ cm}$. Son inertie est suffisante pour qu'il ne se soulève pas immédiatement mais acquiert une vitesse verticale $v_0 = 0,50 \text{ m} \cdot \text{s}^{-1}$. On pose $\alpha = \frac{\lambda}{2m}$. On note Z_i la cote du point M avant le passage du défaut.

- a. On note $Z(t)$ la cote de M . établir l'équation différentielle pour Z après le passage de l'obstacle. Déterminer $Z(t)$ en fonction des données. On remarquera que $\alpha = \Omega$ ou Ω est la pseudo-pulsation.

- b. Les passagers sont sensibles à l'accélération verticale de la voiture, calculer sa valeur maximale. On utilisera le fait que $\alpha = \Omega$.

2. Il faut en fait éviter des oscillations susceptibles de provoquer chez les passagers la mal des transports, en se plaçant dans les conditions critiques. Pour quelle masse par roue est-ce réalisé, k et λ étant inchangés ?

17.7 Gestion du recul d'un canon (★★)

On considère un canon (figure 17.13) de masse $M = 800 \text{ kg}$. Lors du tir horizontal d'un obus de masse $m = 2 \text{ kg}$ avec une vitesse $\vec{v}_0 = v_0 \vec{u}_x$ ($v_0 = 600 \text{ m.s}^{-1}$), le canon acquiert une vitesse de recul $\vec{v}_c = -\frac{m}{M} \vec{v}_0$.

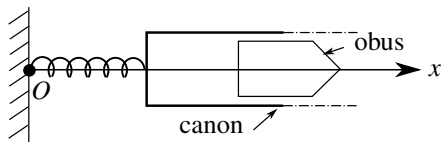


Figure 17.13

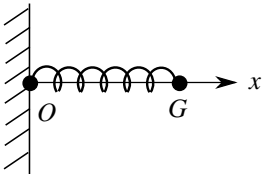
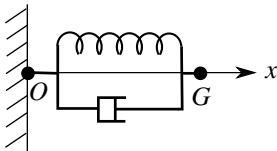


Figure 17.14

Pour limiter la course du canon, on utilise un ressort de raideur k_1 , de longueur à vide L_0 dont l'une des extrémités est fixe et l'autre lié au canon. Le déplacement a lieu suivant l'axe Ox . Dans la suite, le canon est assimilé à son centre de gravité G (figure 17.14).

1. Quelle est la longueur du ressort lorsque le canon est au repos ?
2. En utilisant l'énergie mécanique, déterminer la distance de recul d . En déduire la raideur k_1 pour avoir un recul au maximum égal à d . Application numérique pour $d = 1 \text{ m}$.
3. Retrouver la relation entre k_1 et d en appliquant la relation fondamentale de la dynamique.
4. Quel est l'inconvénient d'utiliser un ressort seul ?

Pour pallier ce problème, on ajoute au système un dispositif amortisseur, exerçant une force de frottement visqueux $\vec{F} = -\lambda \vec{v}$, $\vec{v}$ étant la vitesse du canon.



5. Le dispositif de freinage absorbe une fraction $E_a = 778 \text{ J}$ de l'énergie cinétique initiale. Calculer la nouvelle valeur k_2 de la constante de raideur du ressort avec les données numériques précédentes. Déterminer la pulsation propre ω_0 de l'oscillateur.
6. Déterminer λ pour que le régime soit critique. Application numérique.
7. Déterminer l'expression de l'élongation $x(t)$ du ressort, ainsi que celle de la vitesse $\dot{x}(t)$. En déduire l'instant t_m pour lequel le recul est maximal. Exprimer alors ce recul en fonction de m , v_0 et λ . L'application numérique redonne-t-elle la valeur de d précédente ?

17.8 Interaction entre atomes de gaz nobles (★★★)

On veut comprendre l'évolution de quelques propriétés physiques des gaz nobles : hélium, néon, argon, krypton et xénon notés par la suite par leurs symboles chimiques : He, Ne, Ar, Kr et Xe. Ces corps sont sous forme gazeuse à pression et température ambiante et liquide à des températures très basses. Le tableau 17.1 récapitule certaines de leurs caractéristiques.

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

Pour modéliser l’interaction entre deux atomes d’un gaz noble, on utilise le potentiel de Lennard-Jones (1929) dont l’expression est :

$$E_p = 4\epsilon \left(\left(\frac{\sigma}{r} \right)^{12} - \left(\frac{\sigma}{r} \right)^6 \right),$$

où r représente la distance entre les atomes. Les valeurs de ϵ et de σ sont regroupées dans le tableau 17.1. On a l’habitude de donner la valeur de σ en $\text{\AA} = 10^{-10} \text{ m}$ et celle de ϵ sous la forme $\frac{\epsilon}{k_B}$ qui s’exprime en K avec $k_B = 1,38.10^{-23} \text{ J}\cdot\text{K}^{-1}$, la constante de Boltzman. Dans ce potentiel, le terme en $\frac{1}{r^6}$ correspond aux interactions de Van der Waals ; le terme en $\frac{1}{r^{12}}$ est un potentiel empirique qui traduit le fait que les atomes ne peuvent pas s’interpénétrer. On

atome	He	Ne	Ar	Kr	Xe
Masse molaire ($\text{g}\cdot\text{mol}^{-1}$)	4,0	20	40	83,8	131,3
masse volumique ($\text{kg}\cdot\text{m}^{-3}$)	125	1207	1393	2413	3057
température d’ébullition (K)	4,2	27,1	87,3	119,9	165,0
σ ($\text{\AA}$)	2,63	2,77	3,40	3,60	4,06
ϵ/k_B (K)	5,46	36,8	116,8	164,6	218,2

Tableau 17.1

veut exploiter ce potentiel pour comprendre l’évolution de quelques propriétés physiques des gaz nobles. Pour cela, on considère le modèle suivant :

- l’interaction entre plus proches voisins dérive de l’énergie potentielle de Lennard-Jones tracée figure 17.15 ;
- un atome M de masse m est en interaction avec ses plus proches voisins. Pour simplifier le traitement mathématique, on ne considère que le plus proche voisin distant de r ;
- en phase liquide, les atomes sont mobiles mais ne peuvent pas sortir du puits de potentiel créé par son plus proche voisin ;
- en phase gazeuse, les atomes sont mobiles et libres de s’éloigner indéfiniment ;
- l’énergie cinétique des atomes est reliée à la température par la relation $E_c = \frac{3}{2}k_B T$.

On rappelle la valeur du nombre d’Avogadro : $\mathcal{N}_a = 6,02.10^{23} \text{ mol}^{-1}$.

1. Etude de la température d’ébullition.

a. Exploiter la relation $E_c = \frac{3}{2}k_B T$ pour estimer la température permettant à l’atome étudié de sortir du puits de potentiel.

b. Justifier que cette température est une bonne estimation de la température d’ébullition. La calculer pour chacun des corps considérés. Conclure.

2. Etude de la courbe d’énergie potentielle.

a. La courbe d’énergie potentielle de la figure 17.15 admet un minimum. Déterminer sa position $r_{\min}$ et sa valeur $E_{\min}$.

- b. Déterminer la force d'interaction $\vec{f}$ qui dérive de cette énergie potentielle.
 - c. Dans quelle domaine de distance cette force est-elle attractive, répulsive, nulle ?
3. Etude de la masse volumique.
- a. Justifier que la distance inter-atomique est proche de $r_{\min}$ en phase liquide. En déduire le nombre d'atome par unité de volume en phase liquide.
 - b. En déduire la masse volumique de la phase liquide de chacun des corps considérés. Conclure.
4. Etude qualitative de l'évolution de cette masse volumique avec la température.
- a. On donne le diagramme de phase d'une particule mobile dans le potentiel de Lennard-Jones sur la figure 17.15. Décrire l'évolution des trajectoires de phase lorsque l'énergie mécanique du système augmente ?
 - b. A quel type d'état correspondent les trajectoires de phase fermées ? ouvertes ?
 - c. Dans le cas de trajectoires de phases fermées, comment évolue la distance moyenne du système à son plus proche voisin ?
 - d. En utilisant la relation énergie cinétique - température vue précédemment, quelle prédiction qualitative peut-on faire sur l'évolution de la distance moyenne entre atome dans un liquide en fonction de la température ?
 - e. En déduire l'évolution qualitative de la masse volumique de la phase liquide avec la température prédite par ce modèle ?

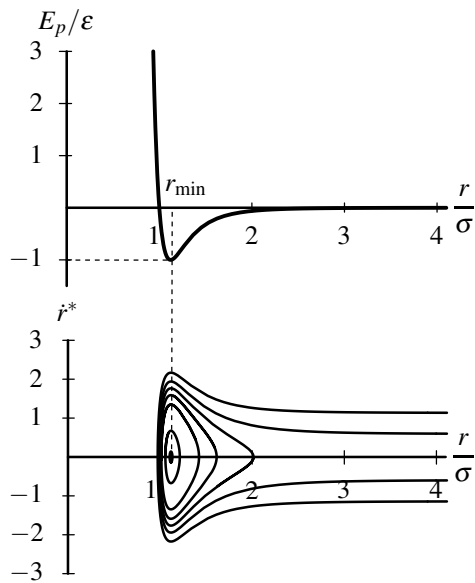


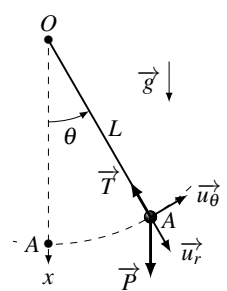
Figure 17.15

CORRIGÉS

17.1 Oscillations d'un pendule simple - partie 1

1. On étudie le mouvement du point A de masse m dans le référentiel terrestre galiléen. Le mouvement du point A est situé dans un plan vertical (Oxy) et la distance de A à O est fixée à L donc le point A se déplace sur un cercle de rayon L et de centre O . On peut alors appliquer les formules de cinématique du point en coordonnées polaires de centre O :

$$\vec{OA} = L\vec{u}_r \quad ; \quad \vec{v} = L\dot{\theta}\vec{u}_\theta \quad ; \quad \vec{a} = -\frac{v^2}{L}\vec{u}_r + L\ddot{\theta}\vec{u}_\theta.$$



Ce point est soumis à son poids et à la tension du fil $\vec{T}$. Le poids est une force conservative, la tension du fil est perpendiculaire au mouvement donc sa puissance est nulle. On est donc dans un cas de conservation de l'énergie mécanique.

On détermine l'énergie potentielle de pesanteur $E_p = mg$ altitude = $-mgL\cos\theta$. L'énergie cinétique vaut $E_c = \frac{1}{2}mv^2 = \frac{1}{2}mL^2\dot{\theta}^2$. On en déduit l'énergie mécanique :

$$E_m = \frac{1}{2}mL^2\dot{\theta}^2 - mgL\cos\theta.$$

L'énergie mécanique étant conservée, sa dérivée est nulle donc

$$\frac{dE_m}{dt} = mL^2\ddot{\theta}\dot{\theta} + mgL\sin\theta\dot{\theta} = 0.$$

Comme, lors d'un mouvement, $\dot{\theta}$ n'est pas nulle à tout instant, on en déduit l'équation du mouvement :

$$\ddot{\theta} + \frac{g}{L}\sin\theta \Leftrightarrow \ddot{\theta} + \omega_0^2\sin\theta = 0 \quad \text{avec } \omega_0^2 = \frac{g}{L}.$$

2. Le principe fondamental de la dynamique projeté sur $\vec{u}_\theta$ en coordonnées polaires donne $mL\ddot{\theta} = -mg\sin\theta$. On retrouve donc la même équation.

3. Le tracé de l'énergie potentielle donne :

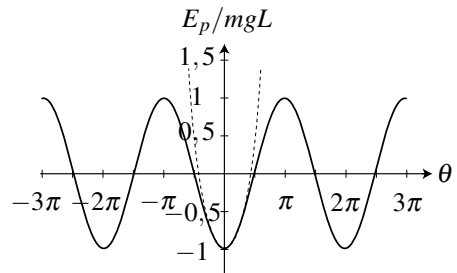


Figure 17.16 – Energie potentielle d'un pendule pesant. En pointillé, l'énergie mécanique minimale nécessaire pour obtenir un mouvement de révolution.

L'énergie potentielle présente donc un minimum en $\theta = 0$ (modulo 2π) qui correspond à une position d'équilibre stable. En cette position, le point A est à la verticale de O en dessous.

L'énergie potentielle présente donc un maximum en $\theta = \pi$ (modulo 2π) qui correspond à une position d'équilibre instable. En cette position, le point A est à la verticale de O au dessus. Cette position d'équilibre n'a donc pas de sens physique car un fil sans rigidité fournit une traction et ne peut en aucun cas soutenir la masse A en s'opposant au poids dans cette position.

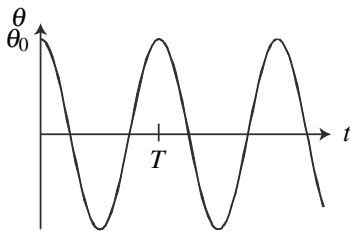
4. En utilisant l'approximation $\sin \theta \simeq \theta$, on déduit $\ddot{\theta} + \frac{g}{L}\theta = 0$, ce qui correspond à l'équation du mouvement d'un oscillateur harmonique de pulsation $\omega_0 = \sqrt{\frac{g}{L}}$ et de période $T_0 = 2\pi\sqrt{\frac{L}{g}}$.

Du point de vue énergétique, cela revient à remplacer l'énergie potentielle E_p par le potentiel harmonique tangent dessiné en pointillé sur la courbe d'énergie potentielle.

5. La solution s'écrit $\theta = A \cos(\omega_0 t) + B \sin(\omega_0 t)$. La détermination des constantes A et B s'obtient à l'aide des conditions initiales soit $\theta(0) = A = \theta_0$ et $\dot{\theta}(0) = B\omega_0 = 0$. Finalement : $\theta = \theta_0 \cos(\omega_0 t)$.

6. La vitesse s'écrit $v = L\dot{\theta}$. Elle est donc maximale si $\dot{\theta}$ est maximal donc lorsque $\dot{\theta} = \theta_0\omega_0$. Elle vaut alors $v_{\max} = L\theta_0\omega_0 = \theta_0\sqrt{Lg}$.

7. L'évolution de θ avec t est de la forme ci-dessous :



17.2 Oscillations d'un pendule simple - partie 2

1. On doit ajouter la force de frottement $\vec{F}_f = -h\vec{v}$ au bilan des forces de l'exercice précédent. Cette force n'est pas conservative et sa puissance vaut $\mathcal{P}(\vec{F}_f) = \vec{F}_f \cdot \vec{v} = -hv^2$. Comme dans l'exercice précédent, la puissance de la tension du fil est nulle. La puissance du poids vaut $\mathcal{P}(\vec{P}) = -\frac{dE_p}{dt}$ car $dE_p = -\delta W(\vec{P}) = -\mathcal{P}(\vec{P})dt$. On applique alors le théorème de la puissance cinétique :

$$\frac{dE_c}{dt} = -\frac{dE_p}{dt} - hv^2 \quad \Leftrightarrow \quad \frac{dE_m}{dt} = -hv^2 = -hL^2\dot{\theta}^2.$$

On a vu dans l'exercice précédent que $\frac{dE_m}{dt} = mL^2\ddot{\theta}\dot{\theta} + mgL\sin\theta\dot{\theta}$, on en déduit l'équation différentielle du mouvement en simplifiant par $\dot{\theta}$ qui n'est pas nul à tout instant :

CHAPITRE 17 – MOUVEMENT DANS UN PUIT DE POTENTIEL

$$\ddot{\theta} + \frac{h}{m} \dot{\theta} + \frac{g}{L} \sin \theta = 0.$$

2. Lorsque θ est faible, on obtient $\ddot{\theta} + \frac{h}{m} \dot{\theta} + \frac{g}{L} \theta = 0$. Son équation caractéristique s'écrit $r^2 + \frac{h}{m} r + \frac{g}{L} = 0$ et son discriminant $\Delta = \left(\frac{h}{m}\right)^2 - 4 \frac{g}{L}$. On a un régime pseudo-périodique si $\Delta < 0$ soit $h < 2m\sqrt{\frac{g}{L}}$.

3. La solution s'écrit $\theta = \left(A \cos \sqrt{\frac{g}{L} - \left(\frac{h}{2m}\right)^2} t + B \sin \sqrt{\frac{g}{L} - \left(\frac{h}{2m}\right)^2} t \right) \exp\left(-\frac{h}{2m} t\right)$.

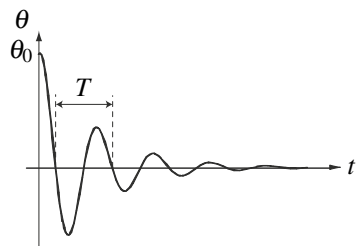
On détermine les constantes A et B avec les conditions initiales :

$$\begin{cases} \theta(0) = A = \theta_0 \\ \dot{\theta}(0) = \frac{B}{2} \sqrt{\frac{g}{L} - \left(\frac{h}{2m}\right)^2} - \frac{h}{2m} A = 0. \end{cases}$$

Finalement on a :

$$\theta = \left(\cos \sqrt{\frac{g}{L} - \left(\frac{h}{2m}\right)^2} t + h \sqrt{\frac{L}{4gm^2 - Lh^2}} \sin \sqrt{\frac{g}{L} - \left(\frac{h}{2m}\right)^2} t \right) \theta_0 \exp\left(-\frac{h}{2m} t\right).$$

4. L'évolution de θ avec t est de la forme ci-dessous :

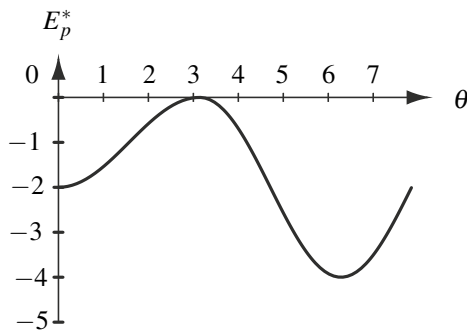


17.3 Mouvement d'un anneau sur une piste circulaire

Le référentiel du laboratoire est supposé galiléen.

1. On prend un axe Oz vertical ascendant d'origine C_2 . L'énergie potentielle est alors donnée par l'expression $E_p = mgz + \text{cte}$. L'énoncé propose de prendre l'énergie potentielle nulle en B soit avec l'origine choisie, pour $z = R_2$, d'où l'expression de l'énergie $E_p = mg(z - R_2)$. On cherche maintenant à exprimer E_p en fonction de θ . Pour la partie (1), on peut écrire $z = R_2 - R_1 - R_1 \cos \theta$ d'où $E_{p1} = -mgR_1(1 + \cos \theta)$. Pour la partie (2), on trouve $z = -R_2 \cos \theta$, soit $E_{p2} = -mgR_2(1 + \cos \theta)$.

2. Les résultats précédents permettent de tracer la courbe $E_p(\theta)$. On choisit l'échelle d'énergie $E_0 = mgR_1$ et pour fixer les idées, on représente $E_p^* = \frac{E_p}{E_0}$ pour $R_2 = 2R_1$.



3. Les positions angulaires d'équilibre apparaissent sur la figure ci-dessus : ce sont les minima (positions stables) en $\theta = 0$ et 2π et le maximum en $\theta = \pi$.

On les retrouve par le calcul : $\frac{dE_p}{d\theta} = mgR_i \sin \theta$, l'indice i étant égal à 1 ou 2 suivant la zone.

On trouve donc des extrémums pour $\sin \theta = 0$. Dans l'intervalle $\left[-\frac{\pi}{2}, \frac{5\pi}{2}\right]$, les solutions sont $\theta \in \{0, \pi, 2\pi\}$.

La dérivée seconde de E_p est égale à $mgR_i \cos \theta$, positive pour $\theta = 0$ et 2π correspondant aux positions stables et négative pour $\theta = \pi$ correspondant à la position instable.

4. Sortie en S :

a. Pour atteindre le puits de potentiel en F depuis A , l'anneau doit passer la barrière de potentiel en B . Un mouvement est possible si l'énergie mécanique vérifie $E_m \geq E_p$, il faut donc donner au point M lorsqu'il est en A une énergie mécanique au moins égale à $E_p(B) = 0$ soit :

$$E_m(A) = E_c(A) + E_p(A) = \frac{1}{2}mV_0^2 - mgR_1,$$

d'où :

$$E_m(A) \geq 0 \Rightarrow V_0 \geq \sqrt{2gR_1}.$$

b. On écrit l'égalité de l'énergie mécanique en A et en F :

$$\frac{1}{2}mV_0^2 - mgR_1 = \frac{1}{2}mV_F^2 - 2mgR_2 \Rightarrow V_F \sqrt{V_0^2 - 2g(R_1 - 2R_2)}.$$

c. Pour sortir de l'anneau en S , il faut que l'anneau ait une énergie mécanique supérieure à $E_p(S)$, ce qui est le cas s'il franchit la barrière de potentiel en B . Donc la condition est la même que pour atteindre F .

17.4 Mouvement d'une bille reliée à un ressort

On étudie la bille dans le référentiel terrestre galiléen. Elle est soumise à son poids, à la force de rappel élastique $\vec{T}$ du ressort et à la réaction $\vec{N}$ du cerceau qui est normale du fait de l'absence de frottement.

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

1. Le triangle OMB est isocèle en O donc les angles des sommets B et M sont égaux. Par ailleurs, la somme des angles d'un triangle vaut π . En explicitant ces deux conditions, on obtient $\alpha = \frac{\pi - \theta}{2}$.

2. Pour calculer la distance MB , on détermine son carré :

$$MB^2 = (\vec{MO} + \vec{OB})^2 = R^2 + R^2 + 2R^2 \cos(\pi - \theta) = 4R^2 \sin^2 \frac{\theta}{2}.$$

On en déduit $MB = 2R \left| \sin \frac{\theta}{2} \right|$.

3. Les forces qui s'appliquent sur le système sont conservatives (poids, force de rappel élastique) ou à puissance nulle (réaction du cerceau). On est donc dans un cas de conservation de l'énergie mécanique.

On détermine l'énergie potentielle dont dérive le poids : $E_{p1} = mgy_M = -mgR \sin \theta$,

et celle dont dérive la force de rappel élastique : $E_{p2} = \frac{1}{2}k\Delta\ell^2 = \frac{1}{2}kMB^2 = 2kR^2 \sin^2 \frac{\theta}{2}$.

L'énergie potentielle totale vaut donc :

$$E_p = -mgR \sin \theta + 2kR^2 \sin^2 \frac{\theta}{2}.$$

Pour la représenter, on remarque que θ ne peut varier qu'entre 0 et π et on fixe $E_0 = mgR$ comme échelle d'énergie. On trace l'énergie potentielle sans dimension :

$$E_p^* = \frac{E_p}{E_0} = -\sin \theta + p \sin^2 \frac{\theta}{2} \quad \text{avec} \quad p = \frac{2kR^2}{mgR}.$$

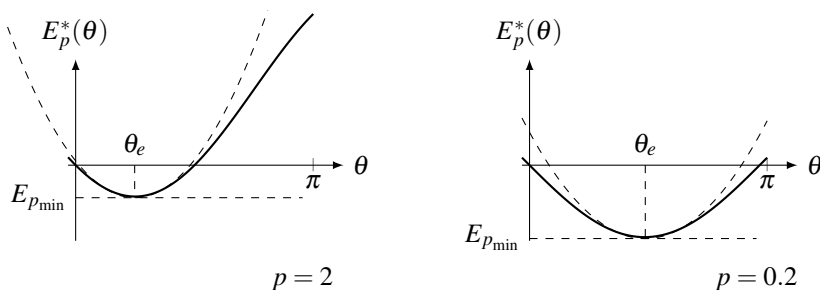


Figure 17.17 – Tracé de l'énergie potentielle pour deux valeurs de p .

L'énergie potentielle présente un minimum en θ_e entre 0 et π . La position de coordonnée θ_e est donc une position d'équilibre stable que l'on peut déterminer en recherchant le point d'annulation de la dérivée :

$$\left(\frac{dE_p}{d\theta} \right)_{(\theta=\theta_e)} = 0 \Rightarrow \cos \theta_e - p \sin \frac{\theta_e}{2} \cos \frac{\theta_e}{2} = 0 \Rightarrow \cos \theta_e - \frac{p}{2} \sin \theta_e = 0,$$

soit :

$$\tan \theta_e = \frac{2}{p} = \frac{mg}{kR}.$$

La position d'équilibre est, comme on pouvait s'y attendre comprise entre 0 et $\frac{\pi}{2}$. Elle tend vers 0 lorsque la raideur du ressort est si grande que le poids de M ne peut pas l'étirer et vers $\frac{\pi}{2}$ lorsqu'elle est si faible que le poids de M l'étire facilement.

4. Si l'on écarte la bille de sa position d'équilibre stable et qu'on la lâche sans vitesse initiale, la bille va osciller dans le puits de potentiel. Si on l'en écarte faiblement, on s'attend à observer des oscillations harmoniques, tout se passant comme si la bille oscillait dans le potentiel harmonique tangent dessiné en pointillé sur la figure ci-dessus.

5. L'énergie cinétique du système vaut $E_c = \frac{1}{2}mv^2 = \frac{1}{2}mR^2\dot{\theta}^2$ puis l'énergie mécanique :

$$E_m = \frac{1}{2}mR^2\dot{\theta}^2 - mgR \sin \theta + 2kR^2 \sin^2 \frac{\theta}{2}.$$

Le mouvement étant conservatif, l'énergie mécanique est conservée et sa dérivée s'annule :

$$mR^2\dot{\theta}\ddot{\theta} - mgR \cos \theta \dot{\theta} + 4kR^2 \sin \frac{\theta}{2} \cos \frac{\theta}{2} \dot{\theta} = 0 \quad \Rightarrow \quad \ddot{\theta} - \frac{g}{R} \left(\cos \theta - \frac{p}{2} \sin \theta \right) = 0.$$

6. On écarte M de sa position d'équilibre et on pose $\theta = \theta_e + \varepsilon$ puis :

$$\begin{aligned} \cos \theta &= \cos(\theta_e + \varepsilon) = \cos \theta_e \cos \varepsilon - \sin \theta_e \sin \varepsilon = \cos \theta_e - \varepsilon \sin \theta_e \\ \sin \theta &= \sin(\theta_e + \varepsilon) = \sin \theta_e \cos \varepsilon + \cos \theta_e \sin \varepsilon = \sin \theta_e + \varepsilon \cos \theta_e, \end{aligned}$$

en utilisant les approximations $\cos \varepsilon \simeq 1$ et $\sin \varepsilon \simeq \varepsilon$.

On injecte alors ces relations dans l'équation du mouvement et on trouve :

$$\ddot{\theta} + \frac{g}{R} \left(\sin \theta_e + \frac{p}{2} \cos \theta_e \right) \varepsilon - \frac{g}{R} \left(\cos \theta_e - \frac{p}{2} \sin \theta_e \right) = 0.$$

$\ddot{\theta} = \ddot{\varepsilon}$ puisque ε et θ diffèrent d'une constante. Le terme $\left(\cos \theta_e - \frac{p}{2} \sin \theta_e \right) = 0$ d'après la question 3. Comme θ_e est compris entre 0 et $\frac{\pi}{2}$, son sinus et son cosinus sont positifs.

$\frac{g}{R} \left(\sin \theta_e + \frac{p}{2} \cos \theta_e \right)$ est une grandeur positive homogène à une pulsation au carré. On pose donc $\omega_0 = \sqrt{\frac{g}{R} \left(\sin \theta_e + \frac{p}{2} \cos \theta_e \right)}$ et l'équation du mouvement devient :

$$\ddot{\varepsilon} + \omega_0^2 \varepsilon = 0.$$

Comme prévu, le mouvement au voisinage d'une position d'équilibre stable est celui d'un oscillateur harmonique. Le calcul effectué permet d'en déterminer la pulsation ω_0 .

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

17.5 Etude d'un oscillateur à l'aide de son portrait de phase

1. Equation du mouvement

a. Dans le référentiel du laboratoire considéré galiléen, la relation fondamentale de la dynamique appliquée à l'oscillateur M donne en projection sur l'axe Ox :

$$m\ddot{x} = -k(x - \ell_0) - \lambda\dot{x} + F_c \Rightarrow \ddot{x} + \frac{\lambda}{m}\dot{x} + \frac{k}{m}x = \frac{k}{m}\ell_0 + \frac{F_c}{m}$$

On peut donc mettre l'équation sous la forme donnée par l'énoncé avec $\omega_0 = \sqrt{k/m}$, $Q = \sqrt{km}/\lambda$ et $X_0 = \ell_0 + F_c/k$.

b. Si la solution est pseudo-périodique, alors les solutions de l'équation caractéristique sont :

$$-\frac{\omega_0}{2Q} \pm i\omega_0\sqrt{1 - \frac{1}{4Q^2}}.$$

On peut définir alors la pseudo pulsation $\Omega = \omega_0\sqrt{1 - \frac{1}{4Q^2}}$ et le temps caractéristique de décroissance des oscillations $\tau = \frac{2Q}{\omega_0}$. L'expression de $x(t)$ est alors :

$$x(t) = A \exp(-t/\tau) \cos(\Omega t + \phi) + X_0$$

où A et ϕ sont des constantes dépendant des conditions initiales.

2. Etude du portrait de phase.

a. Le mobile oscille de part et d'autre de la position finale : il s'agit d'un mouvement pseudo périodique.

b. Au début du mouvement, le point représentatif de l'état du mobile est A et à la fin il s'agit du point B (on rappelle qu'un portrait de phase est décrit dans le sens horaire). Initialement le mobile est donc en $x = 0$ et sa vitesse est $v_0 = 2 \text{ cm}\cdot\text{s}^{-1}$. A la fin M est immobile en $x = 2 \text{ cm}$.

c. La vitesse maximale atteinte correspond au maximum de la courbe, soit $v = 16,5 \text{ cm}\cdot\text{s}^{-1}$ et l'élongation maximale correspond à l'intersection avec l'axe des abscisses la plus à droite soit $x_{\max} \simeq 3,4 \text{ cm}$.

d. Entre deux croisements avec l'axe des abscisses il se passe une demi pseudo-période. Pour avoir une valeur plus précise, on calcule les intervalles de temps entre chaque date donnée dans l'énoncé ce qui donne quatre valeurs de $T/2$ dont on fait la moyenne. Les quatre valeurs obtenues sont : $0,34 \text{ s}$; $0,32 \text{ s}$; $0,33 \text{ s}$ et $0,32 \text{ s}$; ce qui fait une moyenne de $0,3275 \text{ s}$ soit en arrondissant on a $T = 0,65 \text{ s}$ et $\Omega = 2\pi/T = 9,6 \text{ rad}\cdot\text{s}^{-1}$.

e. En utilisant l'expression générale de $x(t)$ déterminée au début de l'exercice et sachant que $x_B = X_0$ (position finale atteinte lorsque $t \rightarrow \infty$, on a :

$$\delta = \frac{1}{n} \ln \left(\frac{A \exp(-t/\tau) \cos(\Omega t + \phi)}{A \exp(-(t+nT)/\tau) \cos(\Omega(t+nT) + \phi)} \right) = \frac{1}{n} \ln \left(\exp(nT/\tau) \right) = \frac{T}{\tau}.$$

Corrigés

Corrigés

Corrigés

Corrigés

Corrigés

Corrigés

Corrigés



Corrigés

Corrigés

Corrigés

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

Avant l'obstacle, le point M se déplace à altitude constante, et la longueur du ressort vaut $Z_i - R$, on peut alors écrire l'équation d'équilibre :

$$0 = -mg - k(Z_i - R - \ell_0).$$

L'équation du mouvement devient alors :

$$\ddot{Z} + 2\alpha\dot{Z} + \omega_0^2 Z = \omega_0^2 (h + Z_i)$$

avec $\alpha = \lambda/2m = 5 \text{ rad.s}^{-1}$ et $\omega_0^2 = k/m = 50 \text{ rad}^2 \cdot \text{s}^{-2}$. Le discriminant de l'équation caractéristique $r^2 + 2\alpha r + \omega_0^2 = 0$ est $\Delta = 4\alpha^2 - 4\omega_0^2 = -100 \text{ rad}^2 \cdot \text{s}^{-2}$. Puisque $\Delta < 0$, le mouvement est pseudo-périodique. Les solutions de l'équation caractéristique sont $-\alpha \pm i\sqrt{-\Delta}/2$. La solution particulière de l'équation différentielle est $h + Z_i$. On en déduit l'expression de $Z(t)$:

$$Z(t) = A \exp(-\alpha t) \cos(\Omega t + \phi) + h + Z_i$$

avec $\Omega = \sqrt{-\Delta}/2 = 5 \text{ rad.s}^{-1}$, A et ϕ sont des constantes à déterminer avec les conditions initiales. A $t = 0$ on a $\dot{Z} = v_0$ et $Z(t) = Z_i$. Le calcul de $\dot{Z}$ donne $\dot{Z} = A \exp(-\alpha t) (-\alpha \cos(\Omega t + \phi) - \Omega \sin(\Omega t + \phi))$ soit avec $\Omega = \alpha$:

$$\dot{Z} = -A\Omega \exp(-\alpha t) (\cos(\Omega t + \phi) + \sin(\Omega t + \phi))$$

d'où les conditions initiales :

$$\begin{cases} A \cos \phi + h + Z_i &= Z_i \\ -A\Omega (\cos \phi + \sin \phi) &= v_0 \end{cases} \Rightarrow \begin{cases} A \cos \phi &= -h \\ A \sin \phi &= -\frac{v_0}{\Omega} + h \end{cases}$$

On en déduit $\tan \phi = \frac{v_0}{h\Omega} - 1 = 1$ soit $\phi = \pi/4$ et $A = -\sqrt{2}h$.

b. Pour déterminer la valeur maximale de l'accélération $\ddot{Z}(t)$ il faut dériver cette expression. On utilise à chaque étape de dérivation le fait que $\Omega = \alpha$. On part de la valeur de $\dot{Z}$ déjà calculée. Il vient :

$$\dot{Z} = A \exp(-\alpha t) ((\alpha^2 - \Omega^2) \cos(\Omega t + \phi) + 2\alpha\Omega \sin(\Omega t + \phi)) = 2A \exp(-\alpha t) \Omega^2 \sin(\Omega t + \phi)$$

puis :

$$\begin{aligned} \ddot{Z} &= 2A\Omega^2 \exp(-\alpha t) (-\alpha \sin(\Omega t + \phi) + \Omega \cos(\Omega t + \phi)) \\ &= 2A\Omega^3 \exp(-\alpha t) (-\sin(\Omega t + \phi) + \cos(\Omega t + \phi)) \end{aligned}$$

$\ddot{Z} = 0$ pour $\tan(\Omega t + \phi) = 1$. Le premier instant pour lequel $\ddot{Z}$ est maximale vérifie soit $\Omega t + \phi = \pi/4$ soit $t = 0$ puisque $\phi = \pi/4$.

On en déduit $\dot{Z}_{\max} = 2A\Omega^2 \sin \phi = 2,5 \text{ m.s}^{-2}$.

2. On se place dans les conditions critiques en rendant le discriminant nul soit $\Delta = 0$ entraîne $\alpha = \omega_0$ soit $m = \frac{\lambda^2}{4k}$. L'application numérique donne $m = 250 \text{ kg}$. Il faut énormément alléger le véhicule.

17.7 Gestion du recul d'un canon

Le référentiel terrestre est supposé galiléen. Le système étudié est le point G représentant le canon.

- 1. Les forces s'exerçant sur le canon sont : son poids (vertical), une réaction de support (verticale) et la tension $\vec{T}$ du ressort qui est la seule force horizontale. Au repos $\vec{T} = \vec{0}$ donc la longueur ℓ du ressort est égale à L_0 .
- 2. Le déplacement de G est horizontal, donc la seule force qui travaille est la tension ; cette force est conservative donc l'énergie mécanique est constante. La longueur du ressort étant égale à x , l'expression de l'énergie mécanique est :

$$E_m = \frac{1}{2}M\dot{x}^2 + \frac{1}{2}k_1(x - L_0)^2.$$

A l'instant initial, juste après le départ de l'obus, le canon n'a pas encore bougé donc $x = L_0$ et d'après l'énoncé $v = v_c$. Lorsque G atteint la distance de recul d sa vitesse est nulle et $x = L_0 - d$. On écrit l'égalité de l'énergie mécanique dans ces deux cas et on utilise $v_c = \frac{m}{M}v_0$ pour obtenir :

$$\frac{1}{2}Mv_c^2 = \frac{1}{2}k_1d^2 \Rightarrow d = \sqrt{\frac{M}{k_1}}v_c = \frac{m}{\sqrt{k_1M}}v_0 \Rightarrow k_1 = \frac{m^2v_0^2}{d^2M}.$$

L'application numérique donne $k = 1800 \text{ N.m}^{-1}$.

- 3. Si on applique la relation fondamentale de la dynamique en projection sur Ox , il vient : $M\ddot{x} = -k_1(x - L_0)$ soit $\ddot{x} + \omega_0^2x = \omega_0^2L_0$, avec $\omega_0 = \sqrt{k_1/M}$. La solution est :

$$x = A\cos(\omega_0t + \phi) + L_0.$$

A $t = 0$, $x = L_0$ ce qui donne $A\cos\phi = 0$ donc on choisit par exemple $\phi = -\pi/2$. Ainsi on peut écrire $x = A\sin\omega_0t + L_0$. Pour la vitesse, on calcule la dérivée : $\dot{x} = A\omega_0\cos\omega_0t$. A $t = 0$, $\dot{x} = -mv_0/M$ d'où $A = \frac{mv_0}{M\omega_0}$ et finalement :

$$x(t) = -\frac{mv_0}{M\omega_0}\sin\omega_0t + L_0 = -\frac{mv_0}{\sqrt{k_1M}}\sin\omega_0t + L_0.$$

La valeur absolue de l'amplitude du sinus est égale à d ce qui donne la même valeur que précédemment.

- 4. Comme le montre la solution de l'équation trouvée dans la question précédente, si l'on utilise un ressort seul, le canon va osciller et donc après le recul, il va repartir vers l'avant. L'amplitude va diminuer petit à petit à cause des frottements inéluctables mais le temps d'immobilisation sera important. On a donc intérêt à ajouter une force de frottement visqueux.
- 5. La variation d'énergie mécanique est maintenant égale à l'énergie absorbée par le dispositif de freinage, soit $\Delta E_m = -E_a$. Or initialement $E_m = E_c = Mv_c^2/2$ et finalement $E_m = E_p = k_2d^2/2$, d'où :

$$\frac{1}{2}k_2d^2 - \frac{1}{2}Mv_c^2 = -E_a \Rightarrow k_2 = \frac{1}{d^2}(Mv_c^2 - 2E_a) = \frac{1}{d^2}\left(\frac{m^2}{M}v_0^2 - 2E_a\right)$$

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

L'application numérique donne $k_2 = 244 \text{ N.m}^{-1}$. La pulsation propre est $\omega_0 = \sqrt{k_2/M}$ ce qui donne $0,55 \text{ rad.s}^{-1}$.

6. La relation fondamentale en projection sur Ox donne : $\ddot{x} + \frac{\lambda}{M}\dot{x} + \omega_0^2 x = \omega_0^2 L_0$.

Le discriminant de l'équation caractéristique $\Delta = \left(\frac{\lambda}{M}\right)^2 - 4\omega_0^2$ doit être nul pour le régime critique d'où $\lambda = 2M\omega_0$. L'application numérique donne $\lambda = 884 \text{ kg.s}^{-1}$.

7. La solution de l'équation est alors :

$$x(t) = \exp\left(-\frac{\lambda t}{2M}\right) (At + B) + L_0.$$

A $t = 0$, $x = L_0$ d'où $B = 0$. On calcule la vitesse : $\dot{x} = A \exp\left(-\frac{\lambda t}{2M}\right) \left(1 - \frac{\lambda}{2M}t\right)$.

Or à $t = 0$, $\dot{x} = v_c$ d'où $A = v_c$. Finalement on trouve :

$$x(t) = v_c t \exp\left(-\frac{\lambda t}{2M}\right) + L_0 = -\frac{m}{M} v_0 t \exp\left(-\frac{\lambda t}{2M}\right) + L_0.$$

Le recul est maximal lorsque la vitesse s'annule soit pour $t_m = 2M/\lambda = 1,8 \text{ s}$. On peut alors écrire la valeur de x à t_m :

$$x(t_m) = -\frac{m}{M} v_0 \frac{2M}{\lambda} e^{-1} + L_0 = -\frac{2mv_0}{\lambda e} + L_0 = L_0 - d \Rightarrow d = \frac{2mv_0}{\lambda e} = 1 \text{ m}.$$

On retrouve la valeur de d précédente.

17.8 Interaction entre atomes de gaz nobles

1. Etude de la courbe d'énergie potentielle.

a. Pour déterminer le minimum de la courbe d'énergie potentielle, on la dérive. Il est alors profitable de poser $u = \frac{\sigma}{r}$ et $E_p^* = \frac{E_p}{\varepsilon}$. On obtient alors : $E_p^* = 4(u^{12} - u^6)$. On remarque alors que :

$$\frac{dE_p^*}{dr} = \frac{dE_p^*}{du} \frac{du}{dr} = -\frac{\sigma}{r^2} \frac{dE_p^*}{du}.$$

Comme $\frac{\sigma}{r^2}$ est non nul, l'annulation de $\frac{dE_p^*}{dr}$ correspond donc à l'annulation de $\frac{dE_p^*}{du}$ qui est beaucoup plus simple à calculer :

$$\frac{dE_p^*}{du} = 4(12u^{11} - 6u^5) \quad \text{donc} \quad \frac{dE_p^*}{du} = 0 \Leftrightarrow 2u^6 - 1 = 0 \quad \text{ou} \quad u = 0.$$

La solution $u = 0$ qui correspond à $r \rightarrow \infty$ n'est pas la solution recherchée. Finalement :

$$\frac{1}{u^6} = 2 \Rightarrow r_{\min} = 2^{\frac{1}{6}} \sigma \quad \text{et} \quad E_{\min} = E_p(r_{\min}) = -\varepsilon.$$

b. On utilise le calcul précédent en appliquant la relation liant la force à l'énergie potentielle :

$$\vec{f} = -\frac{dE_p}{dr} \vec{u}_r = 4\epsilon \frac{\sigma}{r^2} (12u^{11} - 6u^5) \vec{u}_r = 24 \frac{\epsilon}{\sigma} \left(\frac{\sigma}{r}\right)^7 \left(2\left(\frac{\sigma}{r}\right)^6 - 1\right) \vec{u}_r.$$

c. La force est dirigée vers le minimum d'énergie potentielle. Elle est donc répulsive (selon $+\vec{u}_r$) à courte distance lorsque $r < r_{\min}$ et attractive (selon $-\vec{u}_r$) à grande distance lorsque $r > r_{\min}$. Ce résultat de cours permet de vérifier le calcul ci-dessus. A longue distance, le terme en $\frac{1}{r^6}$ correspondant aux force de Van der Waals attractives dominent le terme répulsif en $\frac{1}{r^{12}}$.

2. Etude de la température d'ébullition.

a. Pour qu'un atome puisse sortir du puits de potentiel, il faut que son énergie mécanique soit supérieure au maximum d'énergie potentielle, soit ici 0. Si un atome possède une énergie cinétique supérieure à la profondeur du puits de potentiel ϵ , alors il possède nécessairement l'énergie mécanique lui permettant de sortir de ce puits. La température correspondante est T_1 telle que :

$$\sigma = \frac{3}{2} k_B T_1 \quad \text{soit} \quad T_1 = \frac{2}{3} \frac{\sigma}{k_B}.$$

b. A cette température, les atomes sont libres de s'éloigner indéfiniment les uns des autres et d'occuper tout le volume dont ils disposent. Ils forment donc un gaz. On regroupe les températures T_1 dans le tableau 17.2. On observe que la température T_1 et son évolution selon la nature du gaz noble étudié est en bon accord avec les températures d'ébullition mesurées. Le modèle explique correctement les données expérimentales.

3. Etude de la masse volumique.

a. En phase liquide, le modèle précise que les atomes sont piégés dans le puits de potentiel. Leur mobilité est donc restreinte et, d'après la courbe 17.15, r est compris entre σ et environ 2σ . La distance inter-atomique est donc proche de $r_{\min}$. Chaque atome occupe un volume de l'ordre de $r_{\min}^3$. Le nombre d'atome par unité de volume est donc $n^* = \frac{1}{r_{\min}^3}$.

b. La masse volumique est égale au produit du nombre d'atome par unité de volume n^* par la masse m^* d'un atome avec $m^* = \frac{M}{\mathcal{N}_A}$. Pour l'application numérique, on n'oublie pas que l'unité du système international pour la masse est le kg et on reporte les valeurs dans le tableau 17.2. Là encore, l'accord entre modèle et expérience est satisfaisant mais on voit que le modèle est incomplet car il s'éloigne des valeurs expérimentales pour l'hélium et le xénon.

4. Etude qualitative de l'évolution de cette masse volumique avec la température.

a. Sur diagramme de phase de la figure 17.15, on observe des trajectoires fermées correspondant à des mouvements périodiques dans un état lié. Lorsque l'énergie mécanique est faible, les trajectoires de phases sont quasiment elliptiques au voisinage de $r_{\min}$ puis, lorsque l'énergie mécanique augmente, elles se déforment et s'allongent vers les r croissants. Lorsque l'énergie mécanique est encore plus grande, les trajectoires de phases ne sont plus fermées.

CHAPITRE 17 – MOUVEMENT DANS UN Puits DE POTENTIEL

- b. Les trajectoires de phase fermées correspondent à des états liés (et donc à un liquide d’après ce qui précède) tandis que les trajectoires ouvertes correspondent à des états de diffusion (et donc à un gaz).
- c. Dans le cas de trajectoires de phases fermées, la trajectoire de phase s’étire vers les r croissants lorsque l’énergie mécanique augmente. En moyenne dans le temps, r est donc plus grand et de plus en plus grand à mesure que l’énergie mécanique augmente.
- d. Une augmentation de la température entraîne une augmentation de l’énergie cinétique donc une augmentation de l’énergie mécanique moyenne d’un atome. D’après la question précédente, la distance moyenne entre les atomes d’un liquide augmente donc lorsque la température augmente.
- e. En conséquence, le modèle prévoit une diminution de la masse volumique de la phase liquide avec la température. Ce phénomène s’appelle la dilatation thermique.

atome	He	Ne	Ar	Kr	Xe
masse volumique ($\text{t}\cdot\text{m}^{-3}$)	0,13	1,2	1,4	2,4	3,1
masse volumique modèle ($\text{t}\cdot\text{m}^{-3}$)	0,26	1,1	1,2	2,1	2,3
température d’ébullition (K)	4,2	27,1	87,3	119,9	165,0
T_1 (K)	3.6	25	78	110	146

Tableau 17.2

Mouvement d'une particule chargée dans un champ électrique ou magnétique

18

Le mouvement des particules chargées dans un champ électrique et/ou magnétique est un sujet important du fait du grand nombre d'applications qui l'utilisent. On se limite aux champs uniformes et indépendants du temps qui ne dépendent ni de la position dans l'espace ni de l'instant considérés. On exclut tout développement relativiste, ce qui implique que la vitesse des particules soit négligeable devant la vitesse de la lumière. Cette condition est restrictive car les particules accélérées par des champs électriques peuvent facilement atteindre des vitesses proches de celle de la lumière. Des effets relativistes permettent alors généralement d'expliquer les écarts entre le modèle théorique basé sur la mécanique classique et les observations expérimentales.

1 Force de Lorentz

1.1 Rappel de l'expression

Comme cela a été vu dans le chapitre de dynamique du point, une particule chargée de masse m , de charge q , animée d'une vitesse $\vec{v}$ subit, en présence d'un champ électrique $\vec{E}$ et d'un champ magnétique $\vec{B}$, la **force de Lorentz** dont on rappelle l'expression :

$$\vec{f} = q \left(\vec{E} + \vec{v} \wedge \vec{B} \right).$$

où $\vec{E}$, $\vec{B}$ et $\vec{v}$ dépendent du référentiel. Cette force, qui ne s'applique qu'aux particules chargées, est composée d'un terme électrique $q\vec{E}$ et d'un terme magnétique $q\vec{v} \wedge \vec{B}$.

CHAPITRE 18 – MOUVEMENT D’UNE PARTICULE CHARGÉE DANS UN CHAMP

La force électrique $\vec{F}_{\text{elec}}$ est alignée avec le champ électrique $\vec{E}$. Elle a le même sens si la charge q de la particule est positive et un sens opposé si q est négative.

Pour trouver la direction de la force magnétique $\vec{F}_{\text{mag}}$, on utilise la « règle de la main droite » représentée sur la figure 18.1. Tout comme la force électrique, elle change de sens lorsque le signe de la charge change.

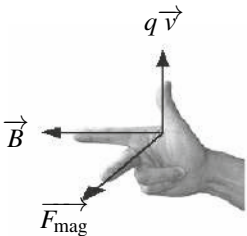


Figure 18.1 –
Règle de la main droite.

1.2 Différence fondamentale entre la composante électrique et la composante magnétique

Du point de vue de la mécanique, une différence fondamentale existe entre les deux composantes de la force de Lorentz :

- la force magnétique est toujours orthogonale à la vitesse. Sa puissance, et donc son travail, sont nuls. Le théorème de l'énergie cinétique implique que l'énergie cinétique d'un corps soumis uniquement à ce type de force est conservée. La norme de sa vitesse est alors une constante du mouvement. Ainsi, la force magnétique ne peut agir que sur la direction du mouvement.
- la force électrique peut délivrer une puissance à une particule chargée. Elle est capable de la mettre en mouvement et de modifier son énergie cinétique. La force électrique peut agir à la fois sur la norme et sur la direction de la vitesse.

La force électrique permet d'accélérer ou de freiner une particule chargée. La force magnétique ne peut que la dévier.

On montre dans le paragraphe 2.3 qu'il est relativement aisé d'accélérer un électron se déplaçant dans le vide jusqu'à une vitesse de $3 \cdot 10^7 \text{ m} \cdot \text{s}^{-1}$. Cette vitesse, qui correspond au dixième de la vitesse de la lumière, est la limite supérieure permettant de négliger les effets relativistes donc de négliger la vitesse de la particule devant celle de la lumière.

1.3 Ordre de grandeur et conséquences

a) Unités et ordre de grandeur des champs électriques et magnétiques

La valeur des champs magnétiques et électriques peut balayer plusieurs ordres de grandeur et il est bon d'avoir en tête que :

- le champ électrique se mesure en volt par mètre ($1 \text{ V} \cdot \text{m}^{-1} = 1 \text{ kg} \cdot \text{m} \cdot \text{s}^{-3} \cdot \text{A}^{-1}$). Un champ électrique de l'ordre de $10^6 \text{ V} \cdot \text{m}^{-1}$ est un champ intense qui peut provoquer la formation d'étincelles dans l'air.
- le champ magnétique se mesure en Tesla ($1 \text{ T} = 1 \text{ kg} \cdot \text{s}^{-2} \cdot \text{A}^{-1}$). Il prend usuellement des valeurs comprises entre $5 \cdot 10^{-5} \text{ T}$ pour le champ magnétique terrestre et 30 T pour le champ magnétique intense d'un spectroscope à résonance magnétique nucléaire (R.M.N.). Les aimants usuels produisent des champs magnétiques de l'ordre de $0,1 \text{ T}$.

b) Comparaison des termes électrique et magnétique

Pour avoir un ordre de grandeur de cette force et de ses composantes électrique et magnétique, on considère le cas d'un électron dont la charge vaut $q = -e = -1,6 \cdot 10^{-19} \text{ C}$, animé d'une vitesse égale à $3 \cdot 10^7 \text{ m} \cdot \text{s}^{-1}$. Pour un champ magnétique peu intense de module $0,1 \text{ T}$, le module de la composante magnétique de la force de Lorentz est de l'ordre de :

$$F_{\text{mag}} = qvB = 1,6 \cdot 10^{-19} \times 3 \cdot 10^7 \times 0,1 = 5 \cdot 10^{-13} \text{ N}.$$

Pour obtenir une composante électrique de la force de Lorentz du même ordre de grandeur, il faut un champ électrique de module :

$$E = \frac{F_{\text{elec}}}{q} \simeq \frac{F_{\text{mag}}}{q} = \frac{5 \cdot 10^{-13}}{1,6 \cdot 10^{-19}} = 3 \cdot 10^6 \text{ V} \cdot \text{m}^{-1}.$$

Il s'agit d'un champ électrique important. Cela signifie que, pour de telles vitesses, la force magnétique d'un champ magnétique peu intense est aussi importante que la force électrique d'un champ électrique intense. C'est la raison principale pour laquelle on utilise très souvent des forces magnétiques lorsque l'on veut dévier des particules chargées allant à grande vitesse.

c) Comparaison au poids de la particule

En prenant le cas d'un proton sur Terre, la masse vaut $m_p = 1,7 \cdot 10^{-27} \text{ kg}$ et le champ de pesanteur $g = 9,8 \text{ m} \cdot \text{s}^{-2}$. L'ordre de grandeur du poids du proton est de $m_p g = 1,7 \cdot 10^{-27} \times 9,8 \simeq 1,7 \cdot 10^{-26} \text{ N}$. La charge du proton valant $e = 1,6 \cdot 10^{-19} \text{ C}$, la force électrique atteint la même valeur que le poids ($eE = m_p g$) pour un champ électrique extrêmement faible de norme $E = 10^{-7} \text{ V} \cdot \text{m}^{-1}$. Le poids est donc largement négligeable devant la force électrique.

Concernant la force magnétique, pour qu'un proton soit soumis à une force magnétique comparable à son poids ($evB = m_p g$), il ne doit pas dépasser la vitesse de $2 \cdot 10^{-3} \text{ m} \cdot \text{s}^{-1}$ dans le champ magnétique terrestre de $5 \cdot 10^{-5} \text{ T}$. C'est une vitesse extrêmement faible. Le poids est donc largement négligeable devant la force magnétique.

Dans les deux cas, pour un électron dont la masse est environ 2000 fois plus faible que celle d'un proton, le rapport de force est encore plus défavorable au poids.

On peut toujours négliger le poids d'une particule chargée soumise à un champ électromagnétique.

2 Mouvement dans un champ électrique uniforme

On s'intéresse dans un premier temps au mouvement d'une particule chargée dans un champ électrique seul.

2.1 Équation du mouvement

On étudie le mouvement de la particule dans le référentiel du laboratoire. On assimile cette particule à un point matériel M de masse m et de charge q et on considère le référentiel du laboratoire comme galiléen. La particule est soumise à :

CHAPITRE 18 – MOUVEMENT D’UNE PARTICULE CHARGÉE DANS UN CHAMP

- la force électrique : $\vec{f} = q \vec{E}$,
- son poids que l’on néglige.

On applique le principe fondamental de la dynamique qui s’écrit :

$$\frac{d(m \vec{v})}{dt} = q \vec{E} \quad \Longleftrightarrow \quad \vec{a} = \frac{q \vec{E}}{m},$$

en notant $\vec{a}$ l’accélération de la particule. Ce type de mouvement, caractérisé par une accélération constante, a été étudié en détail lors du chapitre de cinématique du point.

2.2 Étude de la trajectoire

Comme le champ $\vec{E}$ est indépendant du temps, on peut intégrer vectoriellement soit :

$$\vec{v} = \frac{q \vec{E}}{m} t + \vec{v}_0,$$

et :

$$\vec{OM} = \frac{1}{2} \frac{q}{m} \vec{E} t^2 + \vec{v}_0 t, \tag{18.1}$$

en prenant comme origine O la position initiale de la particule.

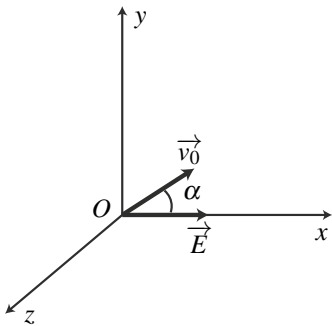


Figure 18.2 –
Conditions initiales du
mouvement.

Le mouvement a lieu dans le plan défini par la position initiale, le champ électrique et le vecteur vitesse initial. Ceci n’a rien de surprenant : la seule force qui s’exerce sur la particule est dans ce plan. On remarque l’analogie avec le mouvement d’un point matériel dans le champ de pesanteur : dans les deux cas, la force appliquée est constante et uniforme.

Pour aller plus loin dans l’analyse du mouvement, on se place en coordonnées cartésiennes. On oriente l’axe (Ox) dans le sens du champ électrique de sorte que : $\vec{E} = E \vec{u}_x$ avec $E > 0$. La vitesse initiale définit avec le champ électrique un plan que l’on prend comme plan (xOy) et on note α l’angle que fait le vecteur vitesse avec le champ électrique dans ce plan (figure 18.2). L’axe (Oz) complète la base directe. La projection de l’équation (18.1) dans cette base, conduit à :

$$\begin{cases} x &= \frac{qE}{2m} t^2 + v_0 t \cos \alpha \\ y &= v_0 t \sin \alpha \\ z &= 0. \end{cases} \tag{18.2}$$

La trajectoire s’obtient en éliminant le temps t dans les relations précédentes ; le plus simple est d’obtenir t en fonction de y . On doit donc distinguer deux cas : le cas où $\alpha = 0$ ou π et le cas où $\alpha \neq 0$, c’est-à-dire le cas où la vitesse initiale est parallèle au champ électrique et celui où elle ne lui est pas parallèle.

a) Trajectoire lorsque la vitesse initiale est parallèle au champ

Dans ce cas, $\sin \alpha = 0$ et $\cos \alpha = 1$ si $\alpha = 0$ et $\cos \alpha = -1$ si $\alpha = \pi$. La loi horaire du mouvement devient :

$$x = \frac{qE}{2m}t^2 \pm v_0t \quad ; \quad y = 0 \quad ; \quad z = 0,$$

ce qui correspond à un mouvement rectiligne le long de l'axe Ox .

b) Trajectoire lorsque la vitesse initiale n'est pas parallèle au champ

Dans ce cas, on a $\alpha \neq 0$ donc $\sin \alpha \neq 0$. On peut alors diviser par $v_0 \sin \alpha$ et obtenir :

$$t = \frac{y}{v_0 \sin \alpha}.$$

En reportant dans l'équation horaire de x , on en déduit :

$$x = \frac{qE}{2mv_0^2 \sin^2 \alpha} y^2 + \frac{\cos \alpha}{\sin \alpha} y.$$

Les trajectoires sont des paraboles passant par l'origine que l'on a représentées sur la figure 18.3. On retrouve le même type de mouvement que pour un point matériel de masse m évoluant sous l'unique action de son poids.

Ces trajectoires sont caractérisées par leur sommet qui correspond à un extremum de x . On cherche donc les valeurs de y qui annulent la dérivée $\frac{dx}{dy}$:

$$\frac{dx}{dy} = \frac{qE}{mv_0^2 \sin^2 \alpha} y + \frac{\cos \alpha}{\sin \alpha} = 0.$$

On obtient $y_s = -\frac{mv_0^2 \sin \alpha \cos \alpha}{qE}$ que l'on reporte dans l'équation de la trajectoire pour déterminer $x_s = -\frac{mv_0^2 \cos^2 \alpha}{2qE}$.

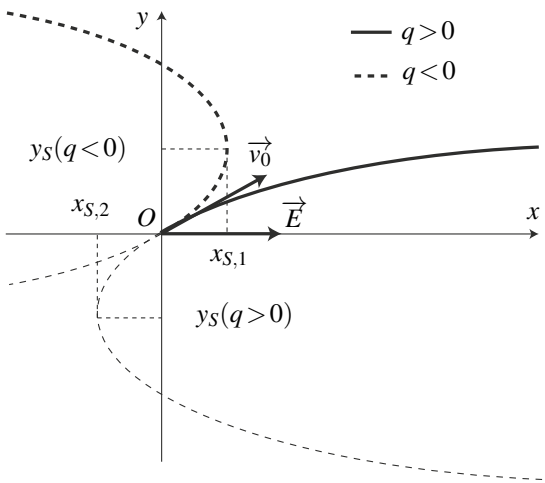


Figure 18.3 – Trajectoires d'une particule chargée positivement ou négativement dans un champ électrique quand la vitesse initiale n'est pas colinéaire au champ.

CHAPITRE 18 – MOUVEMENT D'UNE PARTICULE CHARGÉE DANS UN CHAMP

2.3 Accélération d'une particule chargée par un champ électrique

a) Étude générale

On s'intéresse dans ce paragraphe aux aspects énergétiques afin de montrer qu'un champ électrique peut modifier la norme de la vitesse d'une particule chargée.

Expression de l'énergie potentielle Le champ électrique $\vec{E}$ est supposé uniforme. On a vu dans le chapitre sur l'énergie que la force électrique est alors conservative et dérive d'une énergie potentielle $E_p(x)$ telle que :

$$\delta W(\vec{f}) = \vec{f} \cdot d\vec{OM} = q\vec{E} \cdot d\vec{OM} = qE\vec{u}_x \cdot d\vec{OM} = qEdx = -dE_p.$$

On trouve alors :

$E_p = q(-Ex + C)$

où C est une constante arbitraire.

Définition du potentiel électrique

On définit le **potentiel électrique** noté V par la relation :

$$E_p = qV \quad \text{soit} \quad V = \frac{E_p}{q}.$$

Lorsque le champ électrique est uniforme, dirigé selon (Ox) et d'expression $\vec{E} = E\vec{u}_x$, le potentiel électrique V est une fonction linéaire de x de dérivée $-E$:

$$V(x) = -Ex + C.$$

Le potentiel électrique V qui est défini ici ne dépend pas de la charge q mais uniquement du champ électrique E . Il s'agit du même potentiel que celui que l'on a introduit en électricité. Le choix de la constante arbitraire C correspond au choix de la position du potentiel électrique nul (référence de potentiel) et donc au choix de la masse en électricité.

Le champ électrique est toujours dirigé vers les potentiels décroissants. En effet :

- si $E > 0$, $V(x)$ est une fonction décroissante de x et le champ $\vec{E}$ est dirigé dans le sens des x croissants ;
- si $E < 0$, $V(x)$ est une fonction croissante de x et le champ $\vec{E}$ est dirigé dans le sens des x décroissants.

Remarque

La force électrique est conservative même si le champ électrique n'est pas uniforme. On peut alors définir un potentiel électrique mais son expression n'est plus la même.

Production d'un champ électrique uniforme Pour obtenir un champ électrique uniforme parallèle à l'axe (Ox) , il faut créer un potentiel qui varie linéairement selon l'axe (Ox) . Pour cela, il suffit d'appliquer une différence de potentiel (ou tension) U entre deux grilles métalliques planes perpendiculaires à l'axe (Ox) (et donc parallèles entre elles) et distantes de d (figure 18.4). On admet que le champ créé entre les deux électrodes est alors uniforme.

Pour établir la tension U , on branche un générateur entre les deux électrodes. On peut fixer arbitrairement la masse sur l'électrode placée en $x = \frac{d}{2}$ et le potentiel U en $x = -\frac{d}{2}$. Le potentiel évolue alors avec l'abscisse x selon la relation linéaire $V(x) = -Ex + C$ et vérifie les conditions aux limites suivantes en $x = \pm \frac{d}{2}$:

$$\begin{cases} V\left(-\frac{d}{2}\right) = U = E\frac{d}{2} + C \\ V\left(\frac{d}{2}\right) = 0 = -E\frac{d}{2} + C. \end{cases}$$

On élimine C en calculant la différence entre ces deux équations et on obtient $E = \frac{U}{d}$. Lorsque le potentiel appliqué en $x = -\frac{d}{2}$ est positif, le champ électrique est dirigé vers les x croissants.

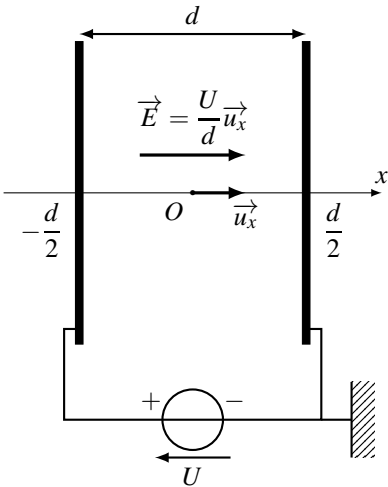


Figure 18.4 – Production d'un champ électrique uniforme.

En appliquant une différence de potentiel U entre deux grilles (ou électrodes) planes, parallèles et distantes de d , on obtient un champ électrique perpendiculaire aux grilles, dirigé vers les potentiels décroissants et de norme :

$$E = \frac{U}{d}.$$

Il est aisé de fabriquer une différence de potentiel de 1 kV et de l'appliquer entre deux plans distants de 1 cm. Cela conduit à un champ électrique de $1 \text{ kV} \cdot \text{cm}^{-1} = 10^5 \text{ V} \cdot \text{m}^{-1}$.

b) Conservation de l'énergie mécanique

La seule force étant la force électrique qui est conservative, le théorème de l'énergie mécanique implique la conservation de l'énergie mécanique d'où :

$$E_m = E_c + E_p = \frac{1}{2}mv^2 + qV = \text{constante}.$$

En indiquant par l'indice i la situation initiale et par f la situation finale, on en déduit la variation d'énergie cinétique :

$$\Delta E_c = \frac{1}{2}mv_f^2 - \frac{1}{2}mv_i^2 = -q(V(x_f) - V(x_i)). \tag{18.3}$$

CHAPITRE 18 – MOUVEMENT D'UNE PARTICULE CHARGÉE DANS UN CHAMP

Pour faire varier l'énergie cinétique et donc la norme de la vitesse d'une particule chargée soumise uniquement à une force électromagnétique, il faut lui faire franchir une différence de potentiel ΔV . La variation de l'énergie cinétique est dans ce cas :

$$\Delta E_c = -q\Delta V.$$

Si $q > 0$, la particule sera accélérée par une différence de potentielle $\Delta V < 0$ et freinée par une différence de potentielle $\Delta V > 0$.

Si $q < 0$, la particule sera accélérée par une différence de potentielle $\Delta V > 0$ et freinée par une différence de potentielle $\Delta V < 0$.

c) Cas particulier où la vitesse initiale est nulle

Dans ce cas, le mouvement est nécessairement rectiligne et accéléré. Le sens de la force $\vec{f} = q\vec{E}$ dépend du signe de la charge considérée.

Accélération de particules chargées positivement On injecte une particule de charge positive sans vitesse initiale à proximité de l'électrode positive : $x_i = -\frac{d}{2}$. Elle est accélérée vers les x croissants jusqu'à l'électrode négative : $x_f = \frac{d}{2}$. L'énergie cinétique finale s'exprime en fonction de la différence de potentiel $V(x_f) - V(x_i) = V\left(\frac{d}{2}\right) - V\left(-\frac{d}{2}\right) = -U$ à l'aide de l'équation (18.3) :

$$\frac{1}{2}mv_f^2 = qU > 0 \quad \implies \quad v_f = \sqrt{\frac{2qU}{m}}.$$

où la différence de potentiel U est appelée tension accélératrice.

Exemple

Pour un proton : $q = 1,6 \cdot 10^{-19}$ C et $m_p = 1,67 \cdot 10^{-27}$ kg. Si la tension accélératrice vaut $U = 2,0$ kV, la vitesse finale est :

$$v_f = \sqrt{\frac{2 \times 1,6 \cdot 10^{-19} \times 2,0 \cdot 10^3}{1,67 \cdot 10^{-27}}} = 6,2 \cdot 10^5 \text{ m} \cdot \text{s}^{-1}.$$

Accélération des particules chargées négativement On injecte une particule de charge négative sans vitesse initiale à proximité de l'électrode négative : $x_i = \frac{d}{2}$. Elle est accélérée vers les x décroissants jusqu'à atteindre l'électrode positive située en $x_f = -\frac{d}{2}$. L'énergie cinétique finale s'exprime en fonction de la différence de potentiel $V(x_f) - V(x_i) = V\left(-\frac{d}{2}\right) - V\left(\frac{d}{2}\right) = U$ à l'aide de l'équation (18.3) :

$$\frac{1}{2}mv_f^2 = -qU > 0 \quad \implies \quad v_f = \sqrt{\frac{-2qU}{m}}.$$

Exemple

Pour un électron : $q = -1,6 \cdot 10^{-19} \text{ C}$ et $m_e = 9,1 \cdot 10^{-31} \text{ kg}$. Si la tension accélératrice vaut $U = 2,0 \text{ kV}$, la vitesse finale est :

$$v_f = \sqrt{\frac{2 \times 1,6 \cdot 10^{-19} \times 2,0 \cdot 10^3}{9,1 \cdot 10^{-31}}} = 2,7 \cdot 10^7 \text{ m} \cdot \text{s}^{-1}.$$

Une nouvelle unité d'énergie adaptée : l'électron-volt La formule (18.3) montre que le produit d'une charge par une différence de potentiel est homogène à une énergie. On définit l'**électron-volt** comme le produit de la charge de l'électron e par le volt. Cette nouvelle unité d'énergie vaut $1 \text{ eV} = 1,6 \cdot 10^{-19} \text{ J}$. Les exemples précédents permettent de comprendre sa signification et son intérêt.

Un électron ou un proton initialement au repos et accéléré sous une tension de 1 V acquiert une énergie cinétique de 1 eV .

Remarque

L'énergie cinétique finale d'un électron initialement au repos et accéléré par une tension accélératrice de 2 kV vaut 2 keV . Sa vitesse est alors de $2,7 \cdot 10^7 \text{ m} \cdot \text{s}^{-1}$ soit environ 10% de la vitesse de la lumière. Une telle différence de potentiel est relativement aisée à réaliser, on peut donc facilement obtenir des électrons relativistes. Dans ce cours, on se limite à des tensions accélératrices inférieures ou égales à 2 kV pour les applications numériques concernant des électrons afin de rester dans l'approximation classique.

3 Mouvement dans un champ magnétique

On s'intéresse au mouvement d'une particule chargée dans un champ magnétique uniforme. Lorsque la vitesse de la particule est initialement perpendiculaire au champ magnétique, on observe expérimentalement qu'elle suit une trajectoire circulaire. On veut expliquer cette trajectoire et déterminer son rayon.

3.1 Le mouvement est uniforme

On étudie le mouvement d'une particule de charge q et de masse m dans un champ magnétique uniforme $\vec{B}$ de norme B . À l'instant initial, la particule est animée d'une vitesse initiale $\vec{v}_0$ de norme v_0 et de direction perpendiculaire à $\vec{B}$.

On modélise la particule par un point matériel M et on étudie son mouvement dans le référentiel du laboratoire considéré comme galiléen. La particule est soumise aux forces suivantes :

- la force magnétique $\vec{f} = q \vec{v} \wedge \vec{B}$,
- son poids qui sera négligé compte tenu des ordres de grandeur.

Dans un premier temps, on s'intéresse uniquement à la norme de la vitesse et on utilise le théorème de l'énergie cinétique. La force magnétique étant perpendiculaire à la vitesse à chaque instant, sa puissance, et par conséquent son travail, sont nuls. L'application du théo-

CHAPITRE 18 – MOUVEMENT D’UNE PARTICULE CHARGÉE DANS UN CHAMP

rème de l’énergie cinétique :

$$\Delta E_c = W(\vec{f}) = 0,$$

montre que l’énergie cinétique $E_c = \frac{1}{2}mv^2$ de la particule est une constante du mouvement. Par conséquent, le module de sa vitesse conserve sa valeur initiale v_0 .

Le mouvement d’une particule chargée dans un champ magnétique est uniforme.

3.2 Étude de la trajectoire

a) Observations expérimentales

La trajectoire des particules est circulaire dans un plan perpendiculaire à $\vec{B}$. Le sens de parcours et la position du centre de la trajectoire dépendent du signe de la charge (figure 18.5).

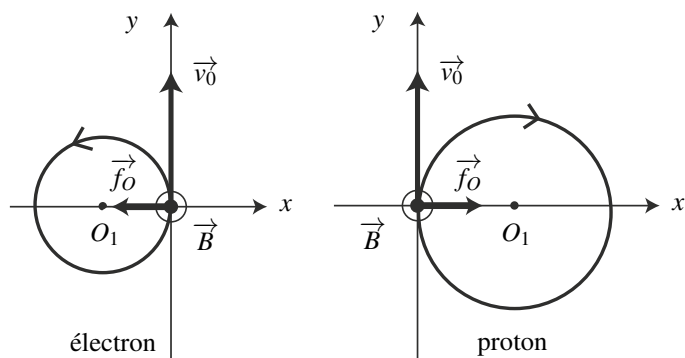


Figure 18.5 – Particules chargées évoluant en présence d’un champ magnétique uniforme. Les rayons ne sont pas tracés avec la même échelle dans le cas de l’électron et du proton.

b) Détermination du rayon et de la vitesse angulaire

On applique le principe fondamental de la dynamique :

$$m \frac{d\vec{v}}{dt} = q \vec{v} \wedge \vec{B}. \tag{18.4}$$

On peut étudier cette trajectoire en coordonnées polaires d’origine O_1 confondue avec le centre de la trajectoire et d’axe (O_1z) orienté par $\vec{B}$. On a montré précédemment que le mouvement est uniforme à la vitesse v_0 , on utilise donc les résultats sur le mouvement circulaire et uniforme établis dans le chapitre de *Cinématique du point* soit :

$$\overrightarrow{O_1M} = R\vec{u}_r \quad ; \quad \vec{v} = R\dot{\theta}\vec{u}_\theta = \pm v_0\vec{u}_\theta \quad ; \quad \vec{a} = -\frac{v_0^2}{R}\vec{u}_r.$$

À ce niveau, le sens de la rotation n'est pas déterminé, ce qui explique la notation $\vec{v} = \pm v_0 \vec{u}_\theta$. L'accélération est dirigée selon $-\vec{u}_r$ (elle est toujours dirigée vers l'intérieur de la trajectoire). La force magnétique étant la seule force, elle est nécessairement parallèle et de même sens que l'accélération d'après le principe fondamental de la dynamique. Cela permet de raisonner sur les normes. Les vecteurs $q\vec{v}$ et $\vec{B}$ sont perpendiculaires entre eux et de norme respectives $|q|v_0$ et B , donc $\|q\vec{v} \wedge \vec{B}\| = |q|v_0B$. Par ailleurs, la norme de l'accélération vaut $\frac{v_0^2}{R}$. Le principe fondamental de la dynamique (18.4) conduit ensuite à :

$$m \frac{v_0^2}{R} = |q|v_0B \quad \Longleftrightarrow \quad v_0 = \frac{|q|B}{m}R.$$

La dimension de $\frac{|q|B}{m}$ est celle d'une pulsation car :

$$[Q \times v \times B] = [F] \Rightarrow \left[\frac{Q \times B}{M} \right] = \left[\frac{F}{M \times v} \right] = \left[\frac{F}{M \times a \times T} \right] = \left[\frac{F}{F \times T} \right] = [T^{-1}].$$

On note alors $\omega_c = \frac{|q|B}{m}$ cette pulsation qui se nomme **pulsation cyclotron**. La relation entre le rayon et la vitesse v_0 devient alors :

$$v_0 = R\omega_c \quad \Longleftrightarrow \quad R = \frac{v_0}{\omega_c}.$$

Une particule de charge q et de vitesse initiale $\vec{v}_0$ perpendiculaire à un champ magnétique uniforme $\vec{B}$ suit un **mouvement circulaire et uniforme** à la vitesse angulaire $\omega_c = \frac{|q|B}{m}$ appelée **pulsation cyclotron**. On trouve le rayon de la trajectoire grâce à la relation $v_0 = R\omega_c$.

Remarque

Lorsque la vitesse initiale n'est pas perpendiculaire au champ magnétique, le mouvement est hélicoïdal, composé d'un mouvement rectiligne et uniforme dans la direction de $\vec{B}$ et d'un mouvement circulaire et uniforme dans le plan perpendiculaire à $\vec{B}$.

c) Position du centre de la trajectoire - Sens de parcours de la trajectoire

On peut tracer la trajectoire dans le cas d'une injection de particule en O avec la vitesse $\vec{v}_0$. On oriente pour cela l'axe (Oy) selon $\vec{v}_0$. La droite (Oy) est alors une tangente à la trajectoire. L'axe (Ox) , perpendiculaire à (Oy) , est donc un diamètre de la trajectoire. La force magnétique est nécessairement dirigée vers le centre de la trajectoire circulaire. Sa direction au passage en O avec une vitesse $\vec{v}_0$ permet donc de situer le centre O_1 de la trajectoire (voir figure 18.5) :

$$\vec{f}_O = q\vec{v}_0 \wedge \vec{B} = qv_0\vec{u}_y \wedge B\vec{u}_z = qv_0B\vec{u}_x.$$

On en déduit que :

- pour une charge positive, $\vec{f}_O$ est dirigée selon $+\vec{u}_x$ donc O_1 a une abscisse positive,
- pour une charge négative, $\vec{f}_O$ est dirigée selon $-\vec{u}_x$ donc O_1 a une abscisse négative.

CHAPITRE 18 – MOUVEMENT D'UNE PARTICULE CHARGÉE DANS UN CHAMP

La connaissance du rayon de la trajectoire permet alors de situer son centre et de la tracer. La vitesse en O permet ensuite de déduire le sens de parcours de la trajectoire (figure 18.5). Il faut noter que le sens du mouvement circulaire dépend du signe de la charge q : les électrons tournent dans le sens direct autour du champ $\vec{B}$ et les protons dans le sens indirect.

d) Ordres de grandeur

On considère des particules préalablement accélérées par une différence de potentiel de 2000 V comme cela a été détaillé plus haut. Un proton possède la vitesse $v_0 = 6,2 \cdot 10^5 \text{ m}\cdot\text{s}^{-1}$ et un électron la vitesse $v_0 = 2,7 \cdot 10^7 \text{ m}\cdot\text{s}^{-1}$. Lorsqu'on les soumet à un champ magnétique de 0,10 T, ces particules de charge $|q| = e = 1,6 \cdot 10^{-19} \text{ C}$ suivent des trajectoires circulaires de rayon :

$$\begin{aligned} R &= \frac{mv_0}{eB} = \frac{1,67 \cdot 10^{-27} \times 6,2 \cdot 10^5}{1,6 \cdot 10^{-19} \times 0,10} = 6,5 \text{ cm pour le proton de masse } m_p = 1,67 \cdot 10^{-27} \text{ kg} \\ &= \frac{9,1 \cdot 10^{-31} \times 2,7 \cdot 10^7}{1,6 \cdot 10^{-19} \times 0,10} = 1,5 \text{ mm pour l'électron de masse } m_e = 9,1 \cdot 10^{-31} \text{ kg.} \end{aligned}$$

La vitesse angulaire s'écrit :

$$\begin{aligned} \omega_c &= \frac{eB}{m} = \frac{1,6 \cdot 10^{-19} \times 0,10}{1,67 \cdot 10^{-27}} = 9,5 \cdot 10^6 \text{ rad}\cdot\text{s}^{-1} \text{ pour le proton} \\ &= \frac{1,6 \cdot 10^{-19} \times 0,10}{9,1 \cdot 10^{-31}} = 1,8 \cdot 10^{10} \text{ rad}\cdot\text{s}^{-1} \text{ pour l'électron.} \end{aligned}$$

L'électron tourne plus rapidement que le proton et sur une trajectoire de rayon plus petit.

4 Quelques applications de ces mouvements

On considère quelques applications des mouvements de particules dans des champs électrique et magnétique. La liste est loin d'être exhaustive et on se limite à l'analyse de quelques expériences intéressantes.

4.1 Expérience de Thomson

À la fin du XIX^e siècle, J.J. Thomson conduit une série d'expériences qui le conduisent à mettre en évidence l'existence de l'électron et à en déterminer le rapport $\frac{e}{m_e}$ de sa charge par sa masse. C'est une avancée décisive dans la connaissance de la matière qui lui vaut d'être récompensé d'un prix Nobel de Physique en 1906. Toutes ces expériences mettent en jeu des mouvements d'électrons dans des tubes à vide.

L'expérience permettant de mesurer ce rapport consiste tout d'abord à accélérer un faisceau d'électrons (charge $-e$) initialement immobiles par un champ électrique $\vec{E}_0$, ce qui leur procure une vitesse $\vec{v}_0$. Ce faisceau est ensuite dévié par un champ électrique $\vec{E}$ appliqué perpendiculairement au mouvement des électrons du faisceau sur une longueur a . On mesure alors la déviation d induite par ce champ électrique sur un écran situé à une distance $L \gg a$. (voir figure 18.6).

L'application du seul champ électrique conduit à une déviation d'un angle θ dont le calcul est présenté au paragraphe suivant et qui vérifie :

$$\tan \theta = \frac{eEa}{m_e v_0^2} \simeq \frac{d}{L}.$$

On en déduit l'expression de la vitesse :

$$v_0 = \sqrt{\frac{eEaL}{m_e d}}.$$

On superpose alors un champ magnétique $\vec{B}$ perpendiculaire à $\vec{E}$ et tel que la force magnétique $-e\vec{v}_0 \wedge \vec{B}$ soit égale à la force électrique $-e\vec{E}$. Le faisceau d'électrons n'est alors plus dévié. L'égalité des forces électrique et magnétique donne : $eE = ev_0B$ soit $v_0 = \frac{E}{B}$.

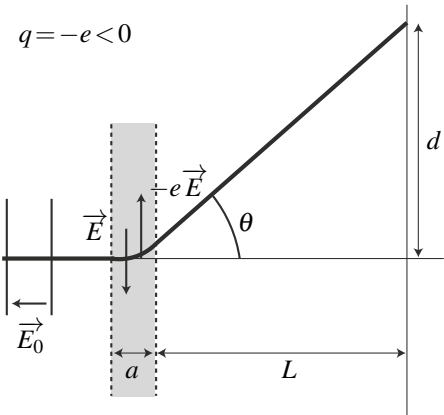


Figure 18.6 – Expérience de Thomson avec le champ électrique seul.

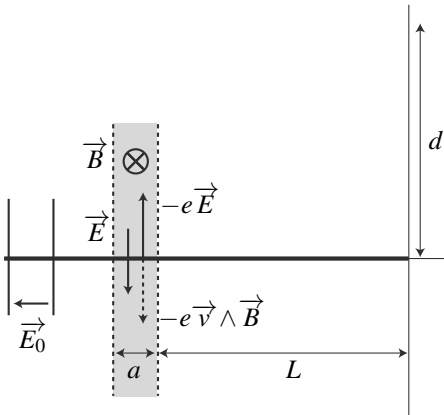


Figure 18.7 – Expérience de Thomson avec champs électrique et magnétique.

En égalant les deux expressions de la vitesse, on obtient :

$$v_0^2 = \frac{E^2}{B^2} = \frac{eEaL}{m_e d} \implies \frac{e}{m_e} = \frac{dE}{aLB^2}.$$

Les paramètres géométriques a et L sont imposés par la configuration géométrique de l'expérience. La mesure de d et des champs E et B permet d'accéder au rapport $\frac{e}{m_e}$.

Dans cette expérience, on utilise un champ électrique pour dévier une particule puis un champ magnétique pour annuler l'effet de cette force. Pour un électron de vitesse $v_0 = 2,7 \cdot 10^7 \text{ m}\cdot\text{s}^{-1}$ et un champ électrique important $E = 10^6 \text{ V}\cdot\text{m}^{-1}$, il suffit d'un champ magnétique faible $B = \frac{E}{v_0} = 0,04 \text{ T}$. Pour dévier un faisceau d'électron, un champ magnétique peu intense est aussi efficace qu'un champ électrique fort.

CHAPITRE 18 – MOUVEMENT D'UNE PARTICULE CHARGÉE DANS UN CHAMP

Lorsque l'on veut dévier des particules chargées, il est souvent préférable d'utiliser des champs magnétiques. C'est d'autant plus vrai que la vitesse des particules est élevée.

Annexe : déviation d'un électron par un champ électrique Pour calculer la déviation, on remarque que la vitesse initiale est perpendiculaire au champ électrique et que le champ ne s'applique que dans une zone limitée de l'espace de largeur a . On se ramène à la situation étudiée au paragraphe 2.2 en choisissant l'axe (Ox) dirigé par le champ $\vec{E}$ et l'axe (Oy) dirigé par la vitesse $\vec{v}_0$ (figure 18.8). Dans la zone comprise entre $y = 0$ et $y = a$, le champ électrique est uniforme et le mouvement des électrons est régi par les équations (18.2) avec $\alpha = 0$ et $q = -e$ soit : $x = \frac{-eE}{2m_e}t^2$ et $y = v_0t$.

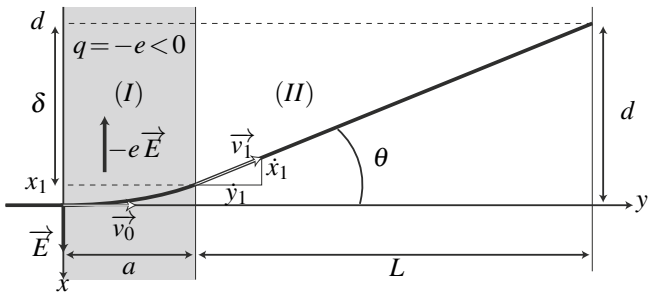


Figure 18.8 – Déviation d'une particule chargée par un champ électrique. Le champ appliqué dans la zone (I) dévie la trajectoire. Le mouvement est ensuite rectiligne et uniforme dans la zone (II).

La particule quitte la zone où règne le champ en $y_1 = a$ à l'instant $t_1 = \frac{a}{v_0}$ en un point d'abscisse $x_1 = \frac{-eEa^2}{2m_ev_0^2}$. Elle suit alors une trajectoire rectiligne et uniforme de vitesse $\vec{v}_1$.

L'angle de déviation de la trajectoire est donné par : $\tan \theta = \tan(\widehat{\vec{v}_0, \vec{v}_1}) = \left| \frac{\dot{x}_1}{\dot{y}_1} \right|$.

Or $\dot{x} = \frac{-eEt}{m_e}$ et $\dot{y} = v_0$. Donc $\dot{x}_1 = \frac{-eEa}{m_ev_0}$ et $\tan \theta = \frac{eEa}{m_ev_0^2}$.

On peut ensuite calculer la déviation $d = x_1 + \delta$ observée sur un écran situé à une distance $L \gg a$ des plaques de déviation car la tangente de l'angle de déviation peut aussi s'exprimer en fonction de L et de δ :

$$\tan \theta = \frac{\delta}{L} \quad \text{d'où} \quad d = \frac{eEa}{m_ev_0^2} \left(\frac{a}{2} + L \right) \simeq \frac{eEaL}{m_ev_0^2} \quad \text{lorsque } a \ll L.$$

4.2 Spectromètre de masse

Le but d'un spectromètre de masse est d'analyser les ions présents dans un faisceau. On peut utiliser leurs différences de masse pour connaître leur proportion respective en supposant que

leur charge est connue. Le montage utilisé comprend deux phases :

- **Phase accélératrice :** On commence par accélérer les ions grâce à une forte différence de potentiel. L'application du théorème de l'énergie cinétique aux ions supposés initialement au repos fournit leur vitesse à la fin de cette phase :

$$v = \sqrt{\frac{2|q|U}{m}}.$$

On se place cependant dans la limite classique en vérifiant que $v \ll c$ pour éviter un traitement relativiste.

- **Phase de déviation :** On applique alors un champ magnétique $\vec{B}$ perpendiculairement au mouvement. En appliquant les relations établies précédemment, on peut affirmer que les ions décrivent une trajectoire circulaire de rayon $R = \frac{mv}{|q|B}$. On en déduit :

$$v = \frac{qBR}{m} = \sqrt{\frac{2|q|U}{m}} \quad \Rightarrow \quad \frac{q^2 B^2 R^2}{m^2} = \frac{2|q|U}{m}$$

puis finalement :

$$\frac{|q|}{m} = \frac{2U}{B^2 R^2}.$$

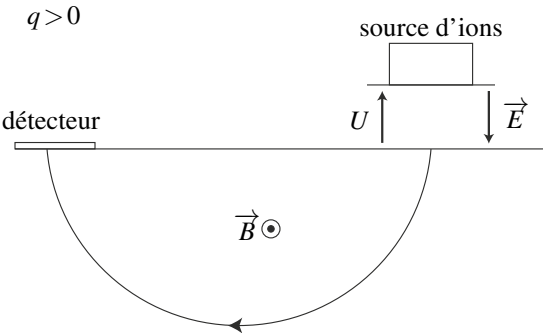


Figure 18.9 – Spectromètre de masse.

L'impact sur le détecteur permet de connaître le rayon de la trajectoire et donc le rapport $\frac{q}{m}$ connaissant la tension accélératrice U et la norme du champ magnétique B . Ce système permet de séparer des éléments ionisés en fonction de leur masse et de leur charge. Si on connaît par ailleurs la charge des ions, on peut en déduire leur masse. Ceci explique le nom de spectromètre de masse donné à ce dispositif.

Ordre de grandeur de la taille du dispositif La charge des ions sera de quelques charges élémentaires donc de l'ordre de 10^{-19} C. Leur masse molaire varie entre quelques centaines et quelques milliers de grammes par mole soit de l'ordre de 10^{-24} kg par ion. On prend une

CHAPITRE 18 – MOUVEMENT D’UNE PARTICULE CHARGÉE DANS UN CHAMP

tension accélératrice de 10 kV et un champ magnétique de 0,1 T. On obtient alors un rayon de l’ordre de :

$$R = \frac{1}{B} \sqrt{\frac{2mU}{|q|}} = \frac{1}{0,1} \sqrt{\frac{2 \cdot 10^{-24} \times 10^4}{10^{-19}}} = 4 \text{ m}.$$

4.3 Cyclotron

L’existence d’une trajectoire circulaire des particules chargées soumises à un champ magnétique permet d’envisager des accélérateurs de particules agissant de manière cyclique. C’est le principe des cyclotrons dont le dispositif est le suivant :

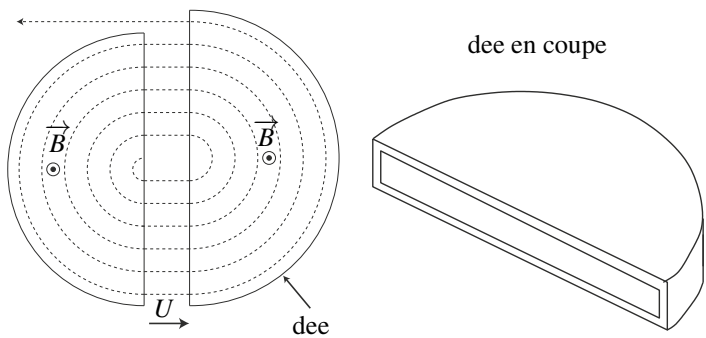


Figure 18.10 – Constitution d’un cyclotron.

Ils sont constitués de deux demi-cylindres creux dénommés « *dees* » du fait de leur forme de « D » majuscules. Chaque dee est soumis à un champ magnétique uniforme, le sens du champ magnétique est le même dans les deux dees. Les deux dees sont séparés par un espace dans lequel on applique une différence de potentiel U . Souvent les particules injectées dans le dispositif ont été préalablement accélérées.

Dans chaque dee, la particule suit un mouvement circulaire lié au champ magnétique. Sa vitesse angulaire vaut :

$$\omega = \omega_c = \frac{qB}{m}$$

C’est la pulsation cyclotron dont on peut déduire la fréquence cyclotron :

$$f_c = \frac{\omega_c}{2\pi} = \frac{qB}{2\pi m}.$$

Quand la particule traverse l’interstice entre les deux dees, elle est accélérée par la présence de la différence de potentiel U . Pour qu’elle soit effectivement accélérée à chaque demi-tour, il faut que les tensions aux instants correspondant à deux passages successifs dans l’interstice soient en opposition de phase. Sinon elle pourra être alternativement accélérée et décélérée. Pour obtenir une telle tension, il suffit de prendre une tension sinusoïdale dont la période est celle du cyclotron. De cette manière, le rayon de la trajectoire augmente à chaque traversée

de l'interstice puisqu'il est proportionnel à la vitesse. On obtient alors le type de trajectoire indiquée en pointillés sur la figure précédente.

L'intérêt d'un tel dispositif est de pouvoir accélérer très fortement les particules sans être contraint d'appliquer une différence de potentiel énorme.

Exemple

Pour le cyclotron de University of Michigan, Ann Arbor (près de Detroit), le champ magnétique vaut $B = 0,10 \text{ T}$ (on peut atteindre au maximum $1,5 \text{ T}$ mais on est alors dans le cadre relativiste), le diamètre des dees est $2,1 \text{ m}$.

Le problème essentiel des cyclotrons est d'obtenir un champ magnétique uniforme sur une surface aussi grande. Il s'agit d'un facteur limitant beaucoup sa taille et donc son utilisation. On a alors une fréquence cyclotron pour les protons de :

$$f_c = \frac{qB}{2\pi m} = \frac{1,6 \cdot 10^{-19} \times 0,10}{2\pi \times 1,67 \cdot 10^{-27}} = 1,5 \text{ MHz}$$

La vitesse maximale est :

$$v_{\max} = R_{\max} \omega = 2\pi \times 1,05 \times 1,5 \cdot 10^6 = 1,0 \cdot 10^7 \text{ m} \cdot \text{s}^{-1}.$$

Pour atteindre une telle vitesse à l'aide d'un seul champ électrique, il faudrait une différence de potentiel de :

$$U = \frac{mv^2}{2|q|} = \frac{1,67 \cdot 10^{-27} \times (10^7)^2}{2 \times 1,6 \cdot 10^{-19}} = 5,3 \cdot 10^5 \text{ V}.$$

Le cyclotron n'est pas le seul type d'accélérateurs de particules. Dans tous les cas, l'accélération est due au champ électrique et le champ magnétique permet de faire revenir les particules au même point et donc de « refermer » les trajectoires.

5 Approche documentaire : limites relativistes en microscopie électronique

Un microscope électronique crée une image très agrandie d'un échantillon. Les grossissements peuvent aller jusqu'à 2 millions, soit environ 1000 fois plus que ce que l'on obtient avec un microscope optique. Néanmoins, la résolution spatiale intrinsèque est toujours plus grande que la longueur d'onde du rayonnement utilisé. Un microscope électronique crée un faisceau d'électrons qu'il dirige vers un échantillon pour l'« éclairer ». La longueur d'onde qui limite la résolution est donc la longueur d'onde de DE BROGLIE du faisceau électronique. Elle est beaucoup plus petite que celle d'un faisceau de lumière visible et d'autant plus petite que les électrons sont rapides. Les électrons sont couramment accélérés jusqu'à des énergies cinétiques E_c comprises entre 10 keV et 1000 keV qui sont telles que l'on doit tenir compte d'effets relativistes.

On doit alors introduire le facteur de Lorentz $\gamma = \frac{1}{\sqrt{1-\beta^2}}$ où $\beta = \frac{v}{c}$ est le rapport de la vitesse de l'électron v à la vitesse de la lumière c . Pour le calculer il faut savoir que, dans le

CHAPITRE 18 – MOUVEMENT D’UNE PARTICULE CHARGÉE DANS UN CHAMP

cadre de la théorie relativiste, l’expression de l’énergie cinétique est modifiée et devient :

$$E_c = (\gamma - 1)m_e c^2,$$

où $m_e c^2$ est l’énergie de masse de l’électron égale à 511 keV. De même, la longueur d’onde de DE BROGLIE est toujours définie à partir de la quantité de mouvement p :

$$\lambda = \frac{h}{p},$$

mais la quantité de mouvement vaut $p = \gamma m_e v$.

On a réunit dans le tableau 18.1 les valeurs des paramètres suivants : la tension accélératrice U , le facteur de Lorentz γ , β et λ . Pour bien se rendre compte de l’importance de tenir compte de la correction relativiste, on a également donné le rapport β_c obtenu dans l’approximation classique, ainsi que la longueur d’onde de DE BROGLIE associée λ_c .

$U(\text{kV})$	γ	β	$\lambda(\text{pm})$	β_c	$\lambda_c(\text{pm})$
10	1,02	0,195	12,19	0,198	12,25
20	1,04	0,272	8,58	0,280	8,66
50	1,10	0,413	5,35	0,442	5,48
100	1,20	0,548	3,70	0,626	3,87
200	1,39	0,695	2,51	0,885	2,74
500	1,98	0,863	1,42	1,40	1,73
1000	2,96	0,941	0,87	1,98	0,122

Tableau 18.1 – Caractéristiques des électrons mis en jeu en microscopie électronique.

Questions À la lecture de ce texte, répondre aux questions suivantes :

1. En quoi les effets relativistes limitent la résolution intrinsèque de l’appareil ?
2. Exprimer les coefficients relativistes γ et β en fonction de la tension accélératrice U , de la charge de l’électron e , de sa masse m_e et de la vitesse de la lumière.
3. Retrouver les longueurs d’onde apparaissant dans le tableau ci-dessus.

Éléments de réponse

1. On voit dans le tableau que la longueur d’onde de DE BROGLIE obtenue en tenant compte des effets relativistes est plus grande que dans le modèle classique. Comme la résolution intrinsèque est limité par cette longueur d’onde, les effets relativistes limitent la résolution intrinsèque de l’appareil.
2. On obtient $\gamma = 1 + \frac{eU}{m_e c^2}$ puis $\beta = \sqrt{1 - \frac{1}{\gamma^2}}$. On en déduit v puis on reporte v et γ dans l’expression de λ pour obtenir :
$$\lambda = \frac{h}{\sqrt{2eUm_e c^2 \left(1 + \frac{eU}{2m_e c^2}\right)}}$$
3. Il n’y plus alors qu’à faire les applications numériques.

SYNTHÈSE

SAVOIRS

- expression de la force de Lorentz
- effet de la partie électrique de cette force
- effet de la partie magnétique de cette force
- lien entre une différence de potentiel et un champ électrique uniforme
- expression de la pulsation cyclotron

SAVOIR-FAIRE

- évaluer les ordres de grandeurs des forces électriques et magnétiques et les comparer aux forces gravitationnelles
- étudier le mouvement d’une particule chargée dans un champ électrique uniforme
- montrer qu’il s’agit d’un mouvement à vecteur accélération constant
- déterminer, à l’aide d’un bilan énergétique, la vitesse d’une particule chargée accélérée par une différence de potentiel
- déterminer le rayon de la trajectoire d’une particule chargée soumise à un champ magnétique uniforme dans le cas où cette trajectoire est circulaire

MOTS-CLÉS

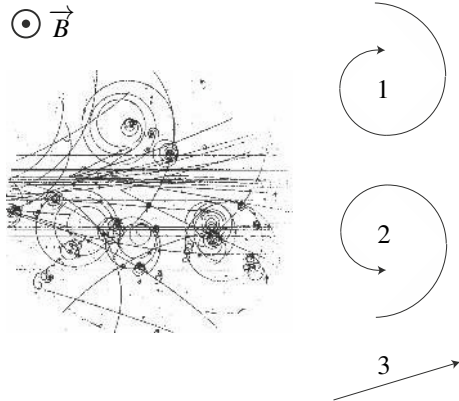
- | | | |
|-----------------------|---------------------------|--------------------|
| • particules chargées | • champ électrique | • champ magnétique |
| • force de Lorentz | • différence de potentiel | |

CHAPITRE 18 – MOUVEMENT D'UNE PARTICULE CHARGÉE DANS UN CHAMP

S'ENTRAÎNER

18.1 Chambre à bulle (★)

Pour visualiser les trajectoires des particules chargées, les premiers détecteurs étaient des « chambre à bulles » dans lesquelles les particules (électrons, protons, neutrons, etc...) déclenchaient la formation de bulles dans un liquide et marquaient ainsi leur passage par une traînée de bulles. La figure ci-contre représente un cliché typique des traces observées lors d'une collision à haute énergie de particules au CERN. Sur le côté droit, on a schématisé les trois types de trajectoires observées avec leur sens de parcourt.



Dans ces chambres à bulles, il règne un champ magnétique uniforme $\vec{B}$. Par ailleurs, le passage dans le liquide conduit à une lente décélération des particules.

1. Déterminer le signe de la charge pour les trois types de trajectoires observées.
2. Expliquer qualitativement pourquoi les trajectoires observées ne sont circulaires mais s'enroulent en spirales dont le rayon diminue.

18.2 Grandeurs cinétiques d'une particule chargée dans un champ magnétique (★)

Une particule de charge q , de masse m et de vitesse $\vec{v}_0$ pénètre dans une région où règne un champ magnétique $\vec{B}$ constant et uniforme avec $\vec{v}_0 \perp \vec{B}$. Comparer les valeurs des quantités suivantes à l'entrée et à la sortie de la zone :

1. énergie cinétique,
2. quantité de mouvement,
3. moment cinétique par rapport au centre de la trajectoire.

18.3 Action d'un champ magnétique sur un proton ou sur un électron (★)

Un électron et un proton de même énergie cinétique décrivent des trajectoires circulaires dans un champ magnétique uniforme. Comparer :

1. leur vitesse,
2. le rayon de leur trajectoire,
3. leur période.

18.4 Questions sur les particules dans champs $\vec{E}$ et $\vec{B}$, d'après A.T.S. 2004 (★)

On considère une particule ponctuelle, de charge q et de masse m , de vitesse initiale $\vec{v}_0$ à l'entrée d'une zone où règnent un champ électrique $\vec{E}$ ou un champ magnétique $\vec{B}$.

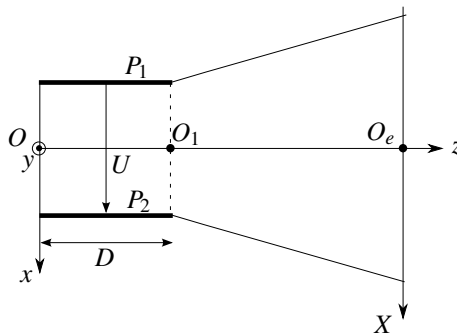
On suppose ces champs uniformes et indépendants du temps, et on néglige toute autre force que celles provoquées par ces champs.

1. La particule décrit une droite et possède une accélération constante a .
 - a. Déterminer la direction et la norme du ou des champs qui provoquent cette trajectoire.
 - b. Déterminer la position du point en fonction du temps.
2. La particule décrit une trajectoire circulaire de rayon R_0 , dans un plan xOy .
 - a. Déterminer la direction du ou des champs qui provoquent cette trajectoire.
 - b. Déterminer l'équation de la trajectoire et la relation entre la norme du champ, v_0 et R_0 .

On suggère d'utiliser les coordonnées polaires.

18.5 Déflexion électrique dans un oscilloscope, d'après Banque PT 2000 (★★)

Dans tout l'exercice on se place dans un référentiel galiléen, associé à un repère cartésien $(O, \vec{u}_x, \vec{u}_y, \vec{u}_z)$. Une zone de champ électrique uniforme (voir figure) est établie entre les plaques P_1 et P_2 (le champ est supposé nul en dehors et on néglige les effets de bord); la distance entre les plaques est d , la longueur des plaques D et la différence de potentiel est $U = V_{P_2} - V_{P_1}$ positive. Des électrons (charge $q = -e$, masse m) accélérés pénètrent en O dans la zone de champ électrique uniforme avec une vitesse $\vec{v}_0 = v_0 \vec{u}_z$ selon l'axe Oz .



1. Etablir l'expression de la force subie par les électrons en fonction de U , q , d et $\vec{u}_x$.
2. Etude du mouvement des électrons
 - a. Déterminer l'expression de la trajectoire $x = f(z)$ de l'électron dans la zone du champ en fonction de d , U et v_0 .
 - b. Déterminer le point de sortie K de la zone de champ ainsi que les composantes de la vitesse en ce point.
 - c. Montrer que dans la zone en dehors des plaques, le mouvement est rectiligne uniforme.
 - d. On note L la distance $O_1 O_e$ (voir figure). Déterminer l'abscisse X_P du point d'impact P de l'électron sur l'écran en fonction de U , v_0 , D , d et L .

APPROFONDIR

18.6 Spectromètre de masse d'après BTS chimiste (★★)

Le spectromètre de masse permet de mesurer la masse des particules chargées avec une telle précision qu'il peut servir à déterminer des compositions isotopiques d'éléments chimiques. Dans cet exercice, on va voir qu'il permet de déterminer la composition isotopique du mercure.

Une source émet des ions mercure $^{200}_{80}\text{Hg}^{2+}$ et $^{202}_{80}\text{Hg}^{2+}$. Ces ions passent dans le spectromètre de masse où ils sont accélérés puis séparés afin de mesurer leur rapport isotopique. Les données numériques nécessaires sont regroupées en fin d'exercice.

1. Accélération des ions.

Des ions de masse m et de charge $q > 0$ sont émis par une source située en F_1 , sans vitesse initiale. Ils sont accélérés entre F_1 et F_2 par une différence de potentiel U appliquée entre les plaques conductrices P_1 et P_2 .

a. Préciser la plaque de potentiel le plus élevé, représenter sur le schéma le champ accélérateur $\vec{E}_0$ qui règne dans l'entrefer séparant F_1 de F_2 . Calculer numériquement la valeur de $\|\vec{E}_0\|$.

b. Etablir l'expression littérale de la vitesse v_0 des ions sur la plaque P_2 .

c. Calculer numériquement v_{01} et v_{02} , les vitesses respectives des ions $^{200}_{80}\text{Hg}^{2+}$ et $^{202}_{80}\text{Hg}^{2+}$ à leur arrivée en F_2 .

Etant donné que l'hypothèse de vitesse nulle en F_1 est difficile à réaliser en pratique, il existe une certaine dispersion des vitesses en F_2 et il est nécessaire de réaliser un filtrage en vitesse pour améliorer les performances de l'appareil.

2. Filtre de vitesse.

Les ions traversent la plaque P_2 par la fente F_2 avec un vecteur vitesse perpendiculaire à P_2 . Ils entrent dans l'espace séparant P_2 et P_3 où règnent :

- un champ $\vec{E}_1$ uniforme situé dans le plan du schéma et parallèle à P_2 ;
- un champ $\vec{B}_1$ uniforme perpendiculaire au plan du schéma.

a. Sous quelle condition les ions peuvent-ils avoir une trajectoire rectiligne les amenant de F_2 à F_3 ?

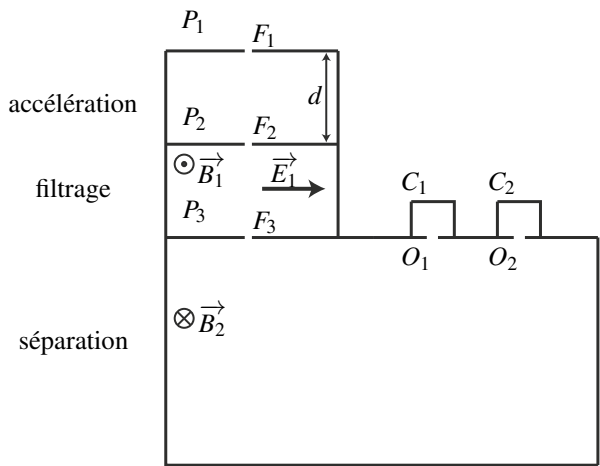
b. En déduire que seuls les ions de vitesse $v_0 = \frac{E_1}{B_1}$ parviennent en F_3 .

c. Calculer numériquement cette vitesse et en déduire quel isotope du mercure parvient en F_3 avec ces réglages.

Pour mesure la composition isotopique du mercure, on règle la valeur de $\vec{E}_1$ pour assurer le passage de $^{200}_{80}\text{Hg}^{2+}$ pendant 1 min puis on change sa valeur pour que les ions $^{202}_{80}\text{Hg}^{2+}$ passent pendant 1 min. Pendant cette opération, la valeur de $\vec{B}_1$ reste constante.

3. Séparation des ions.

Après F_3 , les ions pénètrent dans une région où ne règne qu'un champ magnétique uniforme $\vec{B}_2$ normal au plan du schéma. Ils sont déviés vers les collecteurs C_1 et C_2 .



- a. Montrer que le mouvement d'un ion dans cette région est uniforme.
- b. Sachant que la trajectoire des ions est circulaire, déterminer son rayon R_1 pour les ions $^{200}_{80}\text{Hg}^{2+}$ et R_2 pour les ions $^{202}_{80}\text{Hg}^{2+}$.
- c. Déterminer le collecteur (C_1 ou C_2) qui reçoit $^{200}_{80}\text{Hg}^{2+}$ et celui qui reçoit $^{202}_{80}\text{Hg}^{2+}$.
- d. La distance δ qui sépare les points O_1 et O_2 paraît-elle suffisante pour installer des détecteurs de particules.
- e. Les quantités d'électricité reçues en 1 minute par les collecteurs C_1 et C_2 sont $Q_1 = 1,20 \cdot 10^{-7} \text{ C}$ et $Q_2 = 3,5 \cdot 10^{-8} \text{ C}$. Déterminer la composition du mélange d'ion et en déduire la masse atomique du mercure.

Données numériques utiles : $d = 1,00 \text{ m}$; $U = 1,00 \cdot 10^4 \text{ V}$; $e = 1,60 \cdot 10^{-19} \text{ C}$; unité de masse atomique : $1u = 1,67 \cdot 10^{-27} \text{ kg}$ (masse d'un nucléon) ; $E_1 = 5,30 \cdot 10^4 \text{ V} \cdot \text{m}^{-1}$; $B_1 = 0,383 \text{ T}$; $B_2 = 0,200 \text{ T}$; $F_3O_1 = 1,44 \text{ m}$; $F_3O_2 = 1,45 \text{ m}$.

18.7 Mouvement de gouttelettes chargées, d'après ENAC 2005 (★★)

On disperse un brouillard de fines gouttelettes sphériques d'huile, de masse volumique $\rho_h = 1,3 \cdot 10^3 \text{ kg} \cdot \text{m}^{-3}$, dans l'espace séparant les deux plaques horizontales d'un condensateur plan, distantes de $d = 2 \cdot 10^{-2} \text{ m}$. Les gouttelettes sont chargées négativement et sans vitesse initiale. Toutes les gouttelettes ont même rayon R mais pas forcément la même charge $q < 0$. En l'absence de champ électrique E , une gouttelette est soumise à son poids ($g = 9,81 \text{ m} \cdot \text{s}^{-2}$), à la poussée d'Archimède de l'air ambiant de masse volumique $\rho_a = 1,3 \text{ kg} \cdot \text{m}^{-3}$ et à une force de frottement visqueux $\vec{f} = -k \vec{v}$, avec $k = \alpha R$ et $\alpha = 3,4 \cdot 10^{-4} \text{ S.I.}$. L'accélération de la pesanteur $\vec{g}$ sera prise égale à $9,81 \text{ m} \cdot \text{s}^{-2}$.

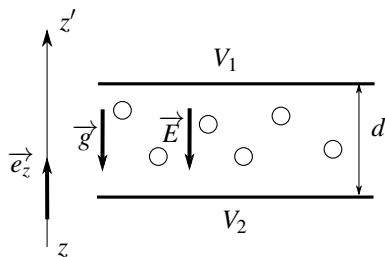


Figure 18.11

CHAPITRE 18 – MOUVEMENT D'UNE PARTICULE CHARGÉE DANS UN CHAMP

1. Mouvement en l'absence de champ électrique.

- Déterminer la vitesse limite $\vec{v}_0$
- Déterminer l'expression de la vitesse des gouttes $\vec{v}(t)$. On fera apparaître un temps caractéristique τ .
- On mesure $v_0 = 2.10^{-4} \text{ m.s}^{-1}$, déterminer la valeur de k .

2. On applique une différence de potentiel $U = V_1 - V_2$ de manière à avoir un champ électrique $\vec{E}$ dirigé vers le bas.

- Déterminer l'expression de $\vec{E}$.
- Une gouttelette est immobilisée pour $U = 3200 \text{ V}$. Calculer la charge q .

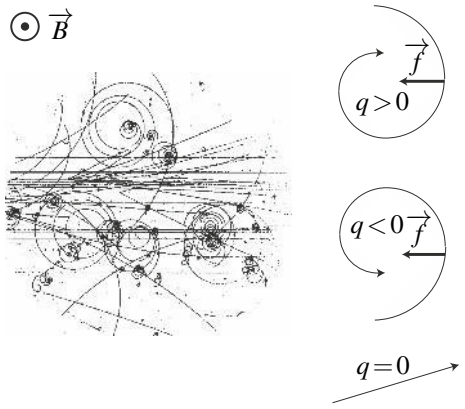
CORRIGÉS

18.1 Chambre à bulle

1. La force magnétique $\vec{f}$ est nécessairement dirigée vers l'intérieur de la concavité de la trajectoire. On peut donc la tracer sur la figure. Or $\vec{f} = q\vec{v} \wedge \vec{B}$, donc le sens de $\vec{B}$ permet de déduire le sens de $q\vec{v}$ grâce à la « règle de la main droite ». Le sens de $\vec{v}$ permet alors de conclure.

La trajectoire 1 correspond à une charge positive, la 2 à une charge négative et la 3 à une particule neutre.

2. En présence d'un champ magnétique seul, une particule chargée suit une trajectoire circulaire et uniforme de rayon $R = \frac{mv}{qB}$ proportionnel à v . La présence du liquide provoque une diminution progressive de l'énergie cinétique et donc de la vitesse de la particule. Il s'en suit une diminution progressive du rayon de la trajectoire.



18.2 Grandeurs cinétiques d'une particule chargée dans un champ magnétique

1. La force magnétique est orthogonale à la vitesse. Sa puissance est donc nulle. D'après le théorème de l'énergie cinétique, l'énergie cinétique de la particule est conservée et la norme de sa vitesse également.

2. Si le module de la vitesse est constant, il n'en est pas de même de sa direction. Par conséquent, la quantité de mouvement peut changer de direction en conservant son module.

3. On a vu que la trajectoire est circulaire et uniforme à la vitesse angulaire $\omega_c = \frac{|q|B}{m}$. Le moment cinétique de la particule par rapport au centre O de la trajectoire circulaire vaut : $\vec{L}_0 = \vec{OM} \wedge m\vec{v} = mR^2\omega_c \vec{u}_z$. Il est donc constant.

18.3 Action d'un champ magnétique sur un proton ou sur un électron

1. L'énergie cinétique s'écrit : $E_c = \frac{1}{2}mv^2$. Comme $m_e < m_p$, on en déduit que la vitesse de l'électron est plus grande que celle du proton : $v_e > v_p$.

2. Le rayon de la trajectoire s'exprime par : $R = \frac{mv}{|q|B}$ donc $\frac{R_e}{R_p} = \frac{m_e v_e}{m_p v_p} = \frac{E_{c_e} v_p}{E_{c_p} v_e} = \frac{v_p}{v_e}$. Le rayon de la trajectoire de l'électron est donc plus petit que celui du proton : $R_e < R_p$.

CHAPITRE 18 – MOUVEMENT D'UNE PARTICULE CHARGÉE DANS UN CHAMP

3. La pulsation cyclotron vaut : $\omega_c = \frac{|q|B}{m}$. La pulsation du mouvement de l'électron est donc plus grande que celle du proton. Comme la pulsation est inversement proportionnelle à la période, la période de l'électron est plus petite que celle du proton : $T_e < T_p$.

18.4 Questions sur les particules dans champs $\vec{E}$ et $\vec{B}$

On étudie la particule M de masse m et de charge q assimilée à un point matériel dans le référentiel du laboratoire supposé galiléen. Cette particule est soumise à la force de Lorentz $\vec{f} = q(\vec{E} + \vec{v} \wedge \vec{B})$.

1. La trajectoire est rectiligne et uniformément accélérée.

a. L'accélération $\vec{a}$ est constante, soit :

$$\frac{d\vec{v}}{dt} = \vec{a} = c\vec{t}e \Rightarrow \vec{v} = \vec{a}t + \vec{v}_0.$$

La norme de $\vec{v}$ varie donc l'énergie cinétique aussi. Or seule la force électrique travaille (la force magnétique $q(\vec{v} \wedge \vec{B})$ est perpendiculaire à $\vec{v}$ donc à la trajectoire), le champ est un champ électrique.

Pour que la trajectoire soit rectiligne, il faut, d'après l'expression de $\vec{v}$ que $\vec{a}$ soit colinéaire à $\vec{v}_0$. Or $\vec{a} = \frac{q}{m}\vec{E}$, donc $\vec{E}$ est colinéaire à $\vec{v}_0$.

b. On note O le centre du repère. On intègre l'expression de la vitesse pour avoir la position $\vec{OM}$ de la particule :

$$\vec{OM} = \frac{q}{2m}t^2\vec{E} + t\vec{v}_0 + \vec{OM}_0,$$

M_0 étant la position initiale du point.

2. La trajectoire est circulaire dans le plan (Oxy) .

a. La trajectoire circulaire est celle d'une charge dans un champ magnétique perpendiculaire à la vitesse initiale. On en déduit que $\vec{B}$ est suivant Oz et que $\vec{v}_0$ est dans le plan xOy .

b. La trajectoire étant circulaire, la vitesse en coordonnées polaires a pour expression $\vec{v} = R_0\dot{\theta}\vec{u}_\theta$, et l'accélération se réduit à $\vec{a} = -R_0\dot{\theta}^2\vec{u}_r + R_0\ddot{\theta}\vec{u}_\theta$. La relation fondamentale de la dynamique donne :

$$m\vec{a} = q\vec{v} \wedge \vec{B},$$

Or $\vec{B} = B\vec{u}_z$ et donc $q\vec{v} \wedge \vec{B} = qR_0\dot{\theta}B\vec{u}_r$ puis, en projetant sur la base $(\vec{u}_r, \vec{u}_\theta)$:

$$\begin{cases} -mR_0\dot{\theta}^2 &= qR_0\dot{\theta}B \\ R_0\ddot{\theta} &= 0. \end{cases}$$

On obtient alors $\dot{\theta} = -\frac{qB}{m}$ = constante. Si la charge est positive elle tourne dans le sens rétrograde (horaire) par rapport à Oz . Puisque $\dot{\theta}$ est constante, le mouvement est circulaire uniforme et $v_0 = R_0|\dot{\theta}|$ d'où $R_0 = \frac{mv_0}{qB}$.

18.5 Déflexion électrique dans un oscilloscope

1. On néglige le poids. Les électrons ne sont soumis qu'à la force électrique. Le potentiel de la plaque P_2 est supérieur à celui de la plaque P_1 puisque $U = V_{P_2} - V_{P_1} > 0$. $\vec{E}$ est de norme $\frac{U}{d}$ et dirigé selon les potentiels décroissants donc selon $-\vec{u}_x$. On en déduit :

$$\vec{E} = -\frac{U}{d}\vec{u}_x \quad \text{puis} \quad \vec{F} = q\vec{E} = -q\frac{U}{d}\vec{u}_x = e\frac{U}{d}\vec{u}_x$$

puisque la charge de l'électron est $q = -e$.

2. Mouvement d'un électron.

a. On applique la relation fondamentale de la dynamique à l'électron dans le référentiel de l'oscilloscope galiléen : $m\vec{a} = \vec{F}$. La force n'a pas de composante sur Oy et comme la vitesse initiale n'a pas de composantes non plus sur Oy on en déduit que le mouvement est dans le plan xOz . La projection de la relation fondamentale avec l'expression de $\vec{F} = \vec{E}$ et les intégrations successives donnent :

$$\begin{cases} m\ddot{x} = e\frac{U}{d} \\ m\ddot{z} = 0 \end{cases} \Rightarrow \begin{cases} \dot{x} = \frac{e}{m}\frac{U}{d}t(+cte_1) \\ \dot{z} = cte_2 = v_0 \end{cases} \Rightarrow \begin{cases} x = \frac{e}{m}\frac{U}{d}\frac{t^2}{2}(+cte_3) \\ z = v_0t(+cte_4) \end{cases}$$

Les constantes entre parenthèses (cte_1 , cte_3 , cte_4) sont déterminées nulles grâce aux conditions initiales. En éliminant t entre les équations paramétriques en x et z , on obtient :

$$x = \frac{e}{m}\frac{U}{d}\frac{z^2}{2v_0^2}.$$

Il s'agit d'une parabole.

b. Le point de sortie correspond à $z_K = D$, soit dans l'équation précédente $x_K = \frac{e}{m}\frac{U}{d}\frac{D^2}{2v_0^2}$. D'après les équations paramétriques obtenues à la question (2), l'instant de passage en K est $t_K = z_K/v_0$ soit D/v_0 . les composantes de la vitesse en K sont obtenues en reportant t_K dans $\dot{x}$ et $\dot{z}$, soit :

$$\begin{cases} \dot{x}_K = \frac{e}{m}\frac{U}{d}\frac{D}{v_0} \\ \dot{z}_K = v_0 \end{cases}$$

c. En dehors des plaques, puisqu'on néglige l'effet du champ de pesanteur et que l'on suppose le champ électrique nul, aucune force ne s'exerce sur l'électron : il est isolé. D'après le principe d'inertie, sa trajectoire est rectiligne uniforme.

d. Dans le triangle O_eJP (figure 18.12), $X_P = O_eP = (JO_1 + L)\tan\theta$. Il faut donc exprimer $\tan\theta$ et JO_1 sachant que le segment JP est tangent à la parabole en K . On peut écrire :

$$\tan\theta = \frac{O_1K}{JO_1} = \frac{x_K}{JO_1} = \frac{e}{m}\frac{U}{d}\frac{D^2}{2v_0^2}\frac{1}{JO_1}$$

CHAPITRE 18 – MOUVEMENT D’UNE PARTICULE CHARGÉE DANS UN CHAMP

et

$$\tan \theta = \left(\frac{dx}{dz} \right)_K = \left(\frac{\dot{x}}{\dot{z}} \right)_K = \frac{e}{m} \frac{U}{d} \frac{D}{v_0^2}.$$

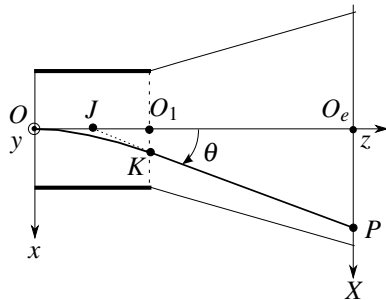


Figure 18.12

e. En combinant les deux équations, on trouve (c’est d’ailleurs une propriété de la parabole) que $JO_1 = D/2$ et donc :

$$X_P = \left(\frac{D}{2} + L \right) \frac{eD}{mdv_0^2} U.$$

La déviation est proportionnelle à la tension U .

18.6 Spectromètre de masse

Dans tout cet exercice, on étudie l’ion de masse m et de charge q dans le référentiel du spectromètre de masse supposé galiléen.

1. Accélération des ions.

a. Pour accélérer un ion de charge positive, le champ $\vec{E}_0$ doit être dirigé de F_1 vers F_2 afin que la force électrique soit elle aussi dans ce sens. Or le champ électrique est dirigé du potentiel le plus élevé vers le potentiel le plus faible donc la plaque P_1 est la plaque qui possède le potentiel le plus élevé.

En notant $\vec{u}_x$ le vecteur unitaire dirigé de F_1 vers F_2 , le champ $\vec{E}_0$ vaut alors :

$$\vec{E}_0 = \frac{U}{d} \vec{u}_x \quad \text{et} \quad \|\vec{E}_0\| = \frac{U}{d} = 1,00.10^4 \text{ V}\cdot\text{m}^{-1}.$$

On peut choisir de manière arbitraire de considérer que le potentiel de P_1 est U et celui de P_2 est nul (choix arbitraire de la masse).

b. On peut négliger le poids de l’ion. La seule force est la force électrique $\vec{f} = q\vec{E}_0$ conservative, qui dérive de l’énergie potentielle $E_p = qV$.

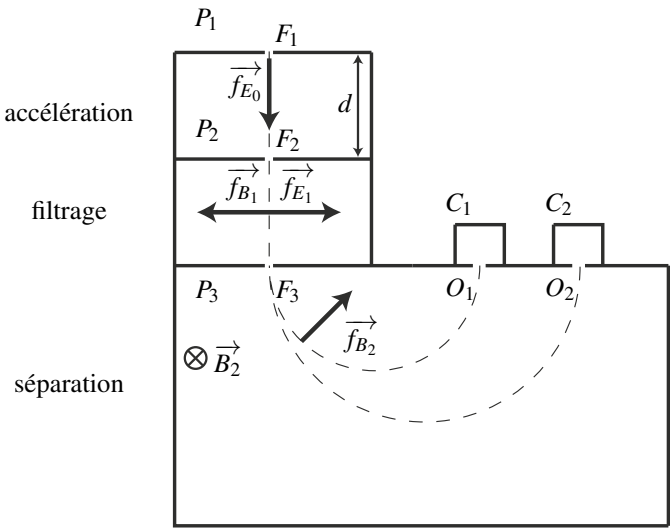
On est dans un cas de conservation de l’énergie mécanique que l’on écrit aux points F_1 de potentiel U et de vitesse nulle et au point F_2 de potentiel 0 et de vitesse v_0 :

$$qU = \frac{1}{2}mv_0^2 \quad \Rightarrow \quad v_0 = \sqrt{\frac{2qU}{m}}.$$

c. Pour l'ion $^{200}_{80}\text{Hg}^{2+}$, $m = 200u$ et $q = 2e$. Pour $^{202}_{80}\text{Hg}^{2+}$, $m = 202u$ et $q = 2e$. On trouve alors :

$$v_{01} = 1,384 \cdot 10^5 \text{ m}\cdot\text{s}^{-1} \quad \text{et} \quad v_{02} = 1,377 \cdot 10^5 \text{ m}\cdot\text{s}^{-1}.$$

2. Filtre de vitesse.



a. Pour que les ions aient une trajectoire rectiligne, il est nécessaire que la force magnétique $\vec{f}_B$ soit opposée à la force électrique $\vec{f}_E$ dans cette zone. Les ions sont alors pseudo-isolés et, d'après le principe d'inertie, ils sont animés d'un mouvement rectiligne et uniforme à la vitesse $\vec{v} = v_0 \vec{u}_x$.

b. Il faut donc que $q\vec{E}_1 + q\vec{v}_0 \wedge \vec{B}_1 = \vec{0}$, ce qui, compte tenu des directions respectives de $\vec{E}_1$, $\vec{B}_1$ et $\vec{v}$ et en notant $E_1 = \|\vec{E}_1\|$ et $B_1 = \|\vec{B}_1\|$ donne :

$$E_1 = v_0 B_1 \quad \Rightarrow \quad v_0 = \frac{E_1}{B_1}.$$

c. On trouve $v_0 = 1,384 \cdot 10^5 \text{ m}\cdot\text{s}^{-1}$. Avec ces réglages, c'est l'isotope $^{200}_{80}\text{Hg}^{2+}$ qui parvient en F_3 . L'isotope $^{202}_{80}\text{Hg}^{2+}$, un peu moins rapide va être dévié sur la droite car la force électrique va l'emporter sur la force magnétique. Il ne pourra pas arriver dans la chambre de séparation en F_3 .

3. Séparation des ions.

a. L'ion n'est soumis qu'à une force magnétique $\vec{f}_{B_2} = q\vec{v} \wedge \vec{B}_2$ créée par un champ magnétique $\vec{B}_2$ uniforme. La puissance de la force magnétique est nulle. Le théorème de la puissance cinétique implique donc que l'énergie cinétique de l'ion et par suite la norme de sa vitesse est conservée. Le mouvement dans cette région est donc uniforme.

CHAPITRE 18 – MOUVEMENT D'UNE PARTICULE CHARGÉE DANS UN CHAMP

b. On étudie cette trajectoire dans un repère polaire centré sur le centre de la trajectoire circulaire et uniforme à la vitesse v_0 . Le principe fondamental de la dynamique projeté sur $\vec{u}_r$ donne :

$$m \frac{v_0^2}{R} = q v_0 B_2 \Rightarrow R = \frac{m v_0}{q B_2}.$$

Numériquement, $R_1 = 0,722$ m et $R_2 = 0,726$ m.

c. Le collecteur C_1 est placé à $2R_1 = 1,44$ m de F_3 . Il reçoit les ions $^{200}_{80}\text{Hg}^{2+}$. Le collecteur C_2 est placé à $2R_2 = 1,45$ m de F_3 . Il reçoit les ions $^{202}_{80}\text{Hg}^{2+}$.

d. La distance de $\delta = 2(R_2 - R_1) = 8$ mm qui sépare les points O_1 et O_2 est suffisante pour installer des détecteurs de particules et permettre des mesures.

e. La quantité d'électricité reçue en 1 min par un collecteur est proportionnelle au nombre d'ions reçus. La proportion d'ion $^{200}_{80}\text{Hg}^{2+}$ vaut donc $\frac{Q_1}{Q_1 + Q_2} = 77,4\%$; celle d'ion $^{202}_{80}\text{Hg}^{2+}$ vaut $\frac{Q_2}{Q_1 + Q_2} = 22,6\%$.

La masse atomique du mercure vaut alors (on rappelle que une mole de nucléon pèse 1 gramme car 1 mole de carbone 12 pèse 12 grammes) :

$$M = 0,774 \times 200 + 0,226 \times 202 = 200,5 \text{ g} \cdot \text{mol}^{-1}.$$

18.7 Mouvement de gouttelettes chargées

1. On étudie la gouttelette assimilée à un point matériel dans le référentiel du laboratoire galiléen.

a. Cette gouttelette est soumise à : son poids, la poussée d'Archimède (opposée du poids du volume d'air déplacé) aux frottements. Le principe fondamental de la dynamique (P.F.D.) appliquée à une goutte donne :

$$m \vec{a} = \frac{4}{3} \pi R^3 \rho_h \vec{g} - \frac{4}{3} \pi R^3 \rho_a \vec{g} - k \vec{v}.$$

La vitesse limite est obtenue lorsque l'accélération devient nulle :

$$\vec{v}_0 = \frac{4}{3} \pi R^3 \frac{1}{k} (\rho_h - \rho_a) \vec{g} \simeq \frac{4}{3} \pi R^3 \frac{1}{k} \rho_h \vec{g} \simeq \frac{4}{3} \pi R^2 \frac{1}{\alpha} \rho_h \vec{g}$$

b. On peut réécrire le P.F.D avec $\vec{v}_0$:

$$m \frac{d\vec{v}}{dt} = k(\vec{v}_0 - \vec{v}) \text{ ou } \frac{d\vec{v}}{dt} + \frac{\vec{v}}{\tau} = \frac{\vec{v}_0}{\tau},$$

où l'on a défini le temps caractéristique $\tau = m/k$.

La solution de cette équation différentielle est :

$$\vec{v}(t) = \vec{A} \exp\left(-\frac{t}{\tau}\right) + \vec{v}_0.$$

A $t = 0$ la vitesse est nulle donc $\vec{0} = \vec{A} + \vec{v}_0$. On en déduit l'expression finale de $\vec{v}$:

$$\vec{v}(t) = \left(1 - \exp\left(-\frac{t}{\tau}\right)\right) \vec{v}_0.$$

- c. Application numérique : avec l'expression de $\vec{v}_0$ on trouve $R = 1,13 \cdot 10^{-6}$ m.
2. a. Le champ électrique $\vec{E}$ étant dirigé vers le bas, on l'écrit $\vec{E} = -E\vec{u}_z$ avec $E > 0$. Or, le champ électrique est suivant les potentiels décroissants, on en déduit que $V_1 > V_2$ et $U > 0$. Par ailleurs, $\|\vec{E}\| = \frac{U}{d}$ d'où on tire :

$$\vec{E} = -\frac{U}{d}\vec{u}_z.$$

- b. La force électrique s'exerçant sur la goutte est $\vec{F} = q\vec{E}$ avec $q < 0$. A l'équilibre, la force de frottement est nulle, donc la force électrique équilibre le poids et la poussée d'Archimède :

$$\vec{0} = q\vec{E} + \frac{4}{3}\pi R^3(\rho_h - \rho_a)\vec{g} \Rightarrow 0 = -q\frac{U}{d} - \frac{4}{3}\pi R^3(\rho_h - \rho_a)g$$

d'où :

$$q = -\frac{d}{U}\frac{4}{3}\pi R^3(\rho_h - \rho_a)g \Rightarrow q = -4,8 \cdot 10^{-19} \text{ C}.$$

CHAPITRE 18 – MOUVEMENT D’UNE PARTICULE CHARGÉE DANS UN CHAMP

Troisième partie

Mécanique 2

Trouver la bibliothèque des livres ici : <https://1024terabox.com/s/1mg0wxl22G8NjM8MA2RzgFw>

Moment cinétique et solide en rotation

19

Le principe fondamental de la dynamique relie la variation de la quantité de mouvement aux forces appliquées. Pourtant, pour les systèmes en rotation, la notion de force n'est pas toujours la plus pertinente. Dans ce chapitre, on introduit deux nouvelles grandeurs mécaniques : le moment cinétique et le moment d'une force. On relie alors la variation du moment cinétique aux moments des forces appliquées au système.

1 Observations préliminaires

1.1 Exemples introductifs

Lorsque l'on cherche à ouvrir une porte, on applique une force sur la poignée. Le positionnement de la poignée, ainsi que l'orientation de la force appliquée ne sont pas laissés au hasard. Pour être efficace, la force doit être appliquée le plus loin possible des gonds de la porte qui matérialisent son axe de rotation. De plus, elle doit être appliquée perpendiculairement à la porte. Cette configuration permet d'augmenter au maximum le bras de levier de la force, qui est la distance séparant l'axe de rotation de la porte de la droite d'application de la force.

Remarque

La droite droite d'application de la force est la droite passant par le point d'application de la force et parallèle à cette dernière.

On retrouve la notion de bras de levier dans de nombreux objets de la vie courante et notamment dans l'un des plus anciens instruments de pesée, la balance romaine (voir figure 19.1). Cet instrument est constitué de deux fléaux de longueurs différentes et d'un contrepoids mobile de masse donnée. La position de la masse à peser est fixe sur le fléau le plus court, celle du contrepoids est variable. On déplace le contrepoids jusqu'à l'obtention de l'équilibre et sa position, repérée grâce à des graduations gravées sur le fléau le plus long, permet de mesurer la masse à peser. La figure 19.1 illustre le fait que l'augmentation du bras de levier permet d'équilibrer une masse de 5 kg ou de 10 kg à l'aide d'un contrepoids identique de 1 kg.

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

1.2 Notion intuitive de bras de levier

Les exemples précédents montrent que la notion de force n’est pas suffisante pour étudier les systèmes en rotation autour d’un axe. Le bras de levier est également déterminant. De fait, la notion pertinente pour étudier ce type de problèmes est le produit « force $\times$ bras de levier ». L’objet de ce chapitre est de mettre en place de nouvelles grandeurs physiques qui permettent d’étudier les systèmes en rotation.

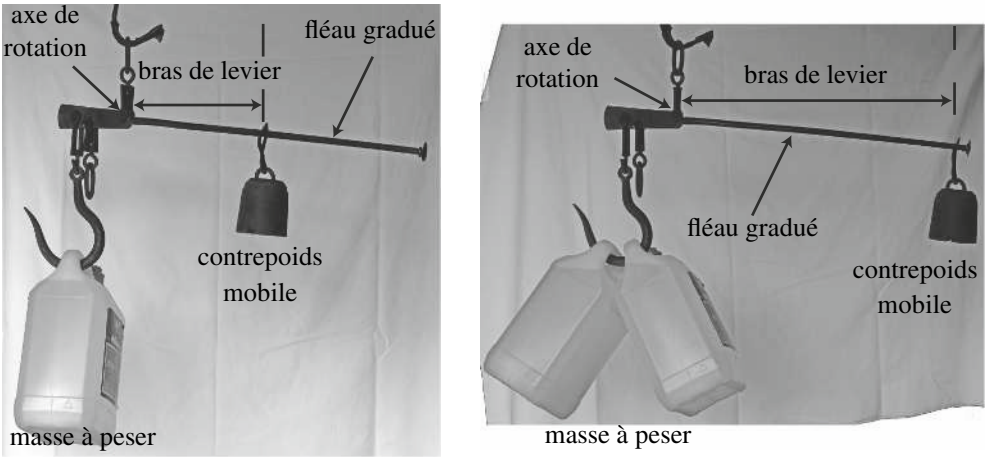


Figure 19.1 – Balance romaine pesant une masse de 5 kg (à gauche) ou 10 kg (à droite).

2 Moment cinétique d’un point matériel

2.1 Définition du moment cinétique

On considère un point matériel M de masse m animé d’une vitesse $\vec{v}$ dans le référentiel $\mathcal{R}$ galiléen. On note $\vec{p} = m\vec{v}$ sa quantité de mouvement. On considère également un axe orienté $\Delta = (O, \vec{u}_\Delta)$ déterminé par un point O et un vecteur unitaire $\vec{u}_\Delta$ dont le sens précise l’orientation de l’axe.

a) Moment cinétique par rapport à un point O

Le **moment cinétique de M par rapport à un point O** est le vecteur défini par le produit vectoriel :

$$\vec{L}_O = \vec{OM} \wedge \vec{p} = m\vec{OM} \wedge \vec{v}.$$

Sa norme se mesure en $\text{kg}\cdot\text{m}^2\cdot\text{s}^{-1} = \text{J}\cdot\text{s}$.

Le moment cinétique par rapport à un point O est défini à partir de la vitesse de M , il dépend donc du référentiel dans lequel on le détermine.

Propriétés du moment cinétique par rapport à un point De par sa définition à partir d'un produit vectoriel, $\vec{L}_O$ est perpendiculaire aux vecteurs $\vec{OM}$ et $\vec{v}$. En conséquence :

- si le mouvement de M est plan et que O appartient au plan du mouvement, le vecteur $\vec{L}_O$ est perpendiculaire à ce plan à tout instant. Sa direction est donc fixe et perpendiculaire au plan du mouvement. La réciproque est vraie.
- si le mouvement de M est rectiligne et inscrit sur une droite $\mathcal{D}$ passant par O , $\vec{L}_O$ est nul à tout instant puisque les vecteurs $\vec{OM}$ et $\vec{v}$ sont colinéaires à tout instant. La réciproque est vraie.

Il faut également noter que le moment cinétique dépend du point par rapport auquel on le calcule et que l'on indique en indice. En effet,

$$\vec{L}_B = \vec{BM} \wedge \vec{p} = (\vec{BA} + \vec{AM}) \wedge \vec{p} = \vec{BA} \wedge \vec{p} + \vec{AM} \wedge \vec{p},$$

que l'on peut écrire :

$$\vec{L}_B = \vec{L}_A + \vec{BA} \wedge \vec{p}. \tag{19.1}$$

b) Moment cinétique par rapport à un axe orienté Δ

Le **moment cinétique de M par rapport à l'axe orienté $\Delta = (O, \vec{u}_\Delta)$** est la projection orthogonale de $\vec{L}_O$ sur l'axe Δ :

$$L_\Delta = \vec{L}_O \cdot \vec{u}_\Delta = m(\vec{OM} \wedge \vec{v}) \cdot \vec{u}_\Delta.$$

Comme le moment cinétique par rapport à un point, le moment cinétique par rapport à un axe dépend du référentiel d'étude et se mesure en J.s. Par contre, il ne dépend pas du choix du point O appartenant à l'axe Δ utilisé pour le calculer mais seulement de la direction et de l'orientation de Δ .

En effet, si l'on considère deux points A et B appartenant à l'axe Δ , le produit scalaire de la relation (19.1) par $\vec{u}_\Delta$ donne :

$$\vec{L}_B \cdot \vec{u}_\Delta = \vec{L}_A \cdot \vec{u}_\Delta + (\vec{BA} \wedge \vec{p}) \cdot \vec{u}_\Delta.$$

Or A et B appartiennent à Δ donc $\vec{AB}$ est colinéaire à $\vec{u}_\Delta$ et $\vec{BA} \wedge \vec{p}$ est perpendiculaire à $\vec{u}_\Delta$. Le terme $(\vec{BA} \wedge \vec{p}) \cdot \vec{u}_\Delta$ s'annule et finalement :

$$\vec{L}_B \cdot \vec{u}_\Delta = \vec{L}_A \cdot \vec{u}_\Delta,$$

ce qui prouve que le point de l'axe choisi pour calculer L_Δ n'a pas d'importance.

c) Cas où le point matériel est en mouvement circulaire

Lorsque M est en mouvement circulaire sur un cercle de centre O et de rayon R . On a intérêt à le repérer en coordonnées cylindriques de centre O , d'axe (Oz) perpendiculaire au plan du

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

cercle, et d'angle polaire θ . On utilise alors les relations établies dans le chapitre *Cinématique du point* :

$$\overrightarrow{OM} = R\overrightarrow{u_r} \quad \text{et} \quad \overrightarrow{v} = R\dot{\theta}\overrightarrow{u_\theta}.$$

On considère $\Delta = (O, \overrightarrow{u_z})$ et on calcule $\overrightarrow{L_O}$ et $L_\Delta = L_{(Oz)}$:

$$\overrightarrow{L_O} = m\overrightarrow{OM} \wedge \overrightarrow{v} = m(R\overrightarrow{u_r}) \wedge (R\dot{\theta}\overrightarrow{u_\theta}) = mR^2\dot{\theta}\overrightarrow{u_z}$$

et

$$L_{(Oz)}(M) = \overrightarrow{L_O} \cdot \overrightarrow{u_z} = mR^2\dot{\theta}.$$

On peut trouver la direction de $\overrightarrow{L_O}$ à l'aide de la « règle de la main droite » représentée sur la figure 19.2. Ainsi :

- lorsque la révolution du point M se fait dans le sens direct autour de $\overrightarrow{u_z}$, $\dot{\theta} > 0$ et $\overrightarrow{L_O}$ est selon $+\overrightarrow{u_z}$;
- lorsque la révolution du point M se fait dans le sens indirect autour de $\overrightarrow{u_z}$, $\dot{\theta} < 0$ et $\overrightarrow{L_O}$ est selon $-\overrightarrow{u_z}$.

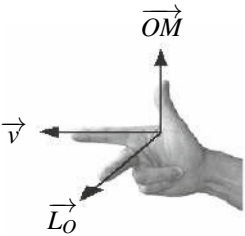


Figure 19.2 – Règle de la main droite.

Le signe de $L_{(Oz)}(M)$ permet donc de déterminer le sens de la révolution circulaire de M :

- si $L_{(Oz)}(M) > 0$, la révolution se fait dans le sens direct dans le plan orienté par $\overrightarrow{u_z}$;
- si $L_{(Oz)}(M) < 0$, la révolution se fait dans le sens indirect dans le plan orienté par $\overrightarrow{u_z}$.

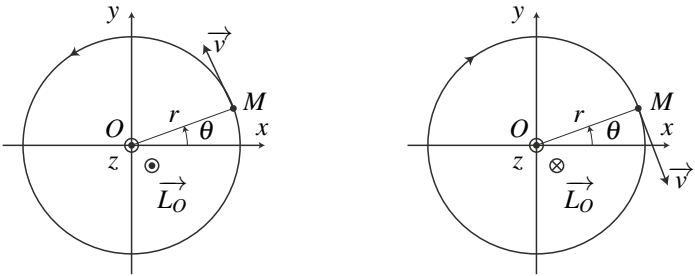


Figure 19.3 – Moment cinétique en O d'un mobile en révolution circulaire de centre O .

Remarque

La définition du moment cinétique en O implique que la trajectoire de M « s'enroule » dans le sens direct autour de la droite orientée $(O, \overrightarrow{L_O})$.

d) Notion de moment d'inertie

Plus généralement, lorsque l'axe Δ est fixe, on peut le faire coïncider avec l'axe (Oz) et repérer M à l'aide de ses coordonnées cylindriques : $\overrightarrow{OM} = r\overrightarrow{u_r} + z\overrightarrow{u_z}$ et $\overrightarrow{v} = \dot{r}\overrightarrow{u_r} + r\dot{\theta}\overrightarrow{u_\theta} + \dot{z}\overrightarrow{u_z}$. On peut alors calculer $L_{(Oz)} = \overrightarrow{L_O} \cdot \overrightarrow{u_z}$ qui correspond à la composante selon $\overrightarrow{u_z}$ du moment cinétique de M par rapport à O :

$$\begin{aligned}\vec{L}_O &= m\vec{OM} \wedge \vec{v} = m(r\vec{u}_r + z\vec{u}_z) \wedge (\dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta + \dot{z}\vec{u}_z) \\ &= m\left(r^2\dot{\theta}\vec{u}_z + (\dot{r}z - r\dot{z})\vec{u}_\theta - zr\dot{\theta}\vec{u}_r\right),\end{aligned}$$

soit :

$$L_{(Oz)} = mr^2\dot{\theta}.$$

En coordonnées cylindriques d'axe (Oz) , on définit le **moment d'inertie** d'un point M de coordonnées (r, θ, z) par rapport à l'axe (Oz) :

$$J_{(Oz)}(M) = mr^2.$$

Le moment cinétique de M par rapport à (Oz) est alors égal au produit du moment d'inertie $J_{(Oz)}(M)$ par la vitesse angulaire $\dot{\theta}$:

$$L_{(Oz)} = J_{(Oz)}(M)\dot{\theta}. \tag{19.2}$$

3 Moment cinétique d'un solide ou d'un système de points

Dans le chapitre de *Cinématique du solide*, on a décrit deux grands types de mouvements de solides : les mouvements de translation et les mouvements de rotation autour d'un axe fixe. Les mouvements de translation s'étudient à l'aide du principe fondamental de la dynamique. L'étude des mouvements de rotation autour d'un axe fixe utilise le moment cinétique par rapport à l'axe de rotation fixe Δ . C'est pourquoi, dans ce paragraphe, on se restreint au moment cinétique L_Δ d'un solide ou d'un système de points par rapport à un axe orienté Δ .

3.1 Cas d'un système déformable

a) Moment cinétique par rapport à un axe orienté

On considère un système constitué de plusieurs points matériels M_i de masses m_i de moments cinétiques par rapport à l'axe orienté Δ : $L_\Delta(M_i)$. Le moment cinétique du système de points est obtenu par sommation des moments cinétiques de chacun des points :

$$L_\Delta = \sum_i L_\Delta(M_i).$$

Les moments cinétiques par rapport à Δ étant algébriques, il faut être rigoureux sur les signes.

b) Cas des coordonnées cylindriques

Lorsque l'axe Δ est fixe, on utilise généralement les coordonnées cylindriques d'axe (Oz) confondu avec Δ . Le moment cinétique de chaque point M_i est donné par la relation (19.2) : $L_{(Oz)}(M_i) = m_i r_i^2 \dot{\theta} = J_{(Oz)}(M_i) \dot{\theta}_i$ où $J_{(Oz)}(M_i) = m_i r_i^2$ est le moment d'inertie du point M_i par

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

rapport à (Oz) et $\dot{\theta}_i$ sa vitesse angulaire :

$$L_{(Oz)} = \sum_i L_{(Oz)}(M_i) = \sum_i J_{(Oz)}(M_i) \dot{\theta}_i. \tag{19.3}$$

Exemple

On considère le cas où deux points M_1 et M_2 de même masse m sont en rotation circulaire et uniforme de centre O et de rayon R dans le plan (Oxy) . Le moment d’inertie de chacun des points vaut alors $J_{(Oz)} = mR^2$. On envisage deux cas :

- le cas où les points M_1 et M_2 ont la même vitesse angulaire $\dot{\theta}$. Le moment cinétique par rapport à (Oz) de chacun des points vaut alors $L_{(Oz)}(M_i) = mR^2 \dot{\theta}$. Celui du système vaut donc $L_{(Oz)} = 2mR^2 \dot{\theta}$;
- le cas où ils ont des vitesses angulaires opposées et $L_{(Oz)}(M_1) = -L_{(Oz)}(M_2)$. Le moment cinétique du système par rapport à (Oz) vaut donc $L_{(Oz)} = 0$.

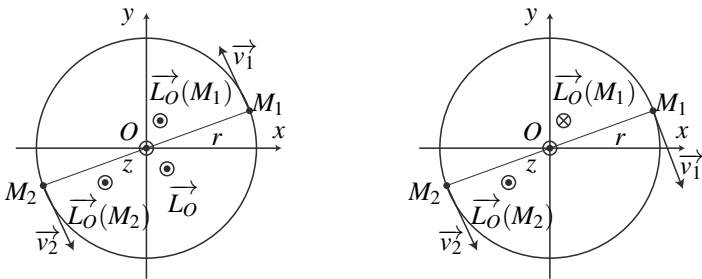


Figure 19.4 – Moment cinétique en O d’un système de deux points en révolution circulaire de centre O .

3.2 Cas d’un solide en rotation par rapport à un axe

On considère un solide en rotation à la vitesse angulaire $\dot{\theta}$ autour d’un axe orienté fixe dans un référentiel $\mathcal{R}$. On choisit l’axe (Oz) pour qu’il coïncide avec cet axe de rotation. On modélise le solide par un ensemble de points matériels M_i de masse m_i repérés en coordonnées cylindriques d’axe (Oz) : $M_i(r_i, \theta_i, z_i)$.

a) Moment d’inertie d’un solide

On rappelle une relation essentielle du chapitre de cinématique du solide : **chaque point d’un solide en rotation autour d’un axe fixe possède la même vitesse angulaire**. Le mouvement du solide est alors un cas particulier du mouvement d’un système de points dans lequel la vitesse angulaire de chacun des points du système est la même. On peut donc factoriser l’expression (19.3) par la vitesse angulaire commune $\dot{\theta}$:

$$L_{(Oz)} = \sum_i L_{(Oz)}(M_i) = \left(\sum_i J_{(Oz)}(M_i) \right) \dot{\theta} = J_{(Oz)} \dot{\theta}.$$

Le **moment d'inertie** $J_{(Oz)}$ du solide par rapport à l'axe (Oz) est défini par la somme des moments d'inertie par rapport à (Oz) de chacun des points le constituant :

$$J_{(Oz)} = \sum_i J_{(Oz)}(M_i).$$

Remarque

On peut également modéliser le solide par une répartition continue de masse. La somme discrète précédente devient alors une intégrale sur le volume du solide.

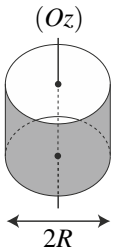
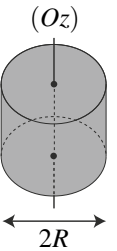
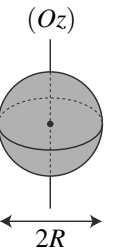
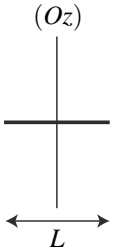
b) Moment cinétique d'un solide par rapport à un axe orienté (Oz)

Le moment cinétique par rapport à un axe (Oz) d'un solide en rotation autour de (Oz) à la vitesse angulaire $\dot{\theta}$ est égal au produit du moment d'inertie $J_{(Oz)}$ du solide par sa vitesse angulaire :

$$L_{(Oz)} = J_{(Oz)} \dot{\theta}.$$

c) Moments d'inertie de quelques solides homogènes

Le moment d'inertie d'un solide par rapport à un axe est une caractéristique intrinsèque que l'on peut mesurer. On peut également le calculer dans certains cas simples. Pour information, on donne les moments d'inertie par rapport à l'axe (Oz) dessiné sur les figures suivantes pour des solides homogènes de masse m :

cylindre vide de rayon R	cylindre plein de rayon R	boule de rayon R	barre de longueur L
mR^2	$\frac{1}{2}mR^2$	$\frac{2}{5}mR^2$	$\frac{1}{12}mL^2$
			

En pratique, en notant d la distance qui sépare l'axe (Oz) du point du solide qui en est le plus éloigné, le moment d'inertie d'un solide de masse m par rapport à (Oz) vaut $J_{(Oz)} = kmd^2$. Dans cette formule, k est un facteur numérique inférieur à 1 qui ne dépend que de la forme du solide et de la manière dont la masse est répartie à l'intérieur de ce dernier.

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

Exemple

On peut modéliser la jante d’une roue de voiture par un cylindre creux de masse m et de rayon R . Son moment d’inertie par rapport au moyeu est donc environ égal à mR^2 .

d) Évolution du moment d’inertie en fonction de la répartition des masses (MPSI)

La contribution d’une masse m au moment d’inertie d’un solide par rapport à un axe (Oz) est égale à mr^2 où r est sa distance à l’axe (Oz) .

Plus une masse m est éloignée de l’axe de rotation (Oz) , plus sa contribution au moment d’inertie par rapport à (Oz) est importante.

Ceci explique pourquoi le moment d’inertie d’un cylindre plein est inférieur à celui d’un cylindre creux de même masse : une partie importante de sa masse est située à faible distance de l’axe et contribue peu à son moment d’inertie.

La mesure du moment d’inertie d’un solide permet également d’obtenir des informations sur la répartition interne des masses.

Exemple

Les mesures astronomiques du moment d’inertie de la Terre par rapport à son axe Nord-Sud montrent qu’il vaut $0,33M_T R_T^2$ où M_T et R_T sont la masse et le rayon de la Terre. Il est inférieur à celui d’une boule homogène de même masse et même rayon qui vaut $0,4M_T R_T^2$. On en déduit que la répartition des masses à l’intérieur de la Terre n’est pas homogène et que la couche profonde située près de son axe de rotation est plus dense que les couches superficielles. Cette couche profonde est le noyau. Connaissant sa taille, on peut estimer sa densité. Elle correspond à celle du fer à haute pression. C’est un des principaux arguments prouvant que le noyau est essentiellement composé de fer.

4 Moment d’une force

On considère maintenant une force $\vec{f}$ qui s’applique en un point M .

4.1 Moment d’une force par rapport à un point O

Le moment en O de la force $\vec{f}$ s’appliquant en M est le vecteur défini par le produit vectoriel :

$$\vec{\mathcal{M}}_O(\vec{f}) = \vec{OM} \wedge \vec{f}.$$

Sa norme se mesure en joule $J = N \cdot m$.

Le moment des forces par rapport à un point est une grandeur additive : si on considère deux forces $\vec{f}_1$ et $\vec{f}_2$ qui s’appliquent sur M :

$$\vec{\mathcal{M}}_O(\vec{f}_1 + \vec{f}_2) = \vec{OM} \wedge (\vec{f}_1 + \vec{f}_2) = \vec{OM} \wedge \vec{f}_1 + \vec{OM} \wedge \vec{f}_2 = \vec{\mathcal{M}}_O(\vec{f}_1) + \vec{\mathcal{M}}_O(\vec{f}_2),$$

Le moment de la somme des forces est la somme des moments.

Remarque

Lorsque l'on s'occupe du mouvement d'un point M , la force s'applique forcément en M . Lorsque l'on s'occupe d'un système de points ou d'un solide, le point M est le point d'application de la force.

4.2 Moment d'une force par rapport à un axe orienté Δ

a) Définition

Le moment de la force $\vec{f}$ par rapport à l'axe orienté $\Delta = (O, \vec{u}_\Delta)$ est la projection orthogonale de $\vec{\mathcal{M}}_O(\vec{f})$ sur Δ :

$$\mathcal{M}_\Delta(\vec{f}) = \vec{\mathcal{M}}_O(\vec{f}) \cdot \vec{u}_\Delta = (\vec{OM} \wedge \vec{f}) \cdot \vec{u}_\Delta.$$

Comme le moment d'une force par rapport à un point, c'est une grandeur additive qui se mesure en joule. On peut également montrer que le moment de la force $\vec{f}$ par rapport à l'axe Δ ne dépend pas du point de l'axe choisi pour le calculer mais seulement de $\vec{f}$ et de la direction et de l'orientation de Δ .

b) Calcul en coordonnées cylindriques

Lorsque l'axe Δ est fixe, on peut faire coïncider l'axe (Oz) avec Δ et repérer M par ses coordonnées cylindriques : $\vec{OM} = r\vec{u}_r + z\vec{u}_z$. On peut alors calculer la composante selon $\vec{u}_z$ du moment $\vec{\mathcal{M}}_O$ qui correspond au moment de $\vec{f}$ par rapport à (Oz) :

$$\begin{aligned} \vec{\mathcal{M}}_O(\vec{f}) &= (r\vec{u}_r + z\vec{u}_z) \wedge (f_r\vec{u}_r + f_\theta\vec{u}_\theta + f_z\vec{u}_z) \\ &= -zf_\theta\vec{u}_r + (zf_r - rf_z)\vec{u}_\theta + rf_\theta\vec{u}_z \end{aligned}$$

d'où :

$$\mathcal{M}_{(Oz)}(\vec{f}) = rf_\theta.$$

(19.4)

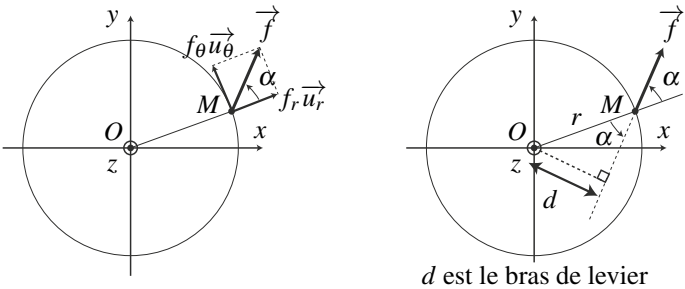


Figure 19.5 – Moment $\mathcal{M}_{(Oz)}(\vec{f})$ de la force $\vec{f}$.

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

Étant donné que ni la composante selon $\vec{u}_z$ du vecteur position ni celle de la force n'interviennent, on représente cette situation sur la figure 19.5 en se restreignant au plan (Oyx) , perpendiculaire à (Oz) . La partie gauche de la figure 19.5 permet de réécrire cette relation $\mathcal{M}_{(Oz)}(\vec{f}) = rf \sin \alpha$ où f est la norme de $\vec{f}$ et α l'angle $(\vec{u}_r, \vec{f})$ défini sur la figure 19.5. On se reporte alors à la partie droite pour observer que $|r \sin \alpha| = d$ est la distance séparant la droite d'action de la force $\vec{f}$ de l'axe (Oz) . La distance d est appelée le bras de levier de la force $\vec{f}$ et on a :

$$|\mathcal{M}_{(Oz)}(\vec{f})| = fd.$$

Le signe de $\mathcal{M}_{(Oz)}$ est :

- positif lorsque $\alpha \in]0, \pi[$. C'est le cas lorsque $\vec{f}$ tend à faire bouger M vers les θ croissants ;
- négatif lorsque $\alpha \in]-\pi, 0[$. C'est le cas lorsque $\vec{f}$ tend à faire bouger M vers les θ décroissants.



On appelle **droite d'action d'une force** $\vec{f}$ appliquée en M est la droite $(M, \vec{f})$ passant par M et dirigée par le vecteur $\vec{f}$.

c) Notion de bras de levier

On appelle **bras de levier** la distance séparant l'axe Δ de la droite d'action $(M, \vec{f})$ de la force $\vec{f}$. La valeur absolue $|\mathcal{M}_{\Delta}(\vec{f})|$ du moment de $\vec{f}$ par rapport à Δ est égale au produit :

$$|\mathcal{M}_{\Delta}(\vec{f})| = \text{norme de la force} \times \text{bras de levier}.$$

Pour déterminer le signe de $\mathcal{M}_{\Delta}(\vec{f})$, on peut au choix :

- rechercher le sens du projeté de $\vec{OM} \wedge \vec{f}$ sur la droite orientée Δ ;
- regarder le sens dans lequel la force $\vec{f}$ tend à faire tourner M autour de Δ :
 - si ce sens est direct, $\mathcal{M}_{\Delta}(\vec{f}) > 0$;
 - si ce sens est indirect, $\mathcal{M}_{\Delta}(\vec{f}) < 0$.

Remarque

Dans le cas de la figure 19.5, le moment de $\vec{f}$ par rapport à l'axe (Oz) orienté est positif car la force $\vec{f}$ tend à faire tourner M dans le sens direct autour de (Oz) .

d) Cas où $\mathcal{M}_{\Delta}(\vec{f})$ est nul

Le moment d'une force $\vec{f}$ non nulle par rapport à un axe orienté Δ est nul dans deux cas :

$$\mathcal{M}_{\Delta}(\vec{f}) = 0 \implies \begin{cases} \vec{f} \parallel \vec{u}_{\Delta} \\ \text{ou} \\ \text{la droite d'action de } \vec{f} \text{ coupe } \Delta. \end{cases}$$

En effet, dans le premier cas, le moment en O de $\vec{f}$ est perpendiculaire à $\vec{f}$ donc à $\vec{u}_{\Delta}$ et sa projection sur Δ est donc nulle. Dans le deuxième cas, le bras de levier est nul.

Remarque

En coordonnées cylindriques d'axe $(Oz) = \Delta$, l'équation (19.4) permet également de démontrer le premier cas puisque $\vec{f} \parallel \vec{u}_z$ donc $f_{\theta} = 0$.

5 Loi du moment cinétique pour un point matériel

On étudie le mouvement d'un point matériel M de masse m dans un référentiel galiléen $\mathcal{R}$. Le point M est soumis à un ensemble de forces $\vec{f}_i$. On note O un point fixe et Δ une droite orientée fixe contenant O . On choisit l'axe (Oz) de telle sorte que $\Delta = (Oz)$. À l'instant t , on note $\vec{OM}$, $\vec{v}$ et $\vec{p} = m\vec{v}$ les vecteurs position, vitesse et quantité de mouvement de M dans $\mathcal{R}$. On note également $\vec{L}_O$ le moment cinétique de M par rapport à O , $L_{(Oz)}$ son moment cinétique par rapport à $\Delta = (Oz)$, $\vec{\mathcal{M}}_O(\vec{f}_i)$ le moment de la force $\vec{f}_i$ par rapport à O et $\mathcal{M}_{(Oz)}(\vec{f}_i)$ son moment par rapport à $\Delta = (Oz)$.

5.1 Loi du moment cinétique par rapport à un point fixe

La dérivée temporelle du moment cinétique de M par rapport à O est égale à la somme des moments des forces calculés par rapport au même point O :

$$\frac{d\vec{L}_O}{dt} = \sum_i \vec{\mathcal{M}}_O(\vec{f}_i). \tag{19.5}$$

On démontre cette loi en remarquant que :

$$\begin{aligned} \frac{d\vec{L}_O}{dt} &= \frac{d\vec{OM} \wedge \vec{p}}{dt} = \vec{OM} \wedge \frac{d\vec{p}}{dt} + \frac{d\vec{OM}}{dt} \wedge \vec{p} \\ &= \vec{OM} \wedge \frac{d\vec{p}}{dt}. \end{aligned}$$

En effet, le point O étant fixe, $\frac{d\vec{OM}}{dt} = \vec{v}$ donc $\frac{d\vec{OM}}{dt} \wedge \vec{p} = \vec{v} \wedge (m\vec{v}) = \vec{0}$. On exprime

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

alors $\frac{d\vec{p}}{dt}$ à l'aide du principe fondamental de la dynamique : $\frac{d\vec{p}}{dt} = \sum_i \vec{f}_i$ pour obtenir :

$$\frac{d\vec{L}_O}{dt} = \vec{OM} \wedge \frac{d\vec{p}}{dt} = \vec{OM} \wedge \sum_i \vec{f}_i = \sum_i \vec{OM} \wedge \vec{f}_i = \sum_i \vec{\mathcal{M}}_O(\vec{f}_i).$$

5.2 Cas de conservation du moment cinétique

Cette loi est particulièrement intéressante lorsque la somme des moments des forces $\vec{f}_i$ par rapport à O s'annule. Cela correspond à $\sum_i \vec{OM} \wedge \vec{f}_i = \vec{OM} \wedge \sum_i \vec{f}_i = \vec{0}$. En pratique, cette situation n'arrive que dans deux cas :

- celui où $\sum_i \vec{f}_i = \vec{0}$ à tout instant. Le point M est alors isolé ou pseudo-isolé. Il est en mouvement rectiligne et uniforme ou immobile ;
- celui où $\sum_i \vec{f}_i \parallel \vec{OM}$ à tout instant. La droite d'action de la résultante des forces passe alors constamment par O . M est soumis à une **force centrale** de centre O . Ce type de force fait l'objet du chapitre *Mouvement dans un champ de force centrale. Champs newtoniens*.

5.3 Loi du moment cinétique par rapport à un axe fixe

La dérivée temporelle du moment cinétique de M par rapport à l'axe orienté fixe (Oz) est égale à la somme des moments des forces calculés par rapport à ce même axe :

$$\frac{dL_{(Oz)}}{dt} = \sum_i \mathcal{M}_{(Oz)}(\vec{f}_i). \tag{19.6}$$

On démontre cette loi en remarquant que : $\frac{dL_{(Oz)}}{dt} = \frac{d(\vec{L}_O \cdot \vec{u}_z)}{dt} = \frac{d\vec{L}_O}{dt} \cdot \vec{u}_z$.

En effet, l'axe (Oz) étant fixe, le vecteur $\vec{u}_z$ est constant. On exprime alors $\frac{d\vec{L}_O}{dt}$ grâce à la relation (19.5) :

$$\frac{dL_{(Oz)}}{dt} = \frac{d\vec{L}_O}{dt} \cdot \vec{u}_z = \sum_i \vec{\mathcal{M}}_O(\vec{f}_i) \cdot \vec{u}_z = \sum_i \mathcal{M}_{(Oz)}(\vec{f}_i).$$

Remarque

La loi (19.6) est une loi scalaire qui ne fournit qu'une seule équation. Utilisée seule, elle ne permet de résoudre que les problèmes à un degré de liberté. En pratique, le mouvement du point M est presque obligatoirement circulaire. Pour les problèmes à deux degrés de liberté, elle doit être complétée par une autre équation mécanique indépendante comme la loi de l'énergie cinétique.

6 Loi du moment cinétique pour un solide en rotation

On s'intéresse au mouvement d'un solide en rotation autour d'un axe orienté (Oz) fixe dans un référentiel galiléen $\mathcal{R}$. Son moment d'inertie par rapport à (Oz) est noté $J_{(Oz)}$. Son mouvement est caractérisé par sa vitesse angulaire $\dot{\theta}$. Son moment cinétique par rapport à (Oz) vaut $L_{(Oz)} = J_{(Oz)} \dot{\theta}$. Il est soumis aux forces extérieures $\vec{f}_i$ de moments $\mathcal{M}_{(Oz)}(\vec{f}_i)$.

6.1 Loi scalaire du moment cinétique pour un solide

Dans un référentiel $\mathcal{R}$ galiléen, la dérivée temporelle du moment cinétique du solide par rapport à son axe de rotation fixe (Oz) est égale à la somme des moments des forces extérieures par rapport à ce même axe :

$$\frac{dL_{(Oz)}}{dt} = \sum_i \mathcal{M}_{(Oz)}(\vec{f}_i). \tag{19.7}$$

Pour un solide en rotation autour de l'axe (Oz) , le moment d'inertie $J_{(Oz)}$ est constant, alors :

$$\frac{dL_{(Oz)}}{dt} = \frac{d(J_{(Oz)} \dot{\theta})}{dt} = J_{(Oz)} \ddot{\theta},$$

et on peut réécrire cette relation :

$$J_{(Oz)} \ddot{\theta} = \sum_i \mathcal{M}_{(Oz)}(\vec{f}_i).$$

On peut noter le parallèle complet entre cette équation, que l'on peut également écrire :

$$J_{(Oz)} \ddot{\theta} = \sum_i \vec{\mathcal{M}}_O(\vec{f}_i) \cdot \vec{u}_z,$$

et le principe fondamental de la dynamique écrit pour un solide en translation rectiligne sur l'axe (Ox) :

$$m\ddot{x} = \sum_i \vec{f}_i \cdot \vec{u}_x.$$

L'accélération linéaire $\ddot{x}$ est remplacée par l'accélération angulaire $\ddot{\theta}$, les forces projetées sur l'axe du mouvement par les moments des forces projetés sur l'axe de rotation et la masse inerte m par le moment d'inertie $J_{(Oz)}$. Cela permet de donner une interprétation physique du moment d'inertie du solide.

Le moment d'inertie d'un solide par rapport à l'axe (Oz) est sa caractéristique intrinsèque qui mesure son aptitude à s'opposer aux variations de vitesse de rotation autour de cet axe.

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

6.2 Cas de conservation du moment cinétique

Le moment cinétique d'un solide en rotation autour d'un axe fixe est conservé si la somme des moments des forces extérieures qui lui sont appliquées est nulle.

En effet, dans ce cas, la loi du moment cinétique appliquée au solide s'écrit :

$$\frac{dL_{(Oz)}}{dt} = 0 \implies L_{(Oz)} = J_{(Oz)}\dot{\theta} = \text{constante}.$$

Le moment d'inertie $J_{(Oz)}$ étant une constante, la vitesse de rotation du solide l'est également.

Cas particulier d'un solide en équilibre Un solide en rotation est à l'équilibre lorsque sa vitesse angulaire reste nulle à tout instant. On obtient donc :

$$\dot{\theta} = 0 \implies L_{(Oz)} = 0 \implies \frac{dL_{(Oz)}}{dt} = 0.$$

La loi du moment cinétique s'écrit alors :

$$\sum_i \mathcal{M}_{(Oz)}(\vec{f}_i) = 0.$$

C'est un cas particulier de la situation précédente pour laquelle non seulement la somme des moments des forces est nulle, mais la vitesse angulaire initiale également.

6.3 Couples

a) Couple de deux forces

Deux forces $\vec{f}_1$ et $\vec{f}_2$ opposées s'appliquant respectivement en A_1 et A_2 forment un couple de forces (figure 19.6). Leur résultante est nulle : $\vec{f}_1 + \vec{f}_2 = \vec{0}$.

Ces forces sont de norme f identique et s'appliquent sur des droites d'action parallèles. La distance d entre ces droites s'appelle le bras de levier du couple et le moment du couple de force par rapport à l'axe orienté (Oz) est égal au produit de la force par le bras de levier :

$$|\mathcal{M}_{(Oz)}| = fd_1 + fd_2 = fd.$$

Par abus de langage, étant donné que la somme des deux forces est nulle et que seul le moment de ces forces est non nul, on désigne souvent par couple, le moment du couple par rapport à (Oz) et on le note Γ .

Il faut remarquer que le moment du couple Γ ne dépend pas de la position de l'axe de rotation.

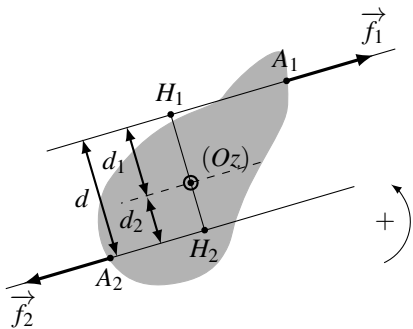


Figure 19.6 – Définition d'un couple de force.

ÉTAPES D'INSCRIPTION



pour voir comment créer un compte

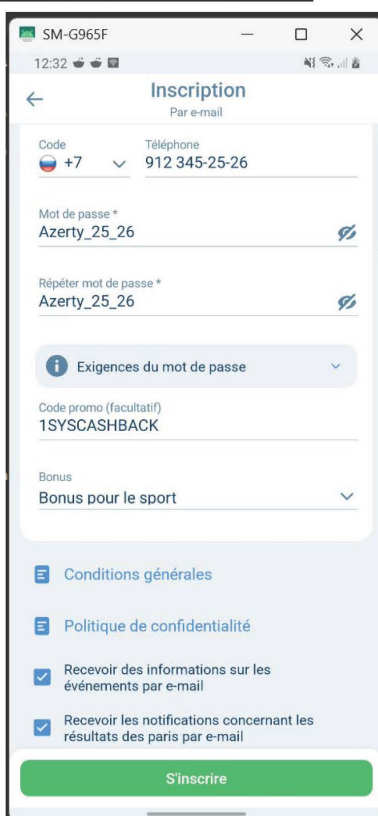
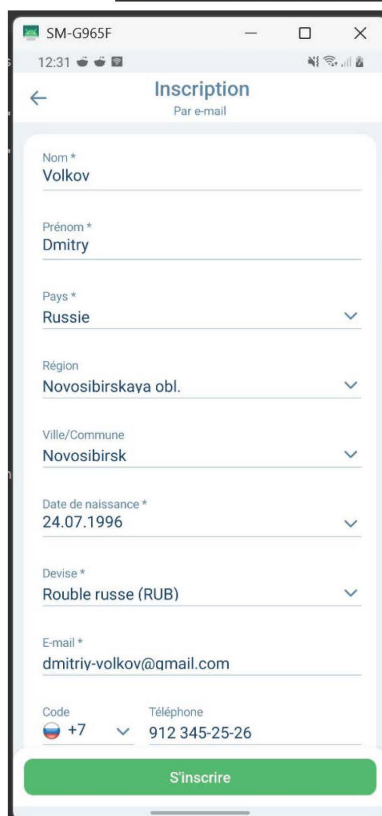
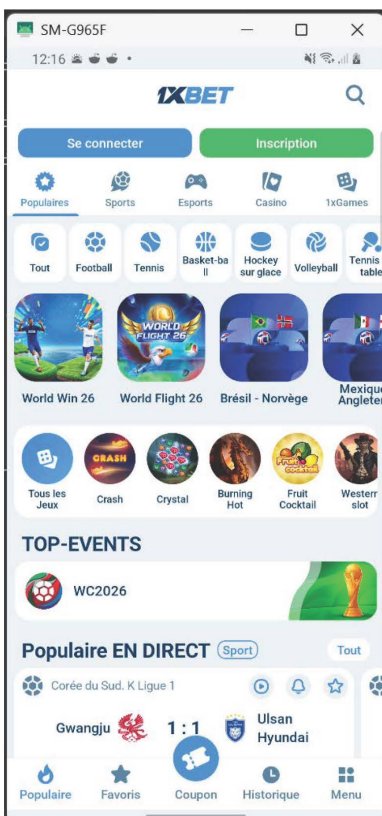
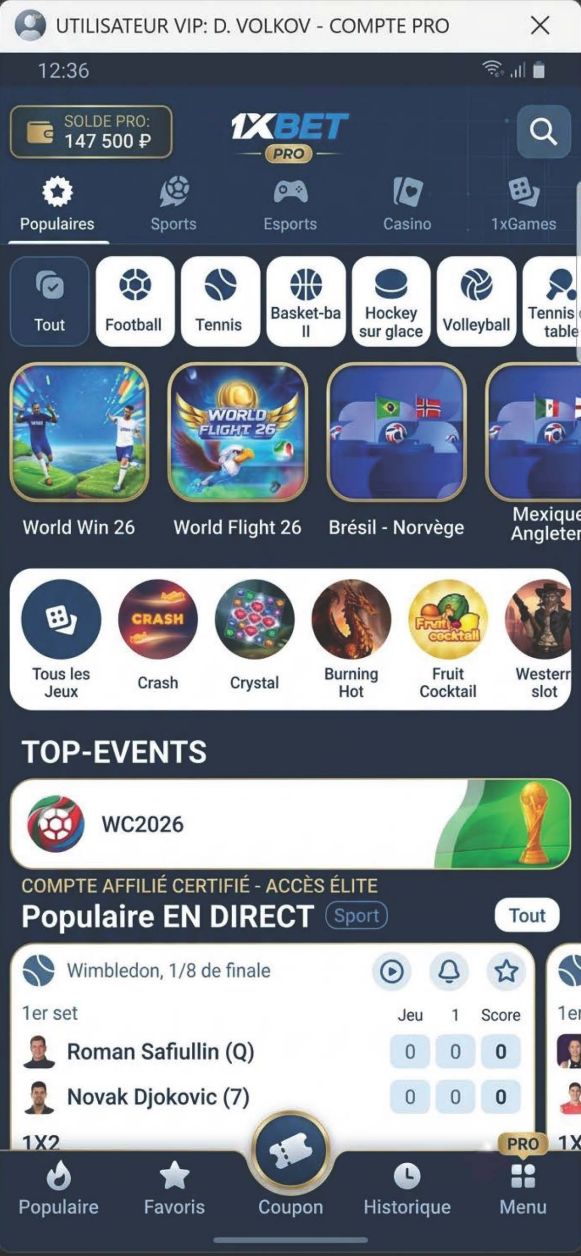
CLICK TO WATCH

1XBET



STREAMING EN DIRECT SUR MOBILE

Profitez du bonus!



Par ailleurs, il est algébrique et on peut trouver son signe en cherchant le sens dans lequel il tend à faire tourner M autour de l'axe orienté (Oz) :

- si ce sens est direct, $\Gamma > 0$;
- si ce sens est indirect, $\Gamma < 0$.

Exemple

Sur le schéma de la figure 19.6, $\Gamma = -fd$.

b) Généralisation

On peut généraliser la notion de couple à tous les cas où la somme des forces est nulle et le moment des forces par rapport à un axe (Oz) n'est pas nul, sans se préoccuper de savoir s'il a fallu deux forces pour réaliser cette situation.

Exemple

La résultante des actions mécaniques qu'un moteur exerce sur un arbre en rotation est un couple.

Remarque

En lien avec le cours de SII, un couple de force est un torseur des actions mécaniques de forme couple.

c) Couple moteur et couple de freinage

On considère un solide en rotation autour d'un axe orienté fixe (Oz) auquel on applique un couple de moment par rapport à (Oz) égal à Γ . La loi du moment cinétique par rapport à l'axe (Oz) appliquée au solide implique que :

$$\frac{dL_{(Oz)}}{dt} = J_{(Oz)} \ddot{\theta} = \Gamma.$$

On suppose que le solide tourne dans le sens direct autour de (Oz) ce qui implique que $\dot{\theta} > 0$ et on distingue deux cas selon le signe de Γ :

- si $\Gamma > 0$, $\ddot{\theta} > 0$ et la vitesse angulaire du solide en rotation augmente. La rotation du solide est accélérée ;
- si $\Gamma < 0$, $\ddot{\theta} < 0$ et la vitesse angulaire du solide en rotation diminue. La rotation du solide est freinée.

On peut tenir le même raisonnement dans le cas où $\dot{\theta} < 0$, et au final :

- lorsque Γ est du même signe que $\dot{\theta}$, la vitesse de rotation du solide augmente en valeur absolue. Le couple est un **couple moteur** ;
- lorsque Γ est du signe opposé à $\dot{\theta}$, la vitesse de rotation du solide diminue en valeur absolue. Le couple est un **couple de freinage**.

7 Application aux dispositifs rotatifs

Un dispositif rotatif est un dispositif dans laquelle un solide indéformable appelé **rotor** est en rotation autour d'un axe fixe par rapport à un solide immobile appelé **stator**. Dans cette

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

partie, on définit la liaison pivot qui permet de restreindre le mouvement relatif entre ces deux solides à un mouvement de rotation autour d'un axe fixe. On montre également qu'il est **nécessaire que le stator impose un couple au rotor** lorsque l'on veut réaliser un moteur ou un frein.

7.1 Liaison pivot d'axe (Oz)

a) Définition

Une **liaison pivot** d'axe (Oz) restreint les possibilités de mouvement du rotor à une rotation d'axe (Oz) par rapport au stator.

Dans ce chapitre, on suppose que la liaison pivot est géométriquement idéale. Dans ce cas, elle assure un guidage parfait de la rotation autour de l'axe de liaison (Oz) et elle bloque toute translation le long de (Oz). La liaison pivot est alors entièrement définie par la direction et la position de l'axe (Oz) que l'on précise systématiquement.

Remarque

C'est la liaison la plus commune dans les systèmes mécaniques. Dans un simple vélo, on en compte plus d'une dizaine puisqu'elle est nécessaire pour relier au cadre les éléments suivant : roues (2), guidon (1), pédalier (1), manettes de frein (2), mâchoires de frein (2). Il y en a également pour relier les pédales au pédalier (2) et encore une bonne demi-douzaine si le vélo possède des dérailleurs.

b) Réalisation pratique d'une liaison pivot

Pour réaliser une liaison pivot d'axe (Oz), la solution technique la plus courante consiste à emboîter deux cylindres de même axe et à réaliser des buttés pour empêcher les cylindres de coulisser le long de leur axe commun (figure 19.7).

Remarque

Le contact entre les deux cylindres conduit à l'existence de frottements que l'on peut réduire fortement en utilisant des roulements à billes ou à aiguilles.

c) Action de liaison et liaison pivot idéale d'axe (Oz)

L'action de liaison résulte des forces exercées par le stator sur le rotor. Elle n'est pas déterminée *a priori*. Si l'on regarde de près les forces de contact au niveau des cylindres emboîtés (coupe (AB) de la figure 19.7), on constate que ces forces sont réparties tout au long de la surface de contact. Si l'on peut négliger les frottements, elles sont normales aux surfaces de contact et coupent l'axe de rotation (Oz). Dans ce cas, le moment par rapport à (Oz) de chacune de ces forces est nul et le moment par rapport à (Oz) de la liaison est égal à 0.

L'action de liaison d'une liaison pivot idéale d'axe (Oz) a un moment par rapport à l'axe (Oz) égal à 0 :

$$\mathcal{M}_{(Oz)}(\text{liaison}) = 0.$$

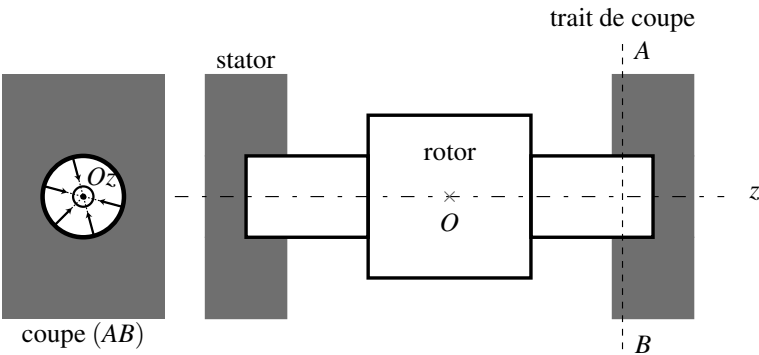


Figure 19.7 – Schéma de principe d'une liaison pivot.



Il faut bien comprendre que seul le moment de l'action de liaison est nul. Sa résultante ne l'est pas puisque c'est elle qui assure le guidage en rotation autour de l'axe (Oz).

Remarque

En pratique, malgré l'utilisation de roulements à bille ou à aiguille, il reste des forces de frottement qui produisent un couple résistant d'axe (Oz).

7.2 Notions sur les moteurs et les freins dans les dispositifs rotatifs (PTSI)

a) Cas du moteur

Un moteur est utilisé pour entraîner une autre pièce mécanique en rotation. Il exerce un couple moteur sur cette pièce et, d'après le principe des actions réciproques, cette dernière exerce un couple de freinage sur le rotor. Pour maintenir le rotor en rotation, il est donc nécessaire de compenser ce couple de freinage par un couple moteur. C'est le stator qui est chargé de produire ce couple.

b) Cas du frein

Un frein est utilisé pour freiner le rotor. Pour cela, il est indispensable de lui appliquer un couple de freinage. C'est le rôle joué dans ce cas par le stator.

Exemple

Lors du freinage d'un vélo, le rotor est la roue du vélo et le stator son cadre. Le freinage est assuré par des patins en caoutchouc qui viennent en appui sur la jante et produisent un couple de freinage sur le rotor par frottements solides.

Au final, lorsque l'on veut appliquer un couple sur le rotor, la présence du stator est indispensable pour deux raisons :

- la première, la plus évidente, est de créer un axe de rotation fixe ;
- la deuxième est de pouvoir appliquer un couple sur le rotor.

8 Pendule pesant

8.1 Position du problème et équation du mouvement

Un **pendule pesant** est un solide de masse m de forme quelconque mobile dans le champ de pesanteur terrestre autour d'un axe horizontal fixe ne passant pas par son centre de gravité G (figure 19.8).

On note (Oz) l'axe de rotation du solide, G son centre de gravité et $J_{(Oz)}$ son moment d'inertie par rapport à l'axe (Oz) . On repère la position du solide par l'angle θ que fait la droite (OG) avec la verticale descendante (Ox) . On suppose que la liaison entre le solide et le référentiel terrestre est une liaison pivot parfaite d'axe (Oz) . Le solide est soumis à :

- l'action exercée par la liaison pivot. On suppose cette liaison pivot idéale, ce qui implique que son moment par rapport à l'axe (Oz) est nul ;
- son poids vertical descendant qui s'applique au centre de gravité G . Son moment par rapport à (Oz) est égal en module au produit $mg \times L$ où $L = d \sin \theta$ est le bras de levier représenté sur la figure 19.8. Cette force tend à ramener le pendule vers sa position d'équilibre. Le signe de son moment par rapport à (Oz) est opposé à celui de $\sin \theta$. En effet, on voit sur la figure 19.8 que $\theta > 0$, $\sin \theta > 0$, et $\mathcal{M}_{(Oz)} < 0$. Au final :

$$\mathcal{M}_{(Oz)} = -mgd \sin \theta.$$

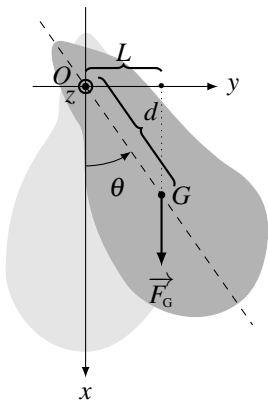


Figure 19.8 – Pendule pesant.

On applique au solide la loi du moment cinétique par rapport à l'axe orienté (Oz) fixe dans le référentiel terrestre $\mathcal{R}$ supposé galiléen :

$$\frac{dL_{(Oz)}}{dt} = -mgd \sin \theta. \quad (19.8)$$

Le moment cinétique de la barre est égal $L_{(Oz)} = J_{(Oz)} \dot{\theta}$ donc $\frac{dL_{(Oz)}}{dt} = J_{(Oz)} \ddot{\theta}$. On injecte alors cette relation dans l'équation (19.8) pour obtenir :

$$J_{(Oz)} \ddot{\theta} = -mgd \sin \theta \quad \Longleftrightarrow \quad \ddot{\theta} + \frac{mgd}{J_{(Oz)}} \sin \theta = 0. \quad (19.9)$$

Cette équation est du type $\ddot{\theta} + \omega_0^2 \sin \theta = 0$ avec $\omega_0 = \sqrt{\frac{mgd}{J_{(Oz)}}}$. Il s'agit de la même équation que celle que l'on a obtenue lors de l'étude du pendule simple.

Les positions d'équilibre $\theta_{eq1} = 0$ et $\theta_{eq2} = \pi$ permettent au moment du poids par rapport à (Oz) d'être égal à 0. Dans ces positions, la somme des moments des forces qui s'appliquent au pendule est nulle. On remarque que dans ce cas, le centre de gravité est sur une droite

verticale partant du point d'attache O . En pratique, seule la position $\theta_{eq_1} = 0$ est réalisable car c'est la seule qui est stable.

Remarque

On met en évidence une méthode qui permet de repérer expérimentalement la position du centre de gravité d'un solide. Si l'on utilise un autre point d'attache, on trouve une autre droite verticale. Le centre de gravité est alors à l'intersection des deux droites et on peut déterminer la distance d .

8.2 Oscillations de faible amplitude

Lorsque les oscillations sont de faible amplitude au voisinage de la position d'équilibre $\theta_{eq} = 0$, $\sin \theta \simeq \theta$ et l'équation (19.9) peut être linéarisée :

$$\ddot{\theta} + \omega_0^2 \theta = 0 \quad \text{où} \quad \omega_0 = \sqrt{\frac{mgd}{J_{(Oz)}}}.$$

On reconnaît une équation d'oscillateur harmonique dont les solutions sont des sinusoides de pulsation ω_0 :

$$\theta(t) = \theta_m \cos(\omega_0 t + \varphi_0).$$

θ_m et φ_0 sont des constantes d'intégration que l'on détermine à partir des conditions initiales. Étant donné que le solide oscille autour d'une position d'équilibre stable, il est normal de trouver des oscillations harmoniques pour les petites oscillations autour de cette position. Ces oscillations présentent la propriété remarquable d'avoir une période indépendante de l'amplitude (isochronisme des petites oscillations). La mesure de cette période permet de déterminer la valeur de $J_{(Oz)}$, connaissant la position du centre de gravité, la masse du solide et la valeur de l'accélération de pesanteur.

8.3 Intégrale première du mouvement et étude qualitative

a) Obtention d'une intégrale première du mouvement

On peut établir une intégrale première du mouvement à partir de l'équation (19.9). Pour cela, on la multiplie par $\dot{\theta}$:

$$J_{(Oz)} \ddot{\theta} \dot{\theta} + mgd \sin \theta \dot{\theta} = J_{(Oz)} \frac{d}{dt} \left(\frac{\dot{\theta}^2}{2} \right) - mgd \frac{d}{dt} (\cos \theta) = 0,$$

puis on l'intègre par rapport au temps :

$$\frac{1}{2} J_{(Oz)} \dot{\theta}^2 - mgd \cos \theta = \text{constante} = E_m.$$

La quantité $\frac{1}{2} J_{(Oz)} \dot{\theta}^2 - mgd \cos \theta = E_m$ est donc une intégrale première du mouvement homogène à une énergie. En effet, $J_{(Oz)}$ se mesure en $\text{kg} \cdot \text{m}^2$ et $\dot{\theta}$ est homogène à l'inverse d'un temps donc $J_{(Oz)} \dot{\theta}^2$ se mesure en $\text{kg} \cdot \text{m}^2 \cdot \text{s}^{-2} = \text{J}$.

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

Le terme $\frac{1}{2}J_{(Oz)}\dot{\theta}^2$ correspond à l'énergie cinétique de la barre en rotation autour de l'axe (Oz) . Le terme $-mgd \cos \theta$ est l'énergie potentielle de pesanteur du solide car $-d \cos \theta$ est l'altitude du centre de gravité G comptée à partir de O (voir figure 19.8). E_m est donc l'énergie mécanique du solide. C 'est une constante qui dépend des conditions initiales du mouvement.

b) Étude qualitative du mouvement

Pour réaliser une étude qualitative du mouvement, on utilise la démarche décrite dans le chapitre *Mouvement dans un puits de potentiel* et on l'adapte au cas d'une variable angulaire.

On trace l'énergie potentielle sans dimension $E_p^* = \frac{E_p}{mgd} = -\cos \theta$ en fonction de θ . E_p^* est obtenue à partir de l'énergie potentielle $E_p = -mgd \cos \theta$ en prenant l'énergie caractéristique $E_0 = mgd$ comme échelle d'énergie.

L'énergie potentielle est minimale lorsque θ est égal à 0 modulo 2π pour lesquels le centre de gravité G du pendule pesant est situé en dessous de O et à sa verticale. Ces minima correspondent à la position d'équilibre stable du pendule pesant. L'énergie potentielle est maximale lorsque θ est un multiple impair de π . Ces maxima correspondent à la position d'équilibre instable du pendule pour laquelle G est à la verticale de O et au dessus.

On trace la droite horizontale d'ordonnée $E_m^* = \frac{E_m}{mgd}$ qui doit être supérieure à E_p^* et doit donc se situer au dessus de la zone grisée.

Deux cas se présentent alors :

- si $-1 < E_m^* = E_{m1}^* < 1$, le pendule pesant est dans un état lié. Il n'a pas l'énergie mécanique suffisante pour sortir du puits de potentiel et on observe un mouvement pendulaire dont on peut déterminer l'amplitude graphiquement ;
- soit $E_m^* = E_{m2}^* > 1$, le pendule pesant est dans un état de diffusion avec une énergie suffisante pour sortir du puits de potentiel. Les mouvements sont alors des mouvements de révolution.

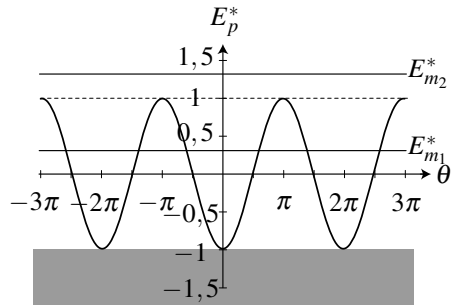


Figure 19.9 – Énergie potentielle d'un pendule pesant. En pointillé, l'énergie mécanique minimale nécessaire pour obtenir un mouvement de révolution.

8.4 Portrait de phase

On peut utiliser l'intégrale première du mouvement pour expliciter la vitesse angulaire $\dot{\theta}$ en fonction de l'angle θ et de l'énergie mécanique E_m . On trouve :

$$\frac{\dot{\theta}}{\omega_0} = \pm \sqrt{2 \left(\frac{E_m}{mgd} + \cos \theta \right)}.$$

On trace alors $\frac{\dot{\theta}}{\omega_0}$ en fonction de θ pour différentes valeurs de $\frac{E_m}{mgd}$ pour obtenir le portrait de phase du pendule pesant dans le plan sans dimension $\left(\theta, \frac{\dot{\theta}}{\omega_0}\right)$. On rappelle que l'énergie cinétique étant positive, l'énergie mécanique ne peut jamais devenir inférieure au minimum d'énergie potentielle, à savoir $-mgd$, ce qui implique que $\frac{E_m}{mgd} > -1$. Ce diagramme de phase est représenté sur la figure 19.10.

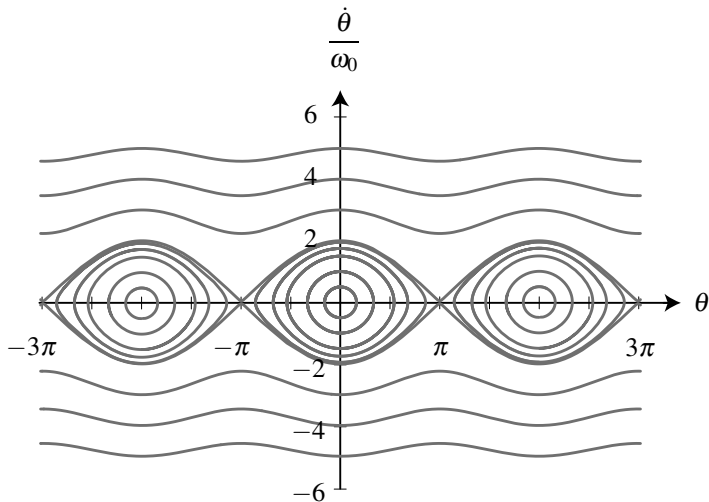


Figure 19.10 – Portrait de phase d'un pendule pesant.

Les mouvements de faible amplitude sont pendulaires : le pendule réalise un mouvement de va et vient entre deux positions extrêmes $\pm\theta_m$. En ces positions extrêmes, l'énergie cinétique s'annule ce qui implique que :

$$E_p(\pm\theta_m) = E_m \quad \text{où } \theta_m \in [0, \pi[.$$

L'angle θ oscille périodiquement dans l'intervalle $]-\theta_m, \theta_m[$. Le portrait de phase est similaire à celui d'un point matériel dans un puits de potentiel quelconque :

- lorsque l'amplitude est faible, les trajectoires de phase sont elliptiques (et même circulaires avec des échelles des abscisses et des ordonnées adaptées), ce qui est la signature des oscillateurs harmoniques ;
- lorsque l'amplitude devient importante, la trajectoire de phase se déforme mais reste fermée.

Pour des énergies plus importantes, le mouvement est révolutif : le pendule effectue une succession de tours complets toujours dans le même sens. L'angle θ n'est plus borné (si l'on néglige les frottements) et varie de 2π à chaque tour. La vitesse est maximale lorsque le centre de gravité du pendule atteint son point le plus bas et minimale au point le plus haut.

La transition entre ces deux types de mouvement a lieu lorsque l'énergie mécanique devient supérieure à l'énergie potentielle maximale, ce qui permet au pendule de s'échapper

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

du puits de potentiel et d’être dans un état de diffusion. On en déduit que :

$$E_{m,transition} = mgd.$$

8.5 Résolution numérique

On peut également résoudre numériquement l’équation (19.9) et trouver l’évolution temporelle de l’angle θ . Pour cela, il faut choisir une position et une vitesse angulaire initiale à procurer au pendule. La figure 19.11 représente les solutions de l’équation (19.9) pour les conditions initiales $\theta(0) = 0$ et $\dot{\theta}(0) = \dot{\theta}_0$, où la vitesse angulaire initiale $\dot{\theta}_0$ est un paramètre ajustable. Sur cette figure, $\frac{\dot{\theta}_0}{\omega_0}$ prend les valeurs : 0, 1 ; 0, 5 ; 1, 0 ; 1, 5 ; 1, 75 ; 1, 9 ; 1, 99 ; 2, 0 et 2, 01.

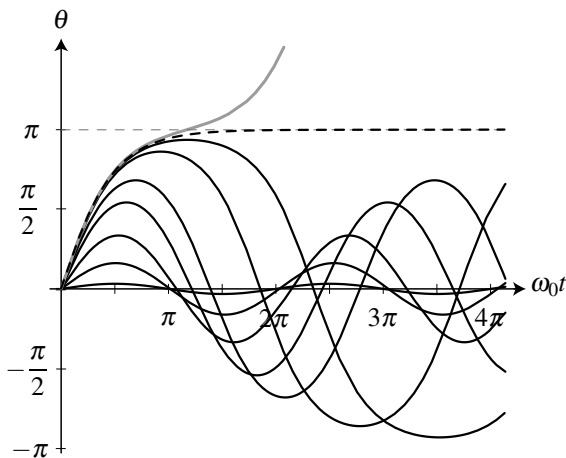


Figure 19.11 – Évolution temporelle de l’angle θ : en trait continu noir, les mouvements sont pendulaires ; en trait continu gris, il est révolutif. Les pointillés correspondent à la transition entre ces deux régimes.

Le mouvement révolutif correspond à une fonction $\theta(t)$ monotone croissante (en trait continu gris obtenue pour $\dot{\theta}_0 = 2,01\omega_0$), tandis que les mouvements pendulaires correspondent à des fonctions $\theta(t)$ périodiques (en trait continu noir obtenues pour $\dot{\theta}_0$ compris entre $0,1\omega_0$ et $1,99\omega_0$). La courbe obtenue pour $\dot{\theta}_0 = 2,0\omega_0$ est la courbe en pointillés qui fait la transition entre ces deux comportements très différents. Les mouvements pendulaires ayant lieu dans un puits de potentiel qui n’est pas harmonique, seules les oscillations de faible amplitude sont sinusoïdales et isochrones.

9 Énergie d’un solide en rotation autour d’un axe fixe

On reprend les mêmes notations que dans le paragraphe 6 : on considère un solide en rotation autour d’un axe orienté fixe noté (Oz) à la vitesse angulaire $\dot{\theta}$ par rapport à un référentiel $\mathcal{R}$ galiléen. Dans ce paragraphe, on établit l’expression de l’énergie cinétique du solide et la loi de l’énergie cinétique pour un solide en rotation.

9.1 Énergie cinétique d'un solide en rotation

On modélise le solide par un ensemble de points matériels M_i de masse m_i repérés en coordonnées cylindriques d'axe (Oz) : $M_i(r_i, \theta_i, z_i)$. On a vu au paragraphe 3.2 que le moment d'inertie du solide par rapport à (Oz) vaut alors $J_{(Oz)} = \sum_i m_i r_i^2$. Par ailleurs, on a vu dans le chapitre *Cinématique du solide* qu'un point M_i quelconque du solide est en mouvement circulaire de rayon r_i à la vitesse angulaire commune $\dot{\theta}$. Sa vitesse est donc $\vec{v}_i = r_i \dot{\theta} \vec{u}_{\theta_i}$ et son énergie cinétique $E_c(M_i) = \frac{1}{2} m_i v_i^2 = \frac{1}{2} m_i r_i^2 \dot{\theta}^2$.

L'énergie cinétique du solide est obtenue par sommation de l'énergie cinétique de chacun des points qui le constituent :

$$E_c = \sum_i E_c(M_i) = \sum_i \frac{1}{2} m_i r_i^2 \dot{\theta}^2 = \frac{1}{2} \left(\sum_i m_i r_i^2 \right) \dot{\theta}^2,$$

où on reconnaît l'expression du moment d'inertie du solide : $J_{(Oz)} = \sum_i m_i r_i^2$.

Un solide de moment d'inertie $J_{(Oz)}$ en rotation autour d'un axe fixe (Oz) à la vitesse angulaire $\dot{\theta}$ possède l'énergie cinétique :

$$E_c = \frac{1}{2} J_{(Oz)} \dot{\theta}^2.$$

C'est bien l'énergie cinétique que l'on a trouvée au paragraphe 8 lors de la détermination de l'intégrale première du mouvement. Il s'agissait alors d'un cas particulier de la loi de l'énergie cinétique pour un solide en rotation que l'on va établir dans le paragraphe suivant.

Remarque

On peut mémoriser l'expression de l'énergie cinétique d'un solide en rotation autour de l'axe (Oz) à partir de celle, bien connue, d'un solide en translation rectiligne sur l'axe (Ox) : $E_c = \frac{1}{2} m \dot{x}^2$. Il suffit de remplacer la masse par le moment d'inertie et la vitesse linéaire par la vitesse angulaire.

9.2 Puissance d'une force appliquée sur un solide en rotation

On considère une force $\vec{f}_i$ qui s'applique au point M_i d'un solide en rotation autour de l'axe (Oz) à la vitesse angulaire $\dot{\theta}$. La puissance de la force $\vec{f}_i$ est égale au produit scalaire de la force $\vec{f}_i$ par la vitesse $\vec{v}_{M_i} = r_i \dot{\theta} \vec{u}_{\theta_i}$ du point M_i sur laquelle elle s'applique d'où :

$$\mathcal{P}(\vec{f}_i) = \vec{f}_i \cdot \vec{v}_{M_i} = \vec{f}_i \cdot r_i \dot{\theta} \vec{u}_{\theta_i} = f_{i\theta} r_i \dot{\theta} = \mathcal{M}_{(Oz)}(\vec{f}_i) \dot{\theta}.$$

En effet, d'après l'équation 19.4, $r_i f_{i\theta} = \mathcal{M}_{(Oz)}$.

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

La puissance de la force $\vec{f}_i$ appliquée en un point M_i d'un solide en rotation autour d'un axe fixe (Oz), est égale au produit du moment par rapport à (Oz) de cette force par la vitesse angulaire $\dot{\theta}$ de rotation du solide autour de cet axe :

$$\mathcal{P}(\vec{f}_i) = \mathcal{M}_{(Oz)}(\vec{f}_i) \dot{\theta}.$$

9.3 Loi de l'énergie cinétique pour un solide indéformable

Dans le référentiel $\mathcal{R}$ galiléen, la dérivée temporelle de l'énergie cinétique d'un solide indéformable en rotation autour d'un axe fixe, est égale à la puissance de l'ensemble des forces extérieures $\vec{f}_i$ qu'on lui applique :

$$\frac{dE_c}{dt} = \sum_i \mathcal{P}(\vec{f}_i).$$

Pour établir cette loi, on dérive l'expression de l'énergie cinétique et on trouve :

$$\frac{dE_c}{dt} = \frac{d}{dt} \left(\frac{1}{2} J_{(Oz)} \dot{\theta}^2 \right) = J_{(Oz)} \ddot{\theta} \dot{\theta}. \quad (19.10)$$

Or, la loi du moment cinétique pour un solide en rotation autour de l'axe (Oz) fixe implique : $J_{(Oz)} \ddot{\theta} = \sum_i \mathcal{M}_{(Oz)}(\vec{f}_i)$ où $\mathcal{M}_{(Oz)}(\vec{f}_i)$ est le moment par rapport à (Oz) de la force extérieure $\vec{f}_i$ appliquée au solide. En insérant cette relation dans l'équation (19.10), on obtient :

$$\frac{dE_c}{dt} = \sum_i \mathcal{M}_{(Oz)}(\vec{f}_i) \dot{\theta} = \sum_i \mathcal{P}(\vec{f}_i),$$

où $\mathcal{P}(\vec{f}_i) = \mathcal{M}_{(Oz)}(\vec{f}_i) \dot{\theta}$ est la puissance de la force $\vec{f}_i$.

La démonstration de la loi de l'énergie cinétique montre qu'elle est équivalente à la loi du moment cinétique autour d'un axe fixe.

Pour étudier un solide en rotation autour d'un axe fixe, les lois de l'énergie cinétique ou du moment cinétique sont équivalentes et donnent des équations identiques.

Par ailleurs, la définition de la puissance confirme la cohérence des notions de couples moteurs et de freinage vues au paragraphe 6.3. En effet, on retrouve que :

- lorsque $\mathcal{M}_{(Oz)}$ est du même signe que $\dot{\theta}$, la puissance de la force est positive, le couple est moteur ;
- lorsque $\mathcal{M}_{(Oz)}$ est du signe opposé à $\dot{\theta}$, la puissance de la force est négative, le couple est résistant.

Remarque

On peut mémoriser l'expression de la puissance d'une force $\vec{f}$ s'exerçant sur un solide en rotation autour de l'axe (Oz) à partir de celle, bien connue, qui s'applique sur un solide en translation rectiligne sur l'axe (Ox) :

$$\mathcal{P} = f_x \dot{x}.$$

Il suffit de remplacer la projection f_x de la force sur l'axe du mouvement (Ox) par la projection $\mathcal{M}_{(Oz)}$ du moment sur l'axe de rotation (Oz) et la vitesse linéaire $\dot{x}$ par la vitesse angulaire $\dot{\theta}$ pour obtenir :

$$\mathcal{P} = \mathcal{M}_{(Oz)} \dot{\theta}.$$

SYNTHÈSE

SAVOIRS

- moment cinétique d'un point matériel par rapport à un point ou à un axe
- moment cinétique d'un système discret de points par rapport à un axe
- moment d'inertie et lien qualitatif à la répartition des masses
- moment cinétique d'un solide par rapport à un axe orienté
- notion de bras de levier
- définition d'un couple
- définition d'une liaison pivot
- loi du moment cinétique en un point fixe
- loi scalaire du moment cinétique pour un solide en rotation autour d'un axe fixe
- pendule pesant
- énergie cinétique d'un solide en rotation

SAVOIR-FAIRE

- relier la direction et le sens du vecteur moment cinétique aux caractéristiques du mouvement
- maîtriser le caractère algébrique du moment cinétique scalaire
- utiliser le bras de levier pour calculer le moment d'une force
- justifier le moment que peut fournir une liaison pivot
- reconnaître les cas de conservation du moment cinétique
- étudier le mouvement d'un pendule pesant. Lire et interpréter son portrait de phase
- établir l'équivalence entre la loi de l'énergie cinétique et la loi scalaire du moment cinétique pour un solide en rotation autour d'un axe fixe

MOTS-CLÉS

- | | | |
|---------------------|-----------------------|-------------------|
| • rotation. | • moment d'une force. | • pendule pesant. |
| • moment cinétique. | • moment d'inertie. | |

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

S'ENTRAÎNER

19.1 Ordres de grandeur des moments cinétiques (★)

1. Le moment d'inertie de la Terre en rotation uniforme autour de l'axe passant par ses pôles vaut $J = 0,33M_T R_T^2$ avec $M_T = 6,0.10^{24}$ kg et $R_T = 6,4.10^3$ km. Calculer le moment d'inertie de la Terre et son moment cinétique par rapport à l'axe de ses pôles.
2. Dans le modèle de Bohr, le mouvement de l'électron autour du noyau est assimilé à un mouvement circulaire et uniforme de centre O confondu avec le noyau. La trajectoire de rayon $r_0 = 53$ pm est parcourue à la fréquence $f = 6,6.10^{15}$ Hz. Calculer le moment cinétique de l'électron. On rappelle sa masse vaut $m_e = 9,1.10^{-31}$ kg.
3. Un tambour de machine à laver de rayon $R = 25$ cm et de masse $m = 5$ kg tourne à la vitesse angulaire de $1000 \text{ tr} \cdot \text{min}^{-1}$. Calculer son moment cinétique par rapport à son axe de rotation sachant que son moment d'inertie par rapport à cet axe vaut $J = mR^2$.

19.2 Etude d'une poulie (★)

Une masse de $m = 5,0$ kg est suspendue à l'extrémité d'une corde enroulée sur une poulie de masse $m_p = 1,0$ kg et de rayon $R = 10$ cm en liaison pivot idéale autour de son axe avec un support fixe (figure 19.12). On prendra $g = 10 \text{ m} \cdot \text{s}^{-2}$. Le moment d'inertie de la poulie par rapport à son axe vaut $I = \frac{1}{2}m_p R^2$.

A - Aspect cinématique : On suppose que la poulie est en rotation uniforme autour de son axe fixe (Oz) à la vitesse angulaire θ .

1. Quel est la vitesse de la masse m ?

B - Aspect statique : Cette même poulie est retenu par un opérateur.

2. Quelle force l'opérateur doit-il exercer sur la poulie pour l'empêcher de tourner ?

C - Aspect dynamique : Avec le même dispositif, l'opérateur lâche la poulie.

3. Déterminer l'accélération angulaire du cylindre, l'accélération linéaire de la masse m et la tension de la corde.

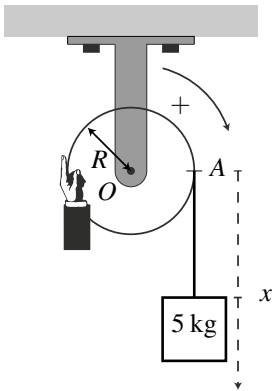
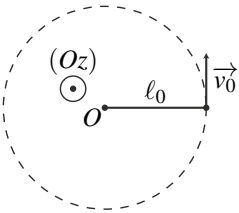


Figure 19.12

19.3 Mouvement d'une sphère attachée au bout d'un fil
(★)

Une sphère de petite taille et de masse $m = 0,10 \text{ kg}$ est attachée à l'extrémité d'un fil sans masse de longueur $\ell_0 = 1,0 \text{ m}$ dont l'autre extrémité est fixée en O . Elle se déplace sur un cercle horizontal de rayon ℓ_0 . Sa vitesse est $v_0 = 1,0 \text{ m}\cdot\text{s}^{-1}$.



1. Déterminer son moment cinétique par rapport à O puis par rapport à (Oz) .
2. On réduit brutalement la longueur du fil à $\ell_1 = 0,50 \text{ m}$. Que devient la vitesse de la sphère ?
3. Comparer l'énergie cinétique avant et après la réduction de la la longueur du fil.
4. Quelle force provoque l'augmentation de l'énergie cinétique de la sphère ? Commenter.

19.4 Chute d'un arbre (★★)

On assimile un arbre à une tige longue et homogène de longueur L et de masse m . On le scie à sa base et l'arbre bascule en tournant autour de son point d'appui au sol. On suppose que le point d'appui reste fixe et ne glisse pas et on repère la position de l'arbre par l'angle θ qu'il fait avec la verticale. A $t = 0$, l'arbre fait un angle $\theta_0 = 5^\circ$ avec la verticale et est immobile. On donne le moment d'inertie par rapport à son extrémité $I = \frac{1}{3}mL^2$.

1. Etablir l'équation du mouvement de chute de l'arbre.
2. Montrer que, lorsque l'arbre fait un angle θ avec la verticale, sa vitesse angulaire vaut :

$$\dot{\theta} = \sqrt{\frac{3g}{L} (\cos \theta_0 - \cos \theta)}.$$

3. Montrer que cette relation peut-être réécrite : $\sqrt{\frac{3g}{L}} dt = \frac{d\theta}{\sqrt{\cos \theta_0 - \cos \theta}}$.
4. Déterminer le temps de chute d'un arbre de 30 m . On prendra $g = 10 \text{ m}\cdot\text{s}^{-2}$. On donne pour $\theta_0 = 5^\circ$: $\int_{\theta_0}^{\frac{\pi}{2}} \frac{d\theta}{\sqrt{\cos \theta_0 - \cos \theta}} = 5,1$.

APPROFONDIR

19.5 Le swing de golf (extrait de Centrale PC 2012) (★★★)

L'objectif du golf est de parvenir, en un nombre de coups le plus faible possible, à envoyer la balle dans chacun des 18 trous du parcours, en la frappant à l'aide d'un instrument appelé « club » : le geste effectué par le joueur avec le club pour frapper la balle est appelé « swing ». On étudie dans cette partie les caractéristiques mécaniques du club de golf.

On notera $\dot{X}$ et $\ddot{X}$ les dérivées temporelles première et seconde d'une fonction $X(t)$ quelconque.

Un club de golf est schématiquement composé de deux parties, rigidement liées entre elles : le manche et la tête de club. Le joueur tient le club avec ses mains par l'extrémité A du manche, et la tête de club est fixée à l'autre extrémité B et entre en contact avec la balle lors de l'impact (voir 19.13). Attention, dans tout l'énoncé, le mot « club » représentera l'ensemble {manche et tête de club}.

Le manche est une tige rectiligne sans épaisseur de longueur $AB = L_c$ et le club possède un centre de masse G_c qu'on considérera situé sur le manche, avec $AG_c = h_c$, une masse totale m_c et un moment d'inertie total J_c par rapport à un axe perpendiculaire passant par A . Il faut bien noter que la tête de club est prise en compte dans G_c , m_c et J_c .

On cherche à déterminer expérimentalement $AG_c = h_c$ et J_c .

1. Expliquer brièvement, en s'appuyant sur un ou deux schéma(s) simple(s) et clair(s), pourquoi, en tendant son index à l'horizontale et en s'arrangeant pour poser le club en équilibre dessus à l'horizontale, on détermine la position de G_c .

2. Afin de mesurer J_c , on réalise à l'aide du club un pendule pesant, en suspendant l'extrémité supérieure du manche (point A) à un axe horizontal (Az) fixe, par une liaison pivot sans frottement. On repère l'écart du club avec la verticale par l'angle φ (voir figure 19.14).

a. A l'aide du théorème du moment cinétique et en négligeant les frottements de l'air, établir l'équation différentielle du mouvement.

b. On mesure la période des petites oscillations : $T_0 = 2,3$ s. Exprimer J_c en fonction de m_c , g , h_c et T_0 . Application numérique avec $m_c = 0,32$ kg, $g = 9,8$ m·s⁻² et $h_c = 80$ cm (données typiques pour un club de type « driver », utilisé pour frapper les coups les plus longs).

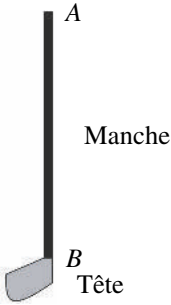


Figure 19.13

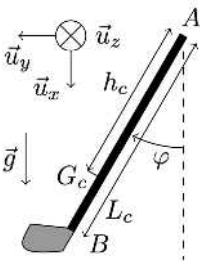


Figure 19.14

CORRIGÉS

19.1 Ordres de grandeur des moments cinétiques

1. Le moment d'inertie de la Terre par rapport à l'axe des pôles vaut $J = 0,33M_T R_T^2 = 8,1.10^{37} \text{ kg}\cdot\text{m}^2$. Son moment cinétique vaut $L = J\dot{\theta}$ avec $\dot{\theta} = \frac{2\pi \text{ rad}}{1 \text{ jour}} = 7,3.10^{-5} \text{ rad}\cdot\text{s}^{-1}$. Numériquement, on trouve : $L \simeq 6.10^{33} \text{ J}\cdot\text{s}$.
2. Le moment cinétique par rapport au centre de la trajectoire circulaire confondue avec le noyau vaut : $L = mr_0^2 \dot{\theta} = mr_0^2 2\pi f$, avec $r_0 = 5,3.10^{-11} \text{ m}$. Numériquement, on trouve $L = 1.10^{-34} \text{ J}\cdot\text{s}$.
3. $L = J\dot{\theta}$ avec $J = mR^2 = 3,1.10^{-1} \text{ kg}\cdot\text{m}^2$ et $\dot{\theta} = \frac{1000 \times 2 \times 3,14}{60} = 105 \text{ rad}\cdot\text{s}^{-1}$ d'où $L = 33 \text{ J}\cdot\text{s}$.

19.2 Etude d'une poulie

1. Lorsque la poulie tourne d'un angle θ , on déroule une longueur de fil $\ell = R\theta$. Avec le choix des orientations de l'angle θ et de l'axe (Ax) , on a donc $\dot{x} = R\dot{\theta}$.
2. On étudie le mouvement de la poulie dans le référentiel terrestre galiléen. Cette poulie est soumise à l'action de la liaison pivot d'axe (Oz) , à la force exercée par l'opérateur et à la traction exercée par le fil. La poulie est à l'équilibre lorsque la somme des moments de toutes ces forces par rapport à l'axe de rotation (Oz) est nulle. Le moment de l'action de la liaison pivot idéale est nul, reste celui des autres forces. Si l'on suppose le fil sans masse, il transmet intégralement le poids $m\vec{g}$ et l'applique en A avec un bras de levier R . L'opérateur exerce une force $\vec{F}$ verticale au point B diamétralement opposé à A avec un bras de levier R également. On trouve alors $FR = mgR$ soit $F = mg = 50 \text{ N}$.
3. Une relation cinématique relie ces deux mouvements puisque la masse est attachée à un fil qui s'enroule autour de la poulie. D'après la question 1, on a :

$$R\dot{\theta} = \dot{x}.$$

On applique le principe fondamentale de la dynamique à la masse m soumise à son poids $m\vec{g}$ et à la tension du fil sur la masse : $\vec{T}_1 = -T_1\vec{u}_x$. En projection sur la verticale, on obtient :

$$m\ddot{x} = mg - T_1.$$

On applique la loi du du moment cinétique à la poulie en rotation autour d'un axe fixe soumis à la tension du fil qui s'applique au point I : $\vec{T}_2 = T_2\vec{u}_x$ dont le moment par rapport à l'axe de rotation vaut RT_2 :

$$I\ddot{\theta} = RT_2.$$

Si l'on suppose le fil sans masse, il transmet intégralement la tension et $T_1 = T_2 = T$. En combinant ces équations et en utilisant l'expression de I fournie dans l'énoncé, on trouve :

$$\ddot{\theta} = \frac{g}{R} \frac{1}{1 + 0,5 \frac{m_p}{m}} \quad ; \quad \ddot{x} = g \frac{1}{1 + 0,5 \frac{m_p}{m}} \quad \Rightarrow \quad T = mg - m\ddot{x} = mg \left(\frac{0,5 \frac{m_p}{m}}{1 + 0,5 \frac{m_p}{m}} \right).$$

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

Le mouvement de la masse m est uniformément accéléré avec une accélération inférieure à g , qui tend vers g si la poulie est de masse négligeable et vers 0 si la poulie est très massive comparativement à m .

19.3 Mouvement d'une sphère attachée au bout d'un fil

On étudie le mouvement de la sphère de petite taille dans le référentiel terrestre galiléen. Elle est soumise à son poids, à la réaction du support assurant le mouvement horizontal et à la tension du fil.

1. Le moment cinétique en O de la sphère vaut $\vec{L}_O = m\ell_0 v_0 \vec{u}_z$. Son moment cinétique par rapport à l'axe Oz vaut $L_{(Oz)} = m\ell_0 v_0$.

2. Le moment du poids et de la réaction du support par rapport à l'axe (Oz) sont nuls car ces deux forces s'appliquent parallèlement à (Oz) . Le moment de la tension du fil est également nul puisque cette force est parallèle à $\vec{OM}$ (force centrale de centre O). On est donc dans un cas de conservation du moment cinétique par rapport à l'axe (Oz) .

Le moment cinétique, après raccourcissement du fil, vaut $m\ell_1 v_1$ puisque la trajectoire est à nouveau circulaire. On en déduit $\ell_0 v_0 = \ell_1 v_1$. La vitesse de la sphère $v_1 = v_0 \frac{\ell_0}{\ell_1} = 2 \text{ m}\cdot\text{s}^{-1}$ double par rapport à la vitesse initiale.

3. L'énergie cinétique $E_c = \frac{1}{2}mv^2$ passe de $E_{c_0} = \frac{1}{2}mv_0^2$ à $E_{c_1} = \frac{1}{2}mv_1^2$. Elle est quadruplée et passe de $E_{c_0} = 0,05 \text{ J}$ à $0,20 \text{ J}$.

4. D'après le théorème de l'énergie cinétique, $\Delta E_c = E_{c_1} - E_{c_0} = W$ où W est le travail des forces appliquées à la bille. Or, le poids et la réaction du support sont perpendiculaires au mouvement donc leur puissance est nulle. Seule la tension du fil a pu procurer cette énergie à la bille. On peut comprendre cela en observant que, pendant le raccourcissement du fil, la vitesse de la bille comporte une composante radiale $\vec{r}\vec{u}_r$ de sorte que la puissance de la tension du fil $\vec{T} = T\vec{u}_r$ vaut $\mathcal{P} = T\dot{r}$. Cette puissance est positive puisque $\dot{r} < 0$ (le rayon de la trajectoire diminue) et $T < 0$ (le fil exerce une traction sur la bille).

C'est en fait la méthode que l'on utilise instinctivement, pour faire tourner une masse accrochée au bout d'un fil de plus en plus vite.

19.4 Chute d'un arbre

On étudie le mouvement de la tige en rotation autour de son point d'appui considéré comme un point fixe. La chute se fait dans un plan vertical (xOy) et on repère la position de l'arbre par l'angle θ qu'il fait avec la verticale (figure 19.15). L'arbre est soumis à son poids $\vec{P} = m\vec{g}$ qui s'applique en son centre de gravité situé au milieu de la tige et à la réaction du sol qui s'applique au point d'appui.

1. On applique le théorème du moment cinétique à l'arbre considéré comme un solide en rotation autour d'un axe fixe (Oz) dans le référentiel terrestre galiléen :

$$\frac{dJ_{(Oz)}}{dt} = \mathcal{M}_{(Oz)}(\text{forces}).$$

Le moment par rapport à (Oz) de la réaction du sol est nul car cette force s'applique sur l'axe de rotation. Le moment du poids est égal au produit de mg par le bras de levier $b = \frac{L}{2} \sin \theta$. Son signe est positif car le poids tend à faire tourner l'arbre dans le sens positif par rapport à l'axe de rotation. Le moment cinétique vaut $J_{(Oz)} = I\dot{\theta} = \frac{1}{3}mL^2\dot{\theta}$. On en déduit :

$$\frac{1}{3}mL^2\ddot{\theta} = mg\frac{L}{2}\sin \theta.$$

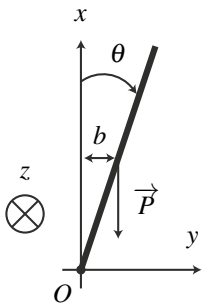


Figure 19.15

2. On cherche une intégrale première de cette équation en la multipliant par $\dot{\theta}$ et en intégrant une fois :

$$\frac{1}{3}mL^2\ddot{\theta}\dot{\theta} = mg\frac{L}{2}\sin(\theta)\dot{\theta} \Rightarrow \frac{1}{3}mL^2\frac{\dot{\theta}^2}{2} = mg\frac{L}{2}(\cos(\theta_0) - \cos(\theta))$$

en considérant que l'arbre fait initialement un angle θ_0 avec la verticale. Etant donné que $\theta > \theta_0$, $\cos(\theta_0) - \cos(\theta) > 0$ et comme $\dot{\theta} > 0$, on obtient :

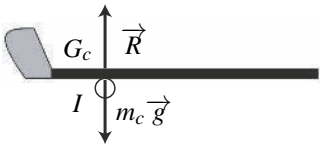
3. On écrit $\dot{\theta} = \frac{d\theta}{dt}$ et on sépare les variables pour obtenir $\sqrt{\frac{3g}{L}}dt = \frac{d\theta}{\sqrt{\cos(\theta_0) - \cos \theta}}$.

4. Le temps de chute de l'arbre est alors le temps nécessaire pour que θ passe de $\theta_0 = 5^\circ = 0,0873 \text{ rad}$ à $\theta_f = \frac{\pi}{2}$ soit :

$$\Delta t = \sqrt{\frac{L}{3g}} \int_{\theta_0}^{\frac{\pi}{2}} \frac{d\theta}{\sqrt{\cos \theta_0 - \cos \theta}} = 5,1 \times \sqrt{\frac{30}{3 \times 10}} = 5,1 \text{ s.}$$

19.5 Le swing de golf

1. Si le club, posé sur le doigt, tient en équilibre, la somme des moments des actions extérieures sur le club, pris au point de contact I entre le doigt et le club est nulle. Or le club n'est soumis qu'à son poids et à l'action de contact du doigt. Le moment en I de cette dernière est nulle puisque c'est une force exercée en I , donc le moment en I du poids du club est aussi nul. On en conclut que le centre d'inertie du club est sur la verticale passant par I .



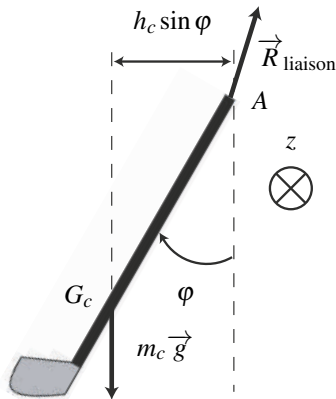
2. Détermination de J_c .

CHAPITRE 19 – MOMENT CINÉTIQUE ET SOLIDE EN ROTATION

- a. Les actions extérieures sur le système {club} sont :
– son poids $m_c \vec{g}$ appliqué en G_c , dont le moment par rapport à l'axe est :

$$\mathcal{M}_{(Az)}(\text{poids}) = -mgh_c \sin \varphi.$$

En effet, le bras de levier vaut $h_c \sin \varphi$ et le signe – provient du fait que le poids tend à faire tourner le club dans le sens indirect autour de l'axe (Az) .
– l'action de liaison, dont le moment par rapport à l'axe est nul puisque c'est une liaison pivot sans frottement.
Le théorème du moment cinétique par rapport à l'axe (Az) , appliqué au système {club} dans le référentiel terrestre galiléen s'écrit :



$$\frac{dL_{(Az)}}{dt} = \mathcal{M}_{(Az)}(\text{poids}) + \mathcal{M}_{(Az)}(\text{liaison}) \Rightarrow J_c \ddot{\varphi} = -m_c g h_c \sin \varphi$$

- b. Pour les petites oscillations, $\varphi \ll \frac{\pi}{2}$ rad donc $\sin \varphi \simeq \varphi$ et l'équation différentielle peut être approchée par :

$$\ddot{\varphi} = -\frac{m_c g h_c}{J_c} \varphi.$$

C'est l'équation d'un oscillateur harmonique de pulsation $\omega_0 = \sqrt{\frac{m_c g h_c}{J_c}}$ et période $T_0 =$

$$\frac{2\pi}{\omega_0} = 2\pi \sqrt{\frac{J_c}{m_c g h_c}}. \text{ On a donc : } J_c = \frac{T_0^2 m_c g h_c}{4\pi^2}. \text{ Numériquement : } J_c = 0,34 \text{ kg} \cdot \text{m}^2.$$

Mouvement dans un champ de force centrale. Champs newtoniens.

20

L'objet de ce chapitre est l'étude du mouvement d'un point matériel soumis à une force centrale conservative et plus particulièrement à une force newtonienne. Les champs newtoniens sont de première importance puisqu'ils régissent les mouvements des planètes autour du Soleil et qu'ils interviennent lors de l'interaction entre deux particules chargées.

1 Force centrale conservative

1.1 Qu'est-ce qu'une force centrale conservative ?

a) Notations

Soit $M(m)$ un point matériel étudié dans le référentiel $\mathcal{R}$ galiléen et O un point fixe de $\mathcal{R}$. On note $\overrightarrow{OM}$ son vecteur position, $\vec{v}$ sa vitesse, $\vec{p} = m\vec{v}$ sa quantité de mouvement, $\vec{a}$ son accélération et $\vec{L}_O = \overrightarrow{OM} \wedge m\vec{v}$ son moment cinétique en O .

M est soumis à la force $\vec{f}$ supposée être une **force centrale conservative** de centre O .

b) Définition

Une force centrale de centre O est une force dont la droite d'action passe constamment par O .

La force $\vec{f}$ est alors colinéaire au vecteur $\overrightarrow{OM}$ et les coordonnées adaptées au problème sont les coordonnées sphériques d'origine O . En posant $\|\overrightarrow{OM}\|$ et $\vec{u}_r = \frac{\overrightarrow{OM}}{r}$, la force s'écrit :

$$\vec{f} = f_r(r)\vec{u}_r.$$

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

Une force conservative dérive d’une énergie potentielle E_p définie par la relation :

$$dE_p = -\delta W(\vec{f}) = -\vec{f} \cdot d\vec{OM}.$$

En coordonnées sphériques, on peut alors écrire :

$$dE_p = -\vec{f} \cdot d\vec{OM} = -f_r(r)\vec{u}_r \cdot d\vec{OM}.$$

$\vec{u}_r \cdot d\vec{OM}$ est le déplacement infinitésimal radial de M dessiné sur la figure 20.1. Il vaut dr . On en déduit :

$$dE_p = -f_r(r)dr \Leftrightarrow f_r(r) = -\frac{dE_p}{dr}.$$

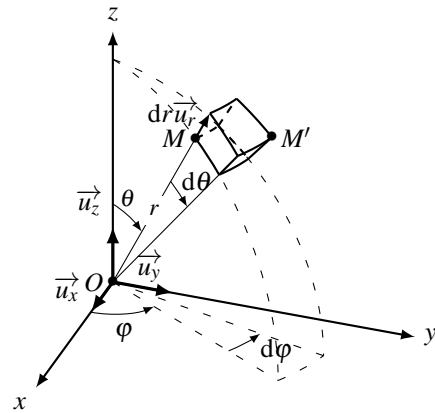


Figure 20.1 – Coordonnées sphériques et déplacement radial. M' est tel que $\vec{MM'} = d\vec{OM}$.

En coordonnées sphériques, une force centrale conservative s’écrit :

$$\vec{f} = f_r(r)\vec{u}_r = -\frac{dE_p}{dr}\vec{u}_r. \tag{20.1}$$

On peut noter que le problème est à symétrie sphérique de centre O et que la force et l’énergie potentielle ne dépendent que de la coordonnée radiale r .

1.2 Exemples de forces centrales conservatives

Dans ce paragraphe, on décrit les forces centrales conservatives les plus courantes.

a) Force d’interaction gravitationnelle

Soit O un point matériel fixe de masse m_1 et M un point matériel de masse m_2 . On note $\mathcal{G} = 6,67 \cdot 10^{-11} \text{ m}^3 \cdot \text{kg}^{-1} \cdot \text{s}^{-2}$ la constante gravitationnelle. La force gravitationnelle exercée par $O(m_1)$ sur $M(m_2)$ est donnée par la relation :

$$\vec{f} = -\mathcal{G} \frac{m_1 m_2}{r^2} \vec{u}_r,$$

où $r = \|\vec{OM}\|$ et $\vec{u}_r = \frac{\vec{OM}}{r}$. Elle est portée par la droite (OM) . Il s’agit donc d’une force centrale de centre O . Elle est également conservative et, d’après l’équation (20.1), elle dérive de l’énergie potentielle E_p telle que :

$$\frac{dE_p}{dr} = \mathcal{G} \frac{m_1 m_2}{r^2} = \frac{d}{dr} \left(-\mathcal{G} \frac{m_1 m_2}{r} \right) \implies E_p = -\mathcal{G} \frac{m_1 m_2}{r},$$

en fixant la référence d’énergie potentielle nulle à l’infini.

b) Force d'interaction coulombienne

Soit O un point matériel fixe de charge q_1 et M un point matériel de charge q_2 . On note $\epsilon_0 = 8,85 \cdot 10^{-12} \text{ F}\cdot\text{m}^{-1}$ la permittivité du vide. La force coulombienne exercée par $O(q_1)$ sur $M(q_2)$ est donnée par la relation :

$$\vec{f} = \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{r^2} \vec{u}_r,$$

où $r = \|\vec{OM}\|$ et $\vec{u}_r = \frac{\vec{OM}}{r}$. Elle est portée par la droite (OM) . Il s'agit donc d'une force centrale de centre O . Elle est également conservative et, d'après l'équation (20.1), elle dérive de l'énergie potentielle E_p telle que :

$$\frac{dE_p}{dr} = -\frac{q_1 q_2}{4\pi\epsilon_0} \frac{1}{r^2} = \frac{d}{dr} \left(\frac{q_1 q_2}{4\pi\epsilon_0 r} \right) \implies E_p = \frac{q_1 q_2}{4\pi\epsilon_0 r}.$$

en fixant la référence d'énergie potentielle nulle à l'infini.
Le module des forces gravitationnelle et coulombienne est proportionnel à l'inverse du carré de la distance OM . On parle de forces en $\frac{1}{r^2}$. Ce sont des forces newtoniennes.

Les forces gravitationnelle et coulombienne sont un type particulier de forces centrales conservatives appelées **forces newtoniennes**. Elles s'expriment en coordonnées sphériques sous la forme :

$$\vec{f} = -\frac{K}{r^2} \vec{u}_r$$

et dérivent de l'énergie potentielle :

$$E_p(r) = -\frac{K}{r}.$$

c) Force de rappel élastique (force harmonique)

Soit O un point fixe et M un point matériel mobile lié à O par une force de rappel élastique :

$$\vec{f} = -k(r - r_0) \vec{u}_r = -m\omega_0^2(r - r_0) \vec{u}_r$$

où $\omega_0^2 = \frac{k}{m}$, $r = \|\vec{OM}\|$ et $\vec{u}_r = \frac{\vec{OM}}{r}$. Cette force est portée par la droite (OM) à tout instant. Elle est donc centrale de centre O . Elle est également conservative et, d'après l'équation (20.1), elle dérive de l'énergie potentielle E_p telle que :

$$\frac{dE_p}{dr} = k(r - r_0) = \frac{d}{dr} \left(\frac{1}{2} k(r - r_0)^2 \right).$$

On choisit généralement la référence d'énergie potentielle nulle en $r = r_0$, ce qui conduit à :

$$E_p = \frac{1}{2} k(r - r_0)^2 = \frac{1}{2} m\omega_0^2(r - r_0)^2.$$

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

1.3 Observations de mouvements à force centrale conservative

La première étude des mouvements à force centrale conservative est l'étude des trajectoires des planètes du système solaire. L'analyse scientifique précise de ces mouvements remonte au XVII^e siècle. Elle est résumée en trois lois énoncées par J. Kepler en 1609 et 1618.

a) Les lois expérimentales de Kepler

Ces lois énoncent les caractéristiques principales des mouvements des planètes soumises à l'attraction gravitationnelle exercée par le Soleil. Elles font appel à la notion d'ellipse qui est définie dans l'annexe mathématique.

Loi des orbites : Les planètes du système solaire décrivent des trajectoires elliptiques dont le Soleil occupe l'un des foyers.

Loi des aires : Les aires balayées par le segment $[SP]$ (où S représente le Soleil et P une planète) pendant des intervalles de temps égaux sont égales.

Loi des périodes : Le carré de la période de révolution d'une planète autour du Soleil est proportionnel au cube du demi-grand axe a de sa trajectoire elliptique :

$$\frac{T^2}{a^3} = Cte.$$

On peut calculer cette constante à l'aide des paramètres du mouvement orbital de la Terre autour du Soleil à savoir :

- sa période $T = 1$ an,
- la distance Terre - Soleil qui définit l'unité astronomique $1 \text{ ua} = 150.10^6 \text{ km}$.

On trouve alors
$$\frac{T^2}{a^3} = \frac{(1 \text{ an})^2}{(1 \text{ ua})^3} = \frac{(365 \times 24 \times 3600)^2}{(150.10^9)^3} = 2,9.10^{-19} \text{ s}^2 \cdot \text{m}^{-3}.$$

La valeur obtenue est valable pour l'ensemble des planètes du système solaire.

Une partie de ces lois s'applique à tous les mouvements à force centrale conservative. La nature de la trajectoire est cependant spécifique de la dépendance en $\frac{1}{r^2}$ de l'attraction gravitationnelle et ne s'applique donc qu'aux forces newtoniennes. Le but de ce cours est de démontrer une partie de ces résultats à partir des caractéristiques de la loi d'attraction gravitationnelle.

b) Quelques données sur les planètes du système solaire

Pour commencer, la culture générale d'un scientifique comporte quelques connaissances sur le système solaire et il est bien de connaître :

- le nom et l'ordre des planètes du système solaire. Pour cela existe de nombreuses phrases mnémotechniques comme « **Me Voilà Tout Mouillé, J'étais Sous Un Nuage** » ;
- la masse de la Terre $M_T = 6,0.10^{24} \text{ kg}$, son rayon $R_T = 6,4.10^3 \text{ km}$, la distance Terre-Soleil $D_{TS} = 1 \text{ ua} = 150.10^6 \text{ km}$ et la masse du Soleil $M_S = 2,0.10^{30} \text{ kg}$.

On peut également noter que les excentricités e des trajectoires des planètes sont faibles. Par conséquent, leurs orbites sont quasi-circulaires. En effet, e caractérise l'écart entre un cercle et une ellipse : l'ellipse est réduite à un cercle lorsque $e = 0$ et son allongement augmente

lorsque e augmente. On peut se rendre compte de cette propriété sur la figure 20.2 et quelques propriétés des ellipses sont détaillées dans l'annexe mathématique. On remarque que, lorsque $e \leq 0,2$, l'ellipse est très proche d'un cercle dont le centre est décalé par rapport à S .

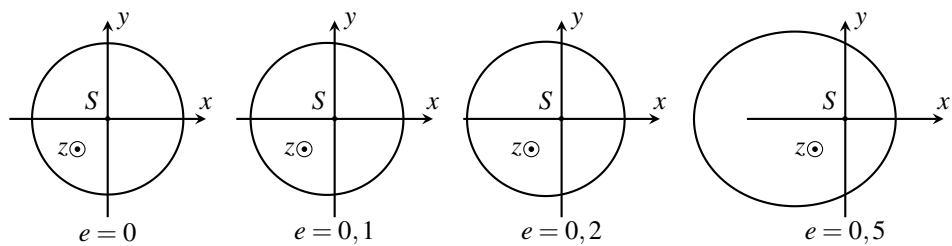


Figure 20.2 – Évolution de la forme d'une l'ellipse de foyer S en fonction de son excentricité e . Les orbites des planètes du système solaire ont une excentricité inférieure ou égale à 0,21 .

	Mercure	Venus	Terre	Mars	Jupiter	Saturne	Uranus	Neptune
d	57,9	108	150	228	778	1430	2870	4500
e	0,206	0,007	0,017	0,093	0,005	0,055	0,048	0,009
a	0,387	0,723	1	1,524	5,202	9,555	19,22	30,11
T_1	0,241	0,615	1	1,882	11,86	29,46	84,02	164,8

Tableau 20.1 – Principales caractéristiques des orbites des planètes du système solaire : d = rayon moyen en million de km, e = excentricité, a demi grand axe en unité astronomique, T_1 = période de révolution sidérale en année.

2 Généralités sur les forces centrales conservatives

2.1 Conséquence du caractère central de la force

Dans ce paragraphe, on montre que le caractère central de la force implique la conservation du moment cinétique en O du système. On en déduit que le mouvement est plan et vérifie la loi des aires.

a) Conservation du moment cinétique

On applique le théorème du moment cinétique au système M de masse m par rapport au point O fixe dans le référentiel $\mathcal{R}$ galiléen :

$$\frac{d\vec{L}_O}{dt} = \vec{\mathcal{M}}_O(\vec{f}).$$

$\vec{f}$ est une force centrale. Par définition, elle est colinéaire au vecteur $\vec{OM}$. D'après les propriétés du produit vectoriel entre vecteurs colinéaires, on trouve $\vec{\mathcal{M}}_O(\vec{f}) = \vec{OM} \wedge \vec{f} = \vec{0}$.

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

Le moment cinétique par rapport au centre de force O :

$$\vec{L}_O = \vec{OM} \wedge m \vec{v}$$

est conservé au cours du mouvement. C'est une intégrale première du mouvement.

On peut déterminer la valeur du moment cinétique en O en utilisant un instant particulier pour lequel on connaît la position et la vitesse. On la détermine souvent à l'aide des conditions initiales qui précisent la position initiale $\vec{OM}_0$ ainsi que la vitesse initiale $\vec{v}_0$.

b) Planéité de la trajectoire

La conservation du moment cinétique implique que le mouvement est plan.

En effet, dans le cas général, les conditions initiales (position $\vec{OM}_0$ et vitesse $\vec{v}_0$) permettent de définir un plan $\mathcal{P} = (O, \vec{OM}_0, \vec{v}_0)$ et d'évaluer le moment cinétique à l'instant initial :

$$\vec{L}_O = \vec{OM}_0 \wedge m \vec{v}_0.$$

Par définition du produit vectoriel, $\vec{L}_O$ est perpendiculaire au plan $\mathcal{P}$ et on peut redéfinir le plan $\mathcal{P}$ comme le plan passant par O et perpendiculaire à $\vec{L}_O$. Or, le moment cinétique $\vec{L}_O$ est conservé au cours du mouvement et, par définition du produit vectoriel, $\vec{OM}$ est perpendiculaire à $\vec{L}_O$. Il s'en suit que le point M est situé dans le plan passant par O et perpendiculaire à $\vec{L}_O$ fixe. M est donc situé dans le plan $\mathcal{P}$ à tout instant. Son mouvement est plan.

c) Cas des coordonnées polaires

En pratique, le moment cinétique $\vec{L}_O$ étant conservé et non nul, on pose $\vec{u}_z = \frac{\vec{L}_O}{\|\vec{L}_O\|}$. La définition de $\vec{L}_O$ et les propriétés du produit vectoriel imposent à $\vec{OM}$ d'être perpendiculaire à $\vec{L}_O$ donc à $\vec{u}_z$. Le point M appartient alors au plan passant par O et perpendiculaire à $\vec{u}_z$, c'est-à-dire au plan (Oxy) . Les coordonnées adaptées à ce type de problème sont les coordonnées polaires d'origine O qui est le centre de la force $\vec{f}$ et on obtient :

$$\vec{OM} = r \vec{u}_r \quad \text{et} \quad \vec{v} = \dot{r} \vec{u}_r + r \dot{\theta} \vec{u}_\theta,$$

d'où on tire :

$$\vec{L}_O = mr^2 \dot{\theta} \vec{u}_z.$$

Remarque

Les forces centrales de centre O sont à symétrie sphérique de centre O : le problème reste inchangé par une rotation quelconque autour de O . La conservation du moment cinétique $\vec{L}_O$ implique que l'axe $(O, \vec{L}_O)$ est un axe de symétrie du problème. On choisit

généralement d'orienter l'axe (Oz) dans le même sens que $(O, \vec{L}_O)$ et les coordonnées cylindriques sont alors les coordonnées naturelles. L'orientation de l'axe (Oz) dans le sens de $\vec{L}_O$ implique que $\dot{\theta}$ est positif donc que θ est croissant. Le mouvement est réalisé dans le sens direct autour de (Oz) .

d) Loi des aires

Le moment cinétique en O ainsi que la masse m sont conservés. La quantité $r^2 \dot{\theta} = \frac{\|\vec{L}_O\|}{m}$ l'est donc également.

On définit la constante des aires : $\mathcal{C} = r^2 \dot{\theta} = \frac{\|\vec{L}_O\|}{m}$.

Cette constante, exprimée en $\text{m}^2 \cdot \text{s}^{-1}$, peut être évaluée à l'instant initial : $\mathcal{C} = \|\vec{OM}_0 \wedge \vec{v}_0\|$. La dimension de $\mathcal{C} dt$ est celle d'une aire que l'on interprète géométriquement à l'aide de la figure 20.3.

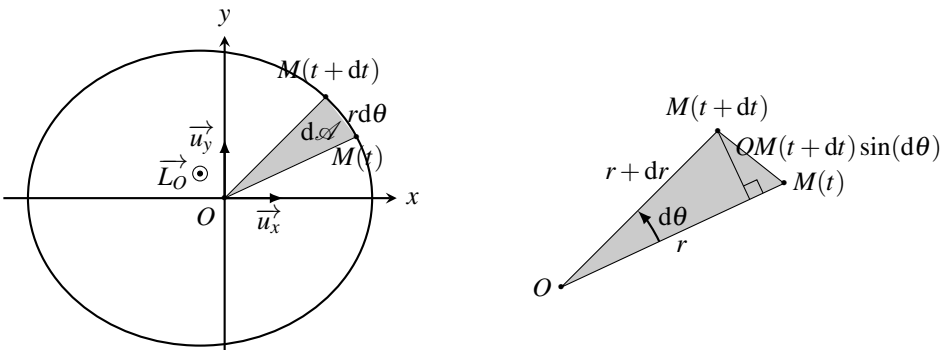


Figure 20.3 – Aire $d\mathcal{A}$ balayée par le rayon vecteur $\vec{OM}$ pendant dt .

L'aire $d\mathcal{A}$ balayée par le rayon vecteur $\vec{OM}$ pendant dt est grisée sur la figure 20.3. Elle correspond à l'aire du triangle $OM(t)M(t+dt)$:

$$\begin{aligned} d\mathcal{A} &= \frac{1}{2} \|\vec{OM}(t)\| \|\vec{OM}(t+dt)\| \sin(d\theta) \\ &= \frac{1}{2} r^2 d\theta, \end{aligned}$$

en assimilant $\sin(d\theta)$ à $d\theta$ et $r+dr$ à r . L'aire balayée pendant dt vaut donc :

$$d\mathcal{A} = \frac{1}{2} \mathcal{C} dt.$$

La constante $\mathcal{V} = \frac{d\mathcal{A}}{dt} = \frac{\mathcal{C}}{2}$ est appelée **vitesse aréolaire**. La loi des aires stipule que la vitesse aréolaire est constante et précise que sa valeur est $\frac{\mathcal{C}}{2}$.

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

La conservation du moment cinétique implique la loi des aires : les aires balayées par le rayon vecteur pendant des intervalles de temps égaux de durée Δt sont égales et valent :

$$\Delta \mathcal{A} = \frac{\mathcal{C}}{2} \Delta t.$$

2.2 Conséquence du caractère conservatif de la force

Dans ce paragraphe, on montre que le caractère conservatif de la force implique la conservation de l'énergie mécanique du système. On en déduit que l'on peut étudier qualitativement le mouvement radial du point M en introduisant une énergie potentielle effective.

a) L'énergie mécanique est conservée

On applique le théorème de l'énergie cinétique au système M de masse m dans le référentiel $\mathcal{R}$ galiléen. M n'est soumis qu'à la seule force $\vec{f}$ qui est conservative et dérive de l'énergie potentielle E_p . On en déduit :

$$\frac{dE_m}{dt} = \frac{dE_c + E_p}{dt} = 0.$$

L'énergie mécanique :

$$E_m = E_c + E_p = \frac{1}{2}mv^2 + E_p$$

est conservée au cours du mouvement. C'est une intégrale première du mouvement.

On détermine la valeur de l'énergie mécanique à l'aide des conditions initiales qui précisent la position initiale et donc l'énergie potentielle initiale E_{p0} , ainsi que la vitesse initiale et donc l'énergie cinétique initiale E_{c0} .

b) Notion d'énergie potentielle effective

On a vu précédemment que le mouvement est plan et qu'on l'étudie en coordonnées polaires. Dans ces conditions, l'énergie potentielle ne dépend que de r et on la note $E_p(r)$. On utilise les relations cinématiques $\vec{OM} = r\vec{u}_r$ et $\vec{v} = \dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta$ pour calculer l'énergie cinétique $E_c = \frac{1}{2}mv^2 = \frac{1}{2}m(\dot{r}^2 + r^2\dot{\theta}^2)$. On aboutit alors à l'équation énergétique :

$$E_m = \frac{1}{2}m(\dot{r}^2 + r^2\dot{\theta}^2) + E_p(r),$$

puis au système d'équations du mouvement :

$$\begin{cases} \mathcal{C} &= r^2\dot{\theta} \\ E_m &= \frac{1}{2}m(\dot{r}^2 + r^2\dot{\theta}^2) + E_p(r). \end{cases}$$

Le problème étant à deux degrés de liberté r et θ , la connaissance de ces deux équations permet de résoudre complètement le problème à condition de connaître la fonction $E_p(r)$.

c) Étude qualitative du mouvement radial

Dans un premier temps, on ne cherche pas à résoudre complètement le problème mais uniquement à trouver l'évolution de la distance radiale $\|\overrightarrow{OM}\| = r$. Pour cela, on élimine θ entre les deux équations du système précédent. On écrit : $\dot{\theta} = \frac{\mathcal{C}}{r^2}$ et on remplace dans l'expression de l'énergie mécanique pour obtenir l'énergie mécanique :

$$E_m = \frac{1}{2}m\left(\dot{r}^2 + r^2\frac{\mathcal{C}^2}{r^4}\right) + E_p(r) = \frac{1}{2}m\dot{r}^2 + \frac{1}{2}m\frac{\mathcal{C}^2}{r^2} + E_p(r).$$

On pose alors :

$$E_{p_{\text{eff}}}(r) = \frac{1}{2}m\frac{\mathcal{C}^2}{r^2} + E_p(r)$$

et on aboutit à l'équation :

$$E_m = \frac{1}{2}m\dot{r}^2 + E_{p_{\text{eff}}}(r).$$

$E_{p_{\text{eff}}}$ est la somme de la partie de l'énergie cinétique liée au mouvement orthoradial et de l'énergie potentielle. Elle est homogène à une énergie et s'apparente à l'énergie potentielle utile pour étudier le mouvement radial de M . On l'appelle énergie potentielle effective.

Le mouvement radial s'apparente à un mouvement conservatif à un degré de liberté dans une **énergie potentielle effective** :

$$E_{p_{\text{eff}}}(r) = \frac{1}{2}m\frac{\mathcal{C}^2}{r^2} + E_p(r).$$

Ce type d'équation a été étudié dans le chapitre sur l'énergie. On applique la même méthode. On utilise la conservation de l'énergie mécanique et le fait que le terme $\frac{1}{2}m\dot{r}^2$ est toujours positif pour étudier qualitativement le mouvement radial et déterminer l'ensemble des positions accessibles. Pour cela, il est nécessaire de connaître l'évolution de l'énergie potentielle en fonction de r . On présente la méthode sur le cas particulier de l'attraction gravitationnelle, mais elle peut être adaptée à d'autres forces.

3 Cas particulier de l'attraction gravitationnelle

3.1 Position du problème

On s'intéresse au mouvement d'un point M de masse m en interaction gravitationnelle avec un astre de masse M_A situé en O . On se place dans la situation où $M_A \gg m$ de sorte que l'action de la masse m sur M_A ne permet pas de mettre M_A en mouvement. Dans ce cas, on a vu au paragraphe 1.2, que le point M est soumis à la force d'attraction gravitationnelle : $\overrightarrow{f} = -\mathcal{G}\frac{mM_A}{r^2}\overrightarrow{u_r}$ qui dérive de l'énergie potentielle $E_p(r) = -\mathcal{G}\frac{mM_A}{r}$.

Cette force est une force centrale conservative. Le mouvement satisfait donc aux conservations du moment cinétique et de l'énergie mécanique. Le mouvement est plan, on l'étudie en

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

coordonnées polaires et on aboutit aux équations du mouvement :

$$\begin{cases} \mathcal{C} &= r^2 \dot{\theta} \\ E_m &= \frac{1}{2} m (\dot{r}^2 + r^2 \dot{\theta}^2) - \mathcal{G} \frac{mM_A}{r}. \end{cases} \quad (20.2)$$

3.2 Étude qualitative du mouvement radial

On écrit l'énergie mécanique en éliminant $\dot{\theta}$ grâce à la constante de aires $\mathcal{C}$:

$$E_m = \frac{1}{2} m \dot{r}^2 + \frac{1}{2} m \frac{\mathcal{C}^2}{r^2} - \mathcal{G} \frac{mM_A}{r} = \frac{1}{2} m \dot{r}^2 + E_{p_{\text{eff}}}(r)$$

où

$$E_{p_{\text{eff}}}(r) = \frac{1}{2} m \frac{\mathcal{C}^2}{r^2} - \mathcal{G} \frac{mM_A}{r}.$$

a) Tracé de la courbe d'énergie potentielle effective

On étudie la courbe $E_{p_{\text{eff}}}(r)$ pour pouvoir la tracer. Sa dérivée s'écrit :

$$\frac{dE_{p_{\text{eff}}}}{dr} = -m \frac{\mathcal{C}^2}{r^3} + \mathcal{G} \frac{mM_A}{r^2}.$$

Elle s'annule en $r_0 = \frac{\mathcal{C}^2}{\mathcal{G}M_A}$ où elle atteint la valeur :

$$E_{\min} = -\frac{m}{2} \frac{\mathcal{G}^2 M_A^2}{\mathcal{C}^2} = -\frac{\mathcal{G} M_A m}{2r_0}.$$

On prend r_0 comme échelle des distances et $E_0 = -E_{\min} = \mathcal{G} M_A m / 2r_0$ comme échelle d'énergie. On rend l'expression de l'énergie potentielle effective adimensionnelle en intro-

duisant : $E_{p_{\text{eff}}}^* = \frac{E_{p_{\text{eff}}}}{E_0}$ et $r^* = \frac{r}{r_0}$. Or, $r_0 = \frac{\mathcal{C}^2}{\mathcal{G}M_A} \Leftrightarrow \frac{\mathcal{C}^2}{r_0} = \mathcal{G}M_A$ donc :

$$\frac{1}{2} m \frac{\mathcal{C}^2}{r^2} = \frac{1}{2} m \frac{\mathcal{C}^2}{r_0^2} \frac{1}{r^{*2}} = \frac{1}{2} m \frac{\mathcal{G}M_A}{r_0} \frac{1}{r^{*2}} = E_0 \frac{1}{r^{*2}}.$$

Par ailleurs, on a $-\mathcal{G} \frac{mM_A}{r} = -\mathcal{G} \frac{mM_A}{r_0} \frac{1}{r^*} = -E_0 \frac{1}{r^*}$. $E_{p_{\text{eff}}}^*(r^*)$ est alors égal à :

$$E_{p_{\text{eff}}}^*(r^*) = \frac{1}{r^{*2}} - \frac{2}{r^*},$$

que l'on a représentée à l'aide d'un logiciel dédié sur la figure 20.4.

b) Analyse de la courbe d'énergie potentielle effective

On distingue alors quatre possibilités :

- si $E_m < E_{\min}$, le mouvement est impossible ;
- si $E_m = E_{\min}$, le mouvement est réalisé à $r = r_0$ fixé. Comme $\mathcal{C} = r_0^2 \dot{\theta}$, on obtient :

$$\dot{\theta} = \frac{\mathcal{C}}{r_0^2} = Cte.$$

r et $\dot{\theta}$ étant fixés, le mouvement est circulaire et uniforme ;

- si $0 > E_m > E_{\min}$, l'ensemble des distances r accessibles est un intervalle borné compris entre les positions $r_1 = r_1^* r_0$ et $r_2 = r_2^* r_0$ où r_1^* et r_2^* sont définies sur la figure 20.4. Le système est dans un **état lié**. Il reste dans le puits de potentiel créé par l'astre A ;
- si $E_m \geq 0$, l'ensemble des distances r accessibles est un intervalle de taille infinie de la forme $[r_3, +\infty[$ avec $r_3 = r_3^* r_0$ où r_3^* est définie sur la figure 20.4. Le système est dans un **état de diffusion**. Il peut sortir du puits de potentiel créé par l'astre A .

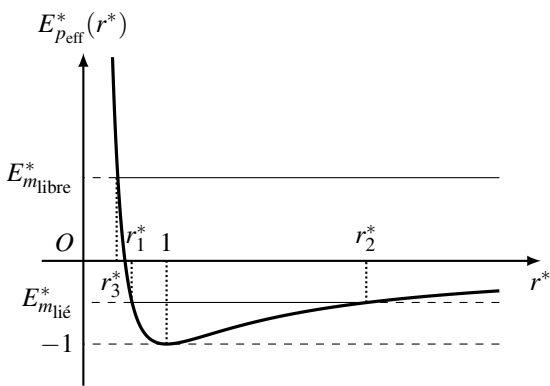


Figure 20.4 – Courbe d'énergie potentielle effective.

4 Étude directe de la trajectoire circulaire

La trajectoire circulaire est très importante car elle permet de comprendre la majorité des phénomènes concernant les états liés tout en restant simple du point de vue mathématique. Par ailleurs, les trajectoires des planètes sont très proches de trajectoires circulaires centrées sur le Soleil. Le caractère circulaire des trajectoires permet d'en faire une étude directe simple.

4.1 Position du problème

On restreint l'étude du problème précédent au cas où la trajectoire de M est un cercle de rayon r_0 et centre O , centre attracteur de la force. Le point M est soumis à la seule force d'attraction gravitationnelle $\vec{f} = -\mathcal{G} \frac{mM_A}{r_0^2} \vec{u}_r$ qui dérive de l'énergie potentielle $E_p(r) = -\mathcal{G} \frac{mM_A}{r}$.

On a démontré la planéité du mouvement mais on peut remarquer que postuler sa circularité

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

implique d’admettre sa planéité puisqu’un cercle est une courbe plane. On étudie alors le mouvement en coordonnées polaires d’origine O dans le plan de la trajectoire circulaire.

À $t = 0$, le point M est situé en M_0 et possède une vitesse $\vec{v}_0$. On choisit la référence des angles et l’orientation du plan comme sur la figure ci-contre de sorte que :

$$\overrightarrow{OM_0} = r_0 \vec{u}_x \quad \text{et} \quad \vec{v}_0 = v_0 \vec{u}_y,$$

avec $v_0 > 0$. Le point M est repéré à l’instant t à l’aide des coordonnées polaires :

$$\begin{cases} \overrightarrow{OM} &= r_0 \vec{u}_r \\ \vec{v} &= r_0 \dot{\theta} \vec{u}_\theta \\ \vec{a} &= -r_0 \dot{\theta}^2 \vec{u}_r + r_0 \ddot{\theta} \vec{u}_\theta = -\frac{v^2}{r_0} \vec{u}_r + r_0 \ddot{\theta} \vec{u}_\theta. \end{cases}$$

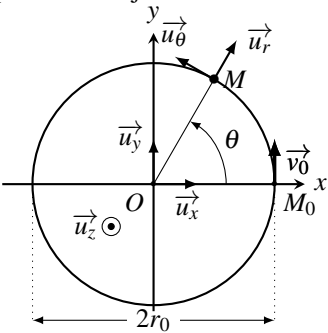


Figure 20.5 –
Mouvement circulaire.

4.2 Étude à partir du principe fondamental de la dynamique

a) Expression de la vitesse sur la trajectoire

On applique le principe fondamental de la dynamique à M de masse m :

$$m \vec{a} = -\mathcal{G} \frac{m M_A}{r_0^2} \vec{u}_r.$$

On projette alors sur $\vec{u}_r$ et on obtient :

$$-m \frac{v^2}{r_0} = -\mathcal{G} \frac{m M_A}{r_0^2}, \quad \text{soit} \quad v = \sqrt{\frac{\mathcal{G} M_A}{r_0}}.$$

La norme de la vitesse $\|\vec{v}\| = \sqrt{\frac{\mathcal{G} M_A}{r_0}}$ est donc constante. Le mouvement est uniforme.

Les trajectoires circulaires de rayon r_0 autour d’un astre attracteur de masse M_A sont parcourues à vitesse uniforme v_0 . Le principe fondamental de la dynamique permet d’établir la relation liant r_0 et v_0 :

$$v_0 = \sqrt{\frac{\mathcal{G} M_A}{r_0}}.$$

Pour mettre en orbite un satellite artificiel sur une trajectoire circulaire autour de la Terre, il faut donc contrôler parfaitement la vitesse de satellisation, à la fois en direction et en norme.

b) Expression des grandeurs énergétiques

On peut établir les expressions des différentes grandeurs énergétiques :

- l'énergie potentielle vaut : $E_p = -\mathcal{G} \frac{mM_A}{r_0}$;
- l'énergie cinétique est égale à : $E_p = \frac{1}{2}mv_0^2 = \frac{1}{2}\mathcal{G} \frac{mM_A}{r_0}$;
- l'énergie mécanique s'écrit donc :

$$E_m = E_c + E_p = -\frac{1}{2}\mathcal{G} \frac{mM_A}{r_0}. \tag{20.3}$$

Dans le cas d'une trajectoire circulaire, toutes ces énergies sont constantes et s'expriment en fonction de r_0 à l'aide de l'énergie caractéristique $E_0 = \mathcal{G} \frac{mM_A}{r_0}$.

Généralisation aux trajectoires elliptiques (MPSI) Dans le cas d'une trajectoire elliptique, les énergies cinétique et potentielle varient au cours du temps et seule l'énergie mécanique est conservée. On admet que l'expression de l'énergie mécanique que l'on vient d'établir reste valable, à condition de remplacer le rayon r_0 de la trajectoire circulaire par le demi grand-axe a de l'orbite elliptique.

L'énergie mécanique d'une planète de masse m en gravitation autour d'un astre de masse M_A vaut :

$$E_m = -\frac{\mathcal{G}M_A m}{2a}$$

où a représente le demi grand-axe de la trajectoire elliptique.

Si l'on connaît l'énergie mécanique d'un satellite, cette relation permettra de déterminer le demi-grand axe de sa trajectoire.

c) Troisième loi de Kepler : loi des périodes

La période de révolution de M autour de O est le temps nécessaire à M pour parcourir la trajectoire de circonférence $2\pi r_0$ à la vitesse constante v_0 :

$$T = \frac{2\pi r_0}{v_0} = \frac{2\pi r_0}{\sqrt{\frac{\mathcal{G}M_A}{r_0}}} \implies \frac{T^2}{r_0^3} = \frac{4\pi^2}{\mathcal{G}M_A}.$$

Généralisation aux trajectoires elliptiques (MPSI) On admet que la relation établie ici reste valable pour les trajectoires elliptiques, à condition de remplacer le rayon r_0 de la trajectoire circulaire par le demi grand-axe a de l'orbite elliptique.

Pour un corps en orbite elliptique autour d'un astre attracteur de masse M_A , le rapport du carré du demi grand-axe a au cube de la période de révolution T ne dépend que du produit $\mathcal{G}M_A$ et vaut :

$$\frac{T^2}{a^3} = \frac{4\pi^2}{\mathcal{G}M_A},$$

indépendamment du corps qui gravite autour de l'astre attracteur.

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

Exemple

Pour un corps qui gravite autour du Soleil, on peut déterminer ce rapport en utilisant les paramètres de l'orbite terrestre $T = 1$ an et $a = 1$ ua = 150.10^6 km. On trouve alors :

$$\frac{T^2}{a^3} = \frac{(1 \text{ an})^2}{(1 \text{ ua})^3} = \frac{(365 \times 24 \times 3600)^2}{(150.10^9)^3} = 2,9.10^{-19} \text{ m}^2 \cdot \text{s}^{-3}.$$

La mesure de $\mathcal{G} = 6,67.10^{-11} \text{ kg} \cdot \text{s}^2 \cdot \text{m}^{-3}$ permet ensuite de déterminer la masse du Soleil :

$$M_S = \frac{4\pi^2}{\mathcal{G}} \frac{a^3}{T^2} = \frac{4 \times 3,14^2}{6,67.10^{-11} \times 2,9.10^{-19}} = 2,0.10^{30} \text{ kg}.$$

Exemple

Pour un corps qui gravite autour de la Terre de masse M_T , on obtient :

$$\frac{T^2}{a^3} = \frac{4\pi^2}{\mathcal{G}M_T}.$$

Le produit $\mathcal{G}M_T$ est égal à $\mathcal{G}M_T = 3,986.10^{14} \text{ m}^3 \cdot \text{s}^{-2}$. La période de révolution de la Lune autour de la Terre (mois lunaire sidéral) vaut $T = 27,3$ jours. On en déduit le demi grand-axe de l'orbite de la Lune que l'on peut assimiler à la distance Terre - Lune :

$$a = \left(\frac{\mathcal{G}M_T T^2}{4\pi^2} \right)^{\frac{1}{3}} = \left(\frac{3,986.10^{14} \times (27,3 \times 24 \times 3600)^2}{4 \times 3,14^2} \right)^{\frac{1}{3}} = 383.10^3 \text{ km}.$$

Les mesures actuelles donnent $a = 384,399.10^3 \text{ km}$.

Ces exemples montrent la puissance de la troisième loi de Kepler pour déterminer un grand nombre de paramètres célestes. À ce titre, c'est une loi fondatrice en astronomie.

4.3 Application aux satellites géostationnaires

a) Caractéristiques des satellites géostationnaires

Un satellite géostationnaire est un satellite artificiel qui reste constamment au-dessus d'un même point de la surface terrestre. Il paraît immobile à un observateur situé sur Terre. Pour maintenir cette caractéristique, il est guidé depuis la Terre à l'aide d'un système de contrôle d'altitude et d'orbite qui permet de compenser diverses perturbations de sa trajectoire. De ce fait, il utilise progressivement ses réserves de carburant qui finissent par s'épuiser et limitent sa durée d'exploitation. Tous les satellites géostationnaires sont placés sur une même orbite circulaire d'altitude $h = 35\,786 \text{ km}$ et de période de révolution égale à un jour sidéral soit 23 heures 56 minutes et 4 secondes. L'orbite géostationnaire est encombrée puisqu'il y a plus de 250 satellites situés sur cet orbite. De ce fait, le positionnement sur l'orbite doit être réalisé avec une précision d'environ 50 km et les opérateurs de ces satellites doivent les évacuer lorsqu'ils arrivent en fin de vie. L'immobilité apparente d'un satellite géostationnaire dans le référentiel terrestre lui permet de faire des observations continues d'une même zone (satellite

météorologique, satellite d’alerte) et d’être visé par une antenne parabolique fixée sur Terre (satellite de télécommunications ou de télédiffusion). Dans ce paragraphe, on va montrer que le maintien d’un satellite en position fixe dans le référentiel terrestre nécessite de réunir trois conditions :

- le plan de l’orbite doit être situé dans le plan de l’équateur ;
- le mouvement du satellite doit être synchrone avec le mouvement de rotation propre de la Terre autour de l’axe des pôles ;
- l’orbite doit être circulaire.

b) Le mouvement d’un satellite géostationnaire est situé dans le plan de l’équateur

L’étude du mouvement du satellite est effectuée dans le référentiel géocentrique supposé galiléen. Dans ce référentiel, on a montré que l’orbite est située dans un plan qui contient le centre d’attraction de la force gravitationnelle. Le plan du mouvement contient donc nécessairement le centre de la Terre.

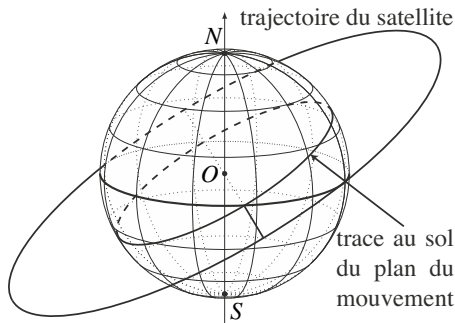


Figure 20.6 – Trajectoire d’un satellite inclinée par rapport au plan équatorial.

La figure 20.6 montre la trajectoire d’un satellite dont le plan orbital est incliné par rapport au plan équatorial terrestre, ainsi que l’intersection du plan du mouvement avec la sphère terrestre. Avec cette inclinaison, le satellite est situé tantôt au dessus de l’hémisphère nord et tantôt au dessus de l’hémisphère sud. Il n’est pas immobile pour un observateur terrestre qui va, au minimum, observer un mouvement apparent d’oscillations nord-sud. Le seul moyen d’éviter ce mouvement est d’annuler l’inclinaison de l’orbite. Le plan de l’orbite d’un satellite géostationnaire coïncide nécessairement avec le plan équatorial de la Terre.

c) Sa vitesse angulaire est constante

Le mouvement se situant dans le plan de l’équateur, le satellite doit rester en permanence à la verticale d’un point *P* fixe de l’équateur. Le déplacement angulaire du satellite pendant une durée quelconque doit être le même que celui du point *P*. La vitesse angulaire du satellite sur son orbite est donc nécessairement égale à la vitesse angulaire de rotation de la Terre.

La rotation propre de la Terre autour de son axe Nord-Sud est uniforme à la vitesse angulaire :

$$\Omega = \frac{2\pi}{T_{\text{sidéral}}},$$

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

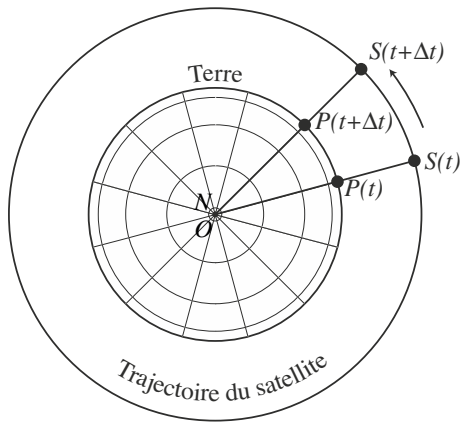


Figure 20.7 – Trajectoire d'un satellite géostationnaire vue du pôle Nord. Le satellite S suit le point de l'équateur P dans son mouvement circulaire et uniforme.

où la période $T_{\text{sidéral}}$ correspond au jour sidéral, durée nécessaire à la Terre pour faire un tour sur elle-même. La vitesse angulaire d'un satellite géostationnaire est constante. Sa période de révolution est égale à un jour sidéral.

d) L'orbite est circulaire

Le mouvement étant à force centrale, il suit la loi des aires et la constante des aires $\mathcal{C}$ vaut :

$$\mathcal{C} = r^2 \dot{\theta}.$$

La vitesse angulaire du satellite étant constante, la distance r l'est également. Le mouvement orbital d'un satellite géostationnaire est circulaire.

e) Durée de la révolution du satellite : notion de jour sidéral

Pour caractériser complètement le mouvement du satellite, on doit déterminer la durée exacte de la rotation propre de la Terre. En première approximation, cette durée correspond à la durée d'un jour soit 24 heures. Quand on veut être plus précis, il est nécessaire de tenir compte de la révolution de la Terre sur son orbite. Ainsi en un jour, la Terre fait un petit peu plus qu'un tour sur elle-même comme on peut le voir sur la figure 20.8.

En $T = 24$ h, la Terre tourne sur elle-même de $2\pi + \alpha$ où α est l'angle dont elle s'est déplacée sur son orbite autour du Soleil, c'est-à-dire $\left(\frac{1}{365,25}\right)^{\text{ième}}$ de tour (le 0,25 provient des années bissextiles). On a alors : $\alpha = \frac{2\pi}{365,25}$ puis

$$T_{\text{sidéral}} = T \frac{2\pi}{2\pi + \alpha} = T \frac{1}{1 + \frac{1}{365,25}} = T \frac{365,25}{366,25} = 23 \text{ h } 56 \text{ min } 4 \text{ s} = 86\,164 \text{ s}.$$

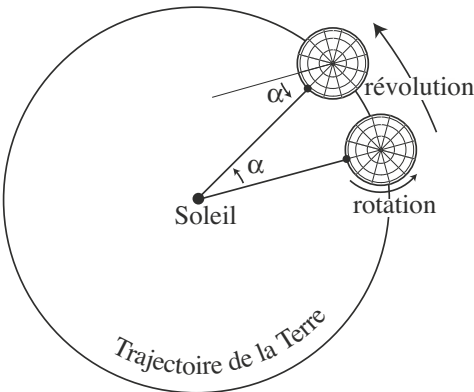


Figure 20.8 – Mouvement de la Terre en 24 heures. Le Soleil est au zénith d'un point de l'équateur terrestre à midi des jours j et $j + 1$. Du fait de la révolution de la Terre sur son orbite autour du Soleil, la Terre tourne sur elle-même d'un angle $2\pi + \alpha$ en 24 h.

f) Rayon de l'orbite géostationnaire

Connaissant la période de révolution d'un satellite géostationnaire : $T_{\text{sidéral}} = 86\,164\text{ s}$, et le produit $\mathcal{G}M_T = 3,9860 \cdot 10^{14}\text{ m}^3 \cdot \text{s}^{-2}$, on peut déterminer le rayon $r_{\text{géo}}$ de son orbite à l'aide de la troisième loi de Kepler : $\frac{T_{\text{sidéral}}^2}{r_{\text{géo}}^3} = \frac{4\pi^2}{\mathcal{G}M_T}$. On trouve donc :

$$r_{\text{géo}} = \left(\frac{T_{\text{sidéral}}^2 \mathcal{G}M_T}{4\pi^2} \right)^{\frac{1}{3}} = \left(\frac{86164^2 \times 3,9860 \cdot 10^{14}}{4 \times 3,14159^2} \right)^{\frac{1}{3}} = 42\,164\text{ km}.$$

En retranchant le rayon de la Terre à l'équateur $R_T = 6378\text{ km}$, on en déduit son altitude :

$$h = r_{\text{géo}} - R_T = 35\,786\text{ km}.$$

Un satellite géostationnaire est fixe dans le référentiel terrestre. Dans le référentiel géocentrique, il suit une trajectoire circulaire et uniforme dans le plan équatorial terrestre à l'altitude $h = 35\,786\text{ km}$ avec une période égale au jour sidéral $T_{\text{sidéral}} = 23\text{h}56\text{min}4\text{s}$. Le centre de cette trajectoire est le centre de la Terre.

Il faut avoir en tête les ordres de grandeur suivant : $T = 24\text{ h}$ et $h = 36\,000\text{ km}$.

4.4 Vitesses cosmiques

a) Vitesse minimale de mise en orbite

On peut remarquer, à l'aide de l'équation (20.3), que les trajectoires circulaires de rayon faible (orbites basses) sont celles pour lesquelles l'énergie mécanique est la plus faible. Pour cette raison, mettre un satellite en orbite basse depuis le sol terrestre nécessite moins d'énergie.

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

La **première vitesse cosmique** est une vitesse limite qui correspond à la vitesse minimale qu'il faut donner à un satellite pour le mettre en orbite autour de la Terre. Si la vitesse du satellite envoyé depuis la Terre est inférieure à cette vitesse, le satellite retombe sur Terre. Cette vitesse est égale à la vitesse v_1 d'un corps en orbite rasante autour de la Terre. On considère alors que l'altitude est négligeable devant le rayon terrestre R_T et que l'orbite est circulaire de rayon $r_0 \approx R_T$. La vitesse du satellite sur la trajectoire vaut donc :

$$v_1 = \sqrt{\frac{\mathcal{G}M_T}{R_T}}.$$

Il faut connaître l'ordre de grandeur de cette vitesse sur Terre :

$$v_1 = \sqrt{\frac{6,7 \cdot 10^{-11} \times 6,0 \cdot 10^{24}}{6,4 \cdot 10^6}} = 8 \text{ km/s} \approx 30\,000 \text{ km/h}.$$

Remarque

Les orbites basses sont très utilisées pour les satellites artificiels. Elles sont contenues dans une zone située entre 200 et 2000 km d'altitude environ. Les orbites plus basses sont inutilisables à cause des frottements dus à l'atmosphère. Leur intérêt est multiple :

- la mise en orbite de satellites proches de la Terre est moins coûteuse en énergie et permet de satelliser de gros objets comme la station spatiale internationale ;
- les orbites proches de la Terre rendent les communications avec les satellites plus faciles et plus rapides (satellites de communication) ;
- la proximité de la Terre permet d'optimiser la résolution des instruments d'observation terrestre embarqués (imagerie terrestre, météorologie ou renseignement).

b) Vitesse de libération

La **deuxième vitesse cosmique** ou **vitesse de libération** est une autre vitesse limite. Elle correspond à la vitesse minimale nécessaire pour quitter l'attraction gravitationnelle de la Terre à partir du sol. Pour la déterminer, on considère que l'on donne au système l'énergie minimale lui permettant de s'éloigner à l'infini de son point de lancement situé sur Terre. Dans ces conditions, la conservation de l'énergie mécanique entre le sol terrestre (vitesse v_2 et $r = R_T$) et l'infini (vitesse v_∞ et $r = r_\infty$) s'écrit :

$$E_m = \frac{1}{2}mv_2^2 - \mathcal{G}\frac{mM_T}{R_T} = \frac{1}{2}mv_\infty^2 - \mathcal{G}\frac{mM_T}{r_\infty}.$$

Comme on cherche la vitesse minimale, on se place dans le cas où la vitesse à l'infini est quasi-nulle et on obtient :

$$E_m = \frac{1}{2}mv_2^2 - \mathcal{G}\frac{mM_T}{R_T} = 0, \quad \text{soit} \quad v_2 = \sqrt{\frac{2\mathcal{G}M_T}{R_T}} = \sqrt{2}v_1.$$

Il faut connaître l'ordre de grandeur de cette vitesse sur Terre :

$$v_2 = \sqrt{\frac{2 \times 6,7 \cdot 10^{-11} \times 6,0 \cdot 10^{24}}{6,4 \cdot 10^6}} = 11 \text{ km/s} \approx 40\,000 \text{ km/h}.$$

4.5 Complément : autres trajectoires envisageables (MPSI)

a) Détermination de la nature de la trajectoire par une méthode numérique

Dans ce paragraphe, on résout numériquement les équations du mouvement pour déterminer les différentes trajectoires possibles.

Équations sans dimension Pour résoudre numériquement le système d'équations du mouvement (20.2), on le rend adimensionnel en utilisant la longueur caractéristique $r_0 = \frac{\mathcal{C}^2}{\mathcal{G}M_A}$ et l'énergie caractéristique $E_0 = \frac{\mathcal{G}M_A m}{2r_0}$ définies au paragraphe 3.2. À partir de r_0 et de la constante des aires, on détermine une vitesse caractéristique $V = \frac{\mathcal{C}}{r_0}$ et un temps caractéristique $\tau = \frac{r_0}{V} = \frac{r_0^2}{\mathcal{C}}$. On pose alors $t^* = t/\tau$, $r^* = r/r_0$ et les équations du mouvement deviennent :

$$\mathcal{C} = \frac{r_0^2}{\tau} r^{*2} \frac{d\theta}{dt^*} \quad \text{et} \quad E_m = \frac{1}{2} m \frac{r_0^2}{\tau^2} \left(\frac{dr^*}{dt^*} \right)^2 + E_0 E_{p_{eff}}^*(r^*).$$

L'énergie potentielle effective adimensionnée $E_{p_{eff}}^*(r^*) = \frac{1}{r^{*2}} - \frac{2}{r^*}$ a été explicitée au paragraphe 3.2 et tracée sur la figure 20.4.

Comme $r_0 = \frac{\mathcal{C}^2}{\mathcal{G}M_A} \Leftrightarrow \frac{\mathcal{C}^2}{r_0} = \mathcal{G}M_A$, on calcule $\frac{1}{2} m \frac{r_0^2}{\tau^2} = \frac{1}{2} m V^2 = \frac{1}{2} m \frac{\mathcal{C}^2}{r_0^2} = \frac{1}{2} m \frac{\mathcal{G}M_A}{r_0} = E_0$.

On pose alors $\alpha = \frac{E_m}{E_0}$ et on obtient les équations adimensionnées suivantes :

$$\frac{d\theta}{dt^*} = \frac{1}{r^{*2}} \quad \text{et} \quad \left(\frac{dr^*}{dt^*} \right)^2 = \alpha - \frac{1}{r^{*2}} + \frac{1}{r^*}.$$

Pour des raisons de techniques numériques, il est plus facile de dériver la deuxième équation et d'effectuer la résolution numérique du système :

$$\frac{d\theta}{dt^*} = \frac{1}{r^{*2}} \quad \text{et} \quad \frac{d^2 r^*}{dt^{*2}} = \frac{1}{r^{*3}} - \frac{1}{r^{*2}}.$$

Conditions initiales et domaine de variation des paramètres La courbe d'énergie potentielle effective de la figure 20.4 permet de limiter l'intervalle accessible à l'énergie mécanique E_m à $[-E_0, +\infty[$. α évolue donc dans l'intervalle $[-1, +\infty[$. Le point où M s'approche le plus de O permet d'établir les conditions initiales du mouvement puisqu'en ce point $\frac{dr^*}{dt^*} = 0$

et r^* vérifie $\alpha - \frac{1}{r^{*2}} + \frac{1}{r^*} = 0 \Rightarrow \alpha r^{*2} + 2r^* - 1 = 0$. On cherche la plus petite solution strictement positive de cette équation qui correspond la distance minimale séparant M de O . Lorsque $\alpha = 0$, la solution est évidemment $r^* = 0,5$.

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

Lorsque $\alpha \neq 0$, le discriminant $\Delta = 4 + 4\alpha$ est positif car $\alpha > -1$ et la plus petite solution positive est $r^* = \frac{-1 + \sqrt{1 + \alpha}}{\alpha}$. Les conditions initiales sont donc :

$$\begin{cases} \frac{dr^*}{dt^*} = 0 \\ r^* = \frac{-1 + \sqrt{1 + \alpha}}{\alpha} \\ \theta = 0 \end{cases} \quad \text{si } \alpha \neq 0 \quad \text{et} \quad \begin{cases} \frac{dr^*}{dt^*} = 0 \\ r^* = 0,5 \\ \theta = 0 \end{cases} \quad \text{si } \alpha = 0.$$

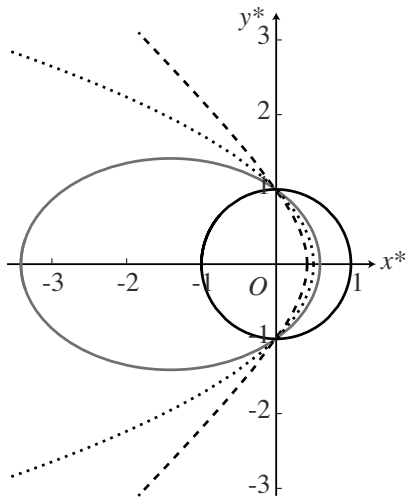


Figure 20.9 – Nature de la trajectoire en fonction de α . La trajectoire est circulaire pour $\alpha = -1$ (trait continu noir), elliptique pour $\alpha = -0,5$ (trait continu gris), parabolique pour $\alpha = 0$ (pointillés) et hyperbolique pour $\alpha = 1$ (tirets).

Tracé et analyse des solutions obtenues en fonction de l'énergie mécanique On trace alors les trajectoires pour α égal à -1, -0,5 et 1 qui correspondent aux énergies mécaniques tracées sur la figure 20.4 et pour une énergie mécanique nulle, c'est-à-dire $\alpha = 0$.

Les courbes obtenues sont des coniques de foyer O dont la nature dépend de la valeur de α :

- si $\alpha = -1 \Leftrightarrow E_m = E_{\min}$, la trajectoire est un cercle de centre O ,
- si $\alpha \in]-1, 0[\Leftrightarrow 0 > E_m > E_{\min}$, la trajectoire est une ellipse de foyer O ,
- si $\alpha = 0 \Leftrightarrow E_m = 0$, la trajectoire est une parabole de foyer O ,
- si $\alpha > 0 \Leftrightarrow E_m > 0$, la trajectoire est une branche d'hyperbole de foyer O .

La valeur de α détermine entièrement la nature de la trajectoire. Les paramètres dimensionnés $\mathcal{G}$, M_A , m , $\mathcal{C}$ et E_m n'interviennent pas indépendamment mais uniquement sous la forme de la combinaison $\alpha = \frac{E_m}{E_0} = \frac{2\mathcal{C}^2 E_m}{\mathcal{G}^2 M_A^2 m}$.

Il faut avoir conscience que la forme de la trajectoire dépend de la nature de la force et que les résultats obtenus ici ne sont valables que pour les forces newtoniennes en $\frac{1}{r^2}$.

b) Expression de l'énergie mécanique pour les trajectoires elliptiques

Dans ce paragraphe, on démontre l'expression de l'énergie mécanique sur une trajectoire elliptique que l'on a admis au paragraphe 4.2. Pour cela, on utilise la courbe d'énergie potentielle effective vue au paragraphe 3 et reproduite figure 20.10. Les trajectoires liées s'effectuent entre les distances r_1 et r_2 . La trajectoire étant une ellipse, on remarque sur le schéma de la figure 20.11 que les paramètres r_1 et r_2 sont liés à la longueur du grand axe de l'ellipse :

$$2a = r_1 + r_2.$$

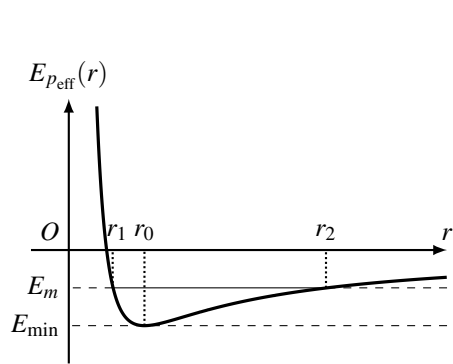


Figure 20.10 – Courbe d'énergie potentielle effective.

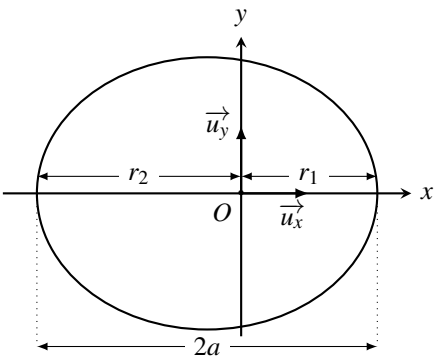


Figure 20.11 – Ellipse montrant r_1 , r_2 et le grand-axe $2a$.

Or r_1 et r_2 sont solutions de l'équation :

$$E_m = E_{\text{eff}}(r) = \frac{1}{2}m\frac{\mathcal{C}^2}{r^2} - \mathcal{G}\frac{mM_A}{r} \implies r^2 + \frac{\mathcal{G}mM_A}{E_m}r - \frac{1}{2}m\frac{\mathcal{C}^2}{E_m} = 0.$$

Cette équation s'écrit également : $(r - r_1)(r - r_2) = r^2 - (r_1 + r_2)r + r_1r_2 = r^2 - Sr + P = 0$, où $S = r_1 + r_2$ représente la somme des racines du polynôme de deuxième ordre et $P = r_1r_2$ son produit. On en déduit que :

$$r_1 + r_2 = 2a = -\frac{\mathcal{G}mM_A}{E_m}, \quad \text{soit} \quad E_m = -\frac{\mathcal{G}M_Am}{2a}.$$

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

SYNTHÈSE

SAVOIRS

- champ de force centrale conservatif
- un mouvement à moment cinétique constant est plan et vérifie la loi des aires
- champ newtonien
- lois de J. Kepler pour les planètes
- définition d'un satellite géostationnaire
- vitesses cosmiques : vitesse en orbite basse et vitesse de libération
- ordre de grandeur des vitesses cosmiques en dynamique terrestre

SAVOIR-FAIRE

- montrer que le moment cinétique est conservé lors d'un mouvement dans un champ de force centrale
- exprimer la conservation de l'énergie mécanique
- construire une énergie potentielle effective pour étudier qualitativement le mouvement radial selon la valeur de l'énergie mécanique
- transposer les lois de J. Kepler au cas des satellites terrestres
- étudier le mouvement circulaire d'une planète autour du Soleil
- établir la troisième loi de J. Kepler pour une planète en révolution circulaire
- exploiter les lois de J. Kepler dans le cas d'une trajectoire elliptique
- calculer l'altitude d'un satellite géostationnaire et justifier que son orbite est située dans le plan équatorial
- exprimer l'énergie mécanique en fonction du rayon de l'orbite dans le cas d'un mouvement circulaire
- exprimer l'énergie mécanique en fonction du demi-grand axe dans le cas d'un mouvement elliptique (MPSI)
- exprimer les vitesses en orbite basse et vitesses de libération (vitesses cosmiques)

MOTS-CLÉS

- | | | |
|---------------------|------------------|-----------------------------|
| • force newtonienne | • loi des aires | • satellite géostationnaire |
| • force centrale | • lois de Kepler | • vitesses cosmiques |

S'ENTRAÎNER

Les données numériques suivantes sont utilisées dans de nombreux exercices :

- Masse de la Terre : $M_T = 5,97.10^{24}$ kg.
- Rayon de la Terre : $R_T = 6,38.10^6$ m.
- Masse du Soleil : $M_S = 1,99.10^{30}$ kg.
- Constante universelle de gravitation : $\mathcal{G} = 6,67.10^{-11}$ N·m²·kg⁻².
- Valeur du champ de pesanteur au niveau du sol : $g = 9,81$ m·s⁻².
- Unité astronomique (distance moyenne Terre-Soleil) : 1 ua = $1,5.10^{11}$ m.

20.1 Paramètre gravitationnel standard de la Terre (★)

1. Montrer que la connaissance de la période d'un satellite terrestre en orbite circulaire à l'altitude z et de la troisième loi de Kepler permet de déterminer $\mu = \mathcal{G}M_T$, le produit de la constante de gravitation universelle par la masse de la Terre.
2. L'expérience de Cavendish permet de déterminer la valeur de $\mathcal{G} = 6,67.10^{-11}$ N·m²·kg⁻². Montrer qu'on peut alors déterminer la masse de la Terre.
3. On rappelle qu'un satellite géostationnaire est en orbite à une altitude approximative de 36.10^3 km et que le rayon de la Terre vaut $R_T = 6,4.10^3$ km. Calculer le paramètre gravitationnel standard de la Terre, sa masse puis sa masse volumique moyenne.
4. Donner une autre méthode pour déterminer μ en utilisant la valeur de champ de gravité terrestre à la surface du globe g .

20.2 Vitesse d'un satellite à son périégée (MPSI) (★)

Lors de son lancement, le satellite d'observation Hipparcos est resté sur son orbite de transfert à cause d'un problème technique. On l'assimile à un point matériel M de masse $m = 1,1$ t. L'orbite de transfert est elliptique et la distance Terre-satellite varie entre $d_P = 200$ km au périégée et $d_A = 35,9.10^3$ km à l'apogée. On rappelle que le périégée est le point de l'orbite le plus proche de la Terre et que l'apogée est le point le plus éloigné. On mesure la vitesse du satellite à son apogée : $v_A = 3,5.10^2$ m·s⁻¹.

1. Faire un schéma de la trajectoire en faisant apparaître la position O du centre de la Terre, l'apogée A et le périégée P .
2. Déterminer le demi-grand axe a de la trajectoire.
3. En déduire l'énergie mécanique et la période du satellite.
4. On note v_A et v_P les vitesses du satellite en A et en P . Exprimer le module moment cinétique calculé au point O du satellite à son apogée puis à son périégée.
5. En déduire la vitesse du satellite à son périégée.

20.3 Energie nécessaire pour mettre un satellite artificiel en orbite (★)

On étudie le mouvement d'un satellite de masse m en orbite circulaire à une altitude z autour de la Terre, ainsi que le lancement d'un satellite artificiel à partir d'un point O de la surface terrestre.

1. Dans quel référentiel se place-t-on pour étudier le mouvement d'un satellite terrestre ?

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

2. Energie d'un satellite artificiel en orbite

- Rappeler l'expression de l'énergie potentielle gravitationnelle du satellite en orbite à une distance r du centre de la Terre. En déduire son expression en fonction de son altitude z .
- Retrouver l'expression de la vitesse en orbite à une altitude z .
- En déduire l'expression de l'énergie cinétique puis de l'énergie mécanique E_m du satellite sur son orbite à l'altitude z .
- Calculer cette énergie mécanique pour $z = 1,0 \cdot 10^3$ km et $m = 6,0$ tonnes.

3. Energie nécessaire au lancement d'un satellite

Pour lancer un satellite, il faut lui communiquer l'énergie $\Delta E_m = E_m - E_{m0}$ où E_{m0} est l'énergie qu'il a au point O .

- Dans le référentiel géocentrique, la Terre peut être assimilée à un solide en rotation autour d'un axe à une vitesse angulaire Ω . Préciser l'axe de rotation. Est-il fixe ? Que vaut la vitesse angulaire ?
- En déduire l'expression de la vitesse du point O dans le référentiel géocentrique $\mathcal{R}_g$ supposé galiléen en fonction de Ω , du rayon terrestre R_T et de la latitude du lieu λ .
- Exprimer alors l'énergie mécanique initiale E_{m0} du satellite posé au sol au point O .
- En déduire les conditions les plus favorables pour le lancement du satellite. Parmi les trois champs de tirs suivants, lequel choisir de préférence ?
 - Baïkonour au Kazakhstan : $\lambda = 46^\circ$;
 - Cap Canaveral aux USA : $\lambda = 28,5^\circ$;
 - Kourou en Guyane française : $\lambda = 5,23^\circ$.
- Calculer l'énergie nécessaire pour mettre le satellite en orbite basse depuis Kourou.
- Calculer numériquement l'énergie gagnée entre Baïkonour et Kourou. Commenter.

20.4 Trajectoire quasi-circulaire d'un satellite - Freinage par l'atmosphère (★★)

On étudie le mouvement d'un satellite artificiel de la Terre dans le référentiel géocentrique supposé galiléen. On néglige les autres interactions que la force de gravitation entre la Terre et le satellite. On note M_T la masse de la Terre, R_T son rayon, m la masse du satellite supposée petite devant M_T et G la constante de gravitation universelle. On note T_0 la période de révolution du satellite.

- Etablir la conservation du moment cinétique du satellite par rapport à la Terre.
- En déduire que le mouvement du satellite est plan.
- Montrer que cela permet de définir une constante des aires C dont on donnera l'expression.
- On suppose que le satellite est en orbite circulaire autour de la Terre. Montrer que son mouvement est uniforme.
- Etablir l'expression de la vitesse v du satellite en fonction de G , M_T et r le rayon de l'orbite du satellite.
- Retrouver la troisième loi de Kepler dans le cas d'une trajectoire circulaire. Faire l'application numérique pour un satellite situé sur une orbite basse à $1,0 \cdot 10^3$ km d'altitude.
- Déterminer l'expression de l'énergie cinétique E_c du satellite en fonction de $\mathcal{G}$, M_T , m et r .
- Même question pour l'énergie potentielle E_p du satellite. Donner la relation entre E_c et E_p .

9. En déduire l'expression de l'énergie mécanique E_m et les relations de E_m avec E_p et E_c .
Les satellites en orbite basses subissent des frottements de la part des hautes couches de l'atmosphère. Ces frottements limitent la durée de vie des satellites en les faisant lentement chuter sur Terre. Un satellite situé sur une orbite à $1,0 \cdot 10^3$ km kilomètres d'altitude descend d'environ 2 m par jour. On cherche à modéliser ces observations.

On modélise l'action des hautes couches de l'atmosphère par une force de frottement proportionnelle à la masse du satellite et sa vitesse au carré : $\vec{f} = -\alpha m v \vec{v}$ où α est un coefficient de frottement. Cette force est suffisamment faible pour que la trajectoire soit quasi-circulaire. Dans ces conditions, les expressions des différentes énergies en fonction de r restent valables mais r varie lentement dans le temps.

10. A l'aide du théorème de l'énergie cinétique, établir l'équation différentielle vérifiée par r .

11. Sans résoudre, montrer que r ne peut que diminuer.

12. En déduire un résultat surprenant sur l'évolution de la vitesse du satellite.

20.5 Modèle de bohr de l'atome d'hydrogène (★★)

Pour expliquer le spectre de raies de l'atome d'hydrogène observées expérimentalement, N. Bohr a proposé un modèle qui s'appuie sur les hypothèses suivantes : dans un référentiel galiléen lié au noyau O ,

- i) l'électron décrit une trajectoire circulaire sur laquelle il ne rayonne pas d'énergie ;
- ii) l'électron échange de l'énergie avec l'extérieur sous forme de lumière lorsqu'il change de trajectoire circulaire ;
- iii) le module du moment cinétique de l'électron est quantifié et ne peut prendre que des valeurs discrètes vérifiant la relation :

$$L_{O_n} = n \frac{h}{2\pi}$$

où n un nombre entier naturel non nul et h la constante de Planck.

Une orbitale électronique correspond à une valeur de l'entier n . Elle est caractérisée par un rayon r_n , une vitesse v_n et une énergie mécanique $E_m(n)$.

Ce modèle semi-classique n'est pas complètement satisfaisant, mais il prédit le spectre de raies de l'atome d'hydrogène. A ce titre, il a eu son heure de gloire et a permis de banaliser l'idée que la quantification des grandeurs physiques est nécessaire à l'échelle atomique.

On rappelle qu'un atome d'hydrogène est constitué d'un noyau (charge e , masse m_p) et d'un électron (charge e , masse m_e), et on donne les valeurs numériques utiles pour cet exercice : masse de l'électron : $m_e = 0,911 \cdot 10^{-30}$ kg ; charge du proton $e = 1,602 \cdot 10^{19}$ C ; constante de Planck : $h = 6,63 \cdot 10^{-34}$ J·s ; célérité de la lumière dans le vide : $c = 3,00 \cdot 10^8$ m·s⁻¹ ; permittivité diélectrique du vide $\epsilon_0 = 8,85 \cdot 10^{-12}$ F·m⁻¹.

1. Rappeler l'expression de la force d'interaction exercée par le noyau sur l'électron et de l'énergie potentielle dont elle dérive.

2. Utiliser le fait que les orbitales sont circulaires pour exprimer le carré v_n^2 de la vitesse de l'électron en fonction de la distance r_n .

3. Utiliser la quantification du moment cinétique pour exprimer le rayon de la trajectoire en fonction de n , h , m_e , e .

4. Calculer sa valeur pour $n = 1$.

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

5. Exprimer l'énergie mécanique $E_m(n)$ de l'électron et montrer qu'elle se met sous la forme $E_m(n) = -\frac{A}{n^2}$. Donner l'expression et la valeur numérique de A , en électron-volts (on rappelle que $1 \text{ eV} = 1,6 \cdot 10^{-19} \text{ J}$).

6. Sachant que le passage d'un niveau d'énergie $E_m(n_1)$ à un autre $E_m(n_2)$ se traduit par l'émission d'un photon de fréquence ν telle que $\Delta E = h\nu$, en déduire que les longueurs d'onde λ émises vérifient :

$$\frac{1}{\lambda} = R_H \left(\frac{1}{n_2^2} - \frac{1}{n_1^2} \right).$$

On rappelle que $\nu = \frac{c}{\lambda}$ où c désigne la vitesse de la lumière dans le vide. On donnera l'expression de la constante de Rydberg R_H ainsi que sa valeur numérique.

7. Quelles sont les longueurs d'onde dans le visible pour les séries de Lyman ($n_2 = 1$) et de Balmer ($n_2 = 2$) ?

APPROFONDIR

20.6 Caractéristique d'un météorite (★★★)

On repère un météorite très éloigné du Soleil et on mesure sa vitesse $\vec{v}_0$. On observe que $\vec{v}_0 = v_0 \vec{u}_x$ est portée par une droite Δ qui est située à une distance b du centre O du Soleil. On suppose qu'à l'instant où on la repère (instant initial), le météorite est si éloigné que son énergie potentielle d'interaction gravitationnelle avec le Soleil est négligeable.

On note m la masse du météorite, M_S celle du Soleil, R_S le rayon du Soleil et $\mathcal{G}$ la constante de gravitation universelle.

1. Faites un schéma de la situation initiale.
2. Montrer que l'énergie mécanique du météorite est une intégrale première du mouvement. Déterminer sa valeur initiale.
3. Montrer que le moment cinétique par rapport à O du météorite est une intégrale première du mouvement. Déterminer sa valeur initiale.
4. Rappeler les conséquences de la conservation du moment cinétique.
5. Définir les coordonnées polaires adaptés et établir l'expression du moment cinétique à un instant t quelconque.
6. Etablir l'expression de l'énergie potentielle effective.
7. Exprimer cette énergie potentielle effective en fonction de m , v_0 , b , du produit $\mathcal{G}M_S$ et de la distance r .
8. Tracer l'allure de la courbe d'énergie potentielle effective. et en déduire la nature bornée ou non de la trajectoire du météorite.
9. Déterminer la distance minimale d'approche $r_{\min}$ en fonction de v_0 , b et du produit $\mathcal{G}M_S$.
10. A quelle condition sur $r_{\min}$ le météorite n'ira pas toucher la surface du Soleil ?

20.7 Changement d'orbite d'un satellite (MPSI) (★★★)

On souhaite transférer un satellite depuis une orbite circulaire rasante de rayon R_T autour de la Terre sur son orbite géostationnaire de rayon R_G . On fera l'étude dans le référentiel

géocentrique dans lequel la Terre tourne sur elle-même à la vitesse angulaire Ω . On néglige les autres interactions que la force de gravitation entre la Terre et le satellite. On note O le centre de la Terre, $M_T = 5,98 \cdot 10^{24}$ kg sa masse, $R_T = 6,37 \cdot 10^6$ m son rayon, $m = 1,5$ t la masse du satellite et G la constante de gravitation universelle.

1. Etablir que la trajectoire du satellite géostationnaire est forcément dans le plan équatorial.
 2. En déduire que les trois orbites appartiennent à un même plan à préciser.
 3. Déterminer la vitesse v_B du satellite sur son orbite basse avant son transfert sur son orbite géostationnaire.
 4. Donner sa valeur numérique.
 5. Exprimer le rayon de la trajectoire géostationnaire.
 6. Donner sa valeur numérique.
 7. Préciser l'altitude de l'orbite géostationnaire.
 8. Calculer la valeur numérique de la vitesse du satellite sur son orbite géostationnaire.
- Le transfert du satellite de son orbite basse à son orbite géostationnaire s'effectue de la manière suivante : on communique au satellite une brusque variation de vitesse en un point P de sa trajectoire basse en éjectant des gaz pendant un intervalle de temps très court dans le sens opposé à la vitesse du satellite. Il suit alors une orbite elliptique et lorsque sa trajectoire croise la droite OP au point A , on lui communique un supplément de vitesse pour le stabiliser sur l'orbite géostationnaire.*
9. Etablir une relation entre les vitesses aux points A et P et les distances $r_A = OA$ et $r_P = OP$.
 10. En utilisant la conservation de l'énergie mécanique sur la trajectoire elliptique, établir l'expression de l'énergie mécanique en fonction de G , M_T , m et a le demi grand axe de l'ellipse.
 11. Donner la valeur numérique de l'énergie mécanique sur l'ellipse.
 12. Etablir l'expression de la vitesse du satellite sur la trajectoire elliptique en fonction de R_G , R_T , r , G et M_T .
 13. Donner la valeur de la variation de vitesse qu'il faut imposer en P .
 14. En déduire la variation de l'énergie mécanique en P .
 15. Donner la valeur de la variation de vitesse qu'il faut imposer en A .
 16. En déduire la variation de l'énergie mécanique en A .
 17. Déterminer la durée de ce transfert.

CORRIGÉS

20.1 Paramètre gravitationnel standard de la Terre

1. La troisième loi de Kepler pour un satellite terrestre en orbite circulaire de rayon r s'écrit

$$\frac{T^2}{r^3} = \frac{4\pi^2}{\mathcal{G}M_T} \Rightarrow \mu = \mathcal{G}M_T = \frac{4\pi^2 r^3}{T^2} = \frac{4\pi^2 (R_T + z)^3}{T^2},$$

où R_T désigne le rayon de la Terre.

2. La connaissance de $\mathcal{G}$ donne accès à $M_T = \frac{\mu}{\mathcal{G}}$.

3. Numériquement : $\mu = \frac{4 \times 3,14^2 \times (42 \cdot 10^6)^3}{86400^2} = 39 \cdot 10^{13} \text{ m}^3 \cdot \text{s}^{-2}$. On en déduit la masse de

la Terre $M_T = 5,9 \cdot 10^{24} \text{ kg}$ puis sa masse volumique moyenne $\rho = \frac{M_T}{\frac{4}{3}\pi R_T^3} = 5,3 \cdot 10^3 \text{ kg} \cdot \text{m}^{-3}$.

A l'époque de Cavendish, ce résultat a surpris car la masse volumique des roches les plus denses est l'ordre de $3 \cdot 10^3$ à $4 \cdot 10^3 \text{ kg} \cdot \text{m}^{-3}$. On sait maintenant que cette valeur est due au noyau terrestre constitué de fer beaucoup plus dense.

4. On peut également trouver μ en exploitant $g = 9,8 \text{ m} \cdot \text{s}^{-2} = \frac{\mathcal{G}M_T}{R_T^2}$ soit :

$$\mu = \mathcal{G}M_T = gR_T^2 = 9,8 \times (6,4 \cdot 10^6)^2 = 40 \cdot 10^{13} \text{ m}^3 \cdot \text{s}^{-2}.$$

20.2 Vitesse d'un satellite à son périégée

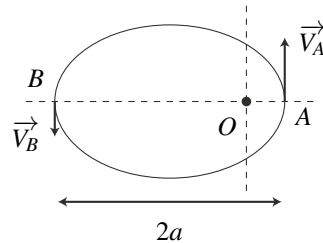
1. Voir ci-contre.

2. $a = \frac{d_P + d_A}{2} = 18,1 \cdot 10^3 \text{ km}$ (voir ci-contre).

3. $E_m = -\frac{\mathcal{G}mM_T}{2a} = -\frac{\mathcal{G}mM_T}{d_P + d_A} = -1,2 \cdot 10^{10} \text{ J}$.

On déduit la période de la troisième loi de Kepler en utilisant un satellite géostationnaire qui effectue une révolution en $T_{\text{geo}} = 1$ jour à l'altitude $z_{\text{geo}} = 36 \cdot 10^3 \text{ km}$ donc pour lequel $a_{\text{geo}} = 42,4 \cdot 10^3 \text{ km}$ pour déterminer la valeur de la constante :

$$T^2/a^3 = T_{\text{geo}}^2/a_{\text{geo}}^3 \Rightarrow T = T_{\text{geo}} \left(\frac{a}{a_{\text{geo}}} \right)^{\frac{3}{2}} = 0,278 \text{ jour} = 6,7 \text{ h}.$$



4. Au points A et P et seulement en ces points, on a $\vec{v}_A \perp \vec{OA}$ (respectivement $\vec{v}_P \perp \vec{OP}$) et $L_T(A) = md_A v_A$ (resp. $L_T(P) = md_P v_P$).

5. Le mouvement est à force centrale de centre O, le moment cinétique en O est une constante du mouvement d'où $d_A v_A = d_P v_P$ puis $v_P = v_A \frac{d_A}{d_P} = 3,5 \cdot 10^2 \times \frac{35,9}{0,2} = 6,3 \cdot 10^4 \text{ m} \cdot \text{s}^{-1}$.

20.3 Energie nécessaire pour mettre un satellite artificiel en orbite

1. Le référentiel d'étude de ce type de mouvements est le référentiel géocentrique.
2. Energie d'un satellite artificiel en orbite

a. $E_p = -\frac{\mathcal{G}mM_T}{r} = -\frac{\mathcal{G}mM_T}{R_T + z}$ où R_T et M_T sont le rayon et la masse de la Terre.

b. Le mouvement circulaire du satellite est étudié en coordonnées polaires dont l'origine est confondue avec le centre de la Terre. Le satellite étudié n'est soumis qu'à l'attraction gravitationnelle de la Terre. On applique le principe fondamental de la dynamique au satellite dans le référentiel géocentrique galiléen en projection sur $\vec{u}_r$:

$$-m\frac{v^2}{r} = -\frac{\mathcal{G}mM_T}{r^2} \Rightarrow v = \sqrt{\frac{\mathcal{G}M_T}{r}} = \sqrt{\frac{\mathcal{G}M_T}{R_T + z}}.$$

c. On trouve $E_c = \frac{1}{2}mv^2 = \frac{1}{2}\frac{\mathcal{G}mM_T}{R_T + z}$ puis $E_m = E_c + E_p = -\frac{1}{2}\frac{\mathcal{G}mM_T}{R_T + z}$. Etant donné que l'énergie potentielle du satellite est en forme de puits dont le maximum est égal à 0 et que le satellite est lié, son énergie mécanique est forcément négative, ce que l'on observe ici.

d. Numériquement : $E_m = -0,5 \frac{6,67 \cdot 10^{-11} \times 6 \cdot 10^3 \times 5,97 \cdot 10^{24}}{7,38 \cdot 10^6} = -1,6 \cdot 10^{11} \text{ J}.$

3. Energie nécessaire au lancement d'un satellite

a. Dans le référentiel géocentrique, la Terre est en rotation autour d'un axe fixe (l'axe de ses pôles) à la vitesse angulaire $\Omega = \frac{2\pi}{T} = 7,3 \cdot 10^{-5} \text{ rad} \cdot \text{s}^{-1}$ en prenant $T = 1 \text{ jour}$.

b. Le point O est en mouvement circulaire de centre H , projeté de O sur l'axe de rotation, de rayon $R = HO = R_T \cos(\lambda)$, à la vitesse angulaire Ω . Sa vitesse est donc $V_O = R\Omega = R_T \cos(\lambda)\Omega$.

c. L'énergie E_{m0} vaut donc :

$$E_{m0} = \frac{1}{2}mR_T^2 \cos^2(\lambda)\Omega^2 - \frac{\mathcal{G}mM_T}{R_T}.$$

d. Plus V_O sera grand, plus l'énergie cinétique au sol du satellite sera grande. Il faut donc maximiser $\cos \lambda$ donc se situer le plus près possible de l'équateur. Le champ de tir de Kourou est donc le site à privilégier.

e. L'énergie à fournir sera la différence $\Delta E_m = E_m - E_{m0}$. On calcule numériquement $E_{m0} = -3,7 \cdot 10^{11} \text{ J}$ puis $\Delta E_m = 2,1 \cdot 10^{11} \text{ J}$.

f. Entre Baïkonour et Kourou, on gagne $E_{m0} = \frac{1}{2}mR_T^2\Omega^2(\cos^2(\lambda_{\text{Kourou}}) - \cos^2(\lambda_{\text{Baïkonour}}))$. Numériquement on trouve une économie de $3,3 \cdot 10^8 \text{ J}$. Le gain d'énergie relatif de l'ordre de 0,15% est relativement faible, mais comme chaque tir nécessite plusieurs centaines de tonnes de carburant, cela n'est pas tout à fait négligeable.

20.4 Trajectoire quasi-circulaire d'un satellite - Freinage par l'atmosphère

On étudie le satellite dans le référentiel géocentrique supposé galiléen. Dans ce référentiel, le satellite n'est soumis qu'à la force de gravitation $\vec{F}_G$ qui est une force centrale de centre O confondu avec le centre de la Terre.

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

1. La résultante des forces est colinéaire au vecteur position. Son moment $\vec{\mathcal{M}}_O$ par rapport à O est nul. Le théorème du moment cinétique implique que $\frac{d\vec{L}_O}{dt} = \vec{\mathcal{M}}_O = \vec{0}$. Par conséquent que $\vec{L}_O$ est une constante du mouvement.

2. Si la valeur du moment cinétique est nulle, le mouvement est rectiligne. Sinon, on pose $\vec{u}_z = \frac{\vec{L}_O}{\|\vec{L}_O\|}$ et comme $\vec{OM}$ est perpendiculaire à $\vec{L}_O = \vec{OM} \wedge m\vec{v}$ (propriété du produit vectoriel), il est perpendiculaire à $\vec{u}_z$. Le point M se déplace donc dans le plan xOy .

3. Dans ce plan, on utilise les coordonnées polaires et on a : $\vec{OM} = r\vec{u}_r$ et $\vec{v} = \dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta$. On en déduit l'expression du moment cinétique en O : $\vec{L}_O = \vec{OM} \wedge m\vec{v} = r\vec{u}_r \wedge m(\dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta) = mr^2\dot{\theta}\vec{u}_z$. On peut définir la constante des aires par $C = r^2\dot{\theta}$.

4. Le mouvement étant circulaire, r est une constante. $C = r^2\dot{\theta}$ est aussi une constante donc $\dot{\theta}$ l'est également. Le mouvement est circulaire et uniforme.

5. Pour un mouvement circulaire et uniforme, l'étude cinématique donne $\vec{OM} = r\vec{u}_r$, $\vec{v} = r\dot{\theta}\vec{u}_\theta$ et $\vec{a} = -\frac{v^2}{r}\vec{u}_r$. La projection du principe fondamental de la dynamique sur $\vec{u}_r$ donne $-m\frac{v^2}{r} = -\frac{GM_T m}{r^2}$ soit $v = \sqrt{\frac{GM_T}{r}}$.

6. La vitesse étant constante, $v = \frac{2\pi r}{T} = \sqrt{\frac{GM_T}{r}}$ puis $\frac{T^2}{r^3} = \frac{4\pi^2}{GM_T}$ (troisième loi de Kepler).

7. L'énergie cinétique vaut $E_c = \frac{1}{2}mv^2 = \frac{GmM_T}{2r}$.

8. Quant à l'énergie potentielle, elle vaut $E_p = -\frac{GmM_T}{r} = -2E_c$.

9. On en déduit l'énergie mécanique $E_m = E_c + E_p = -\frac{GmM_T}{2r} = -E_c = \frac{E_p}{2}$.

10. L'application du théorème de l'énergie cinétique au satellite donne :

$$\frac{dE_c}{dt} = \mathcal{P}(\vec{F}_G) + \mathcal{P}(\vec{f}) = -\frac{dE_p}{dt} + \mathcal{P}(\vec{f}) \Rightarrow \frac{dE_m}{dt} = \mathcal{P}(\vec{f}) = \vec{f} \cdot \vec{v}.$$

Or $\frac{dE_m}{dt} = \frac{GmM_T}{2r^2} \frac{dr}{dt}$ et $\mathcal{P}(\vec{f}) = \vec{f} \cdot \vec{v} = -\alpha mv^3 = -\alpha m \left(\frac{GM_T}{r} \right)^{\frac{3}{2}}$. En combinant ces deux équation on aboutit à :

$$\frac{dr}{dt} = -2\alpha (GM_T)^{\frac{1}{2}} r^{\frac{1}{2}}.$$

11. On aboutit à $\dot{r} < 0$. r diminue au cours du temps. Le satellite tombe sur la Terre.

12. Au cours de sa chute, l'énergie cinétique $E_c = \frac{GmM_T}{2r}$ augmente car $\frac{1}{r}$ est une fonction décroissante. La vitesse du satellite augmente. L'existence d'une force de frottement provoque une augmentation de la vitesse, ce qui est contre intuitif.

20.5 Modèle de bohr de l'atome d'hydrogène

On étudie le mouvement de l'électron dans le référentiel lié au noyau supposé galiléen. L'électron est animé d'un mouvement circulaire et uniforme de centre O . On l'étudie en coordonnées polaires d'origine O . Dans ce système de coordonnées :

$$\overrightarrow{OM} = r\overrightarrow{u_r} \quad \text{puis} \quad \overrightarrow{v} = r\dot{\theta}\overrightarrow{u_r} \quad \text{et} \quad \overrightarrow{a} = -\frac{v^2}{r}\overrightarrow{u_r}$$

où r est la distance séparant le noyau de l'électron.

1. La seule force dont il faut tenir compte est la force d'interaction coulombienne qu'exerce le noyau de charge $+e$ sur l'électron de charge $-e$: $\overrightarrow{f} = -\frac{e^2}{4\pi\epsilon_0 r^2}\overrightarrow{u_r}$ qui dérive de l'énergie potentielle $E_p = -\frac{e^2}{4\pi\epsilon_0 r}$. On peut négliger l'interaction gravitationnelle.

2. On projette le principe fondamental de la dynamique sur le vecteur $\overrightarrow{u_r}$ et on obtient :

$$-m_e \frac{v_n^2}{r_n} = -\frac{e^2}{4\pi\epsilon_0 r_n^2} \quad \Rightarrow \quad v_n^2 = \frac{e^2}{4\pi\epsilon_0 m_e r_n}.$$

3. Le module du moment cinétique s'écrit : $L_{O_n} = m_e r_n v_n = \frac{nh}{2\pi}$. On écrit alors le carré de la vitesse de deux manière et on obtient : $v^2 = \frac{e^2}{4\pi\epsilon_0 m_e r} = \frac{n^2 h^2}{4\pi^2 m_e^2 r^2}$ soit $r = \frac{n^2 h^2 \epsilon_0}{\pi m_e e^2}$.

4. Application numérique pour $n = 1$: $r = 53 \text{ pm}$.

5. L'énergie mécanique vaut : $E_m(n) = Ec + Ep = \frac{1}{2}m_e v_n^2 - \frac{e^2}{4\pi\epsilon_0 r_n} = -\frac{e^2}{8\pi\epsilon_0 r_n}$. De plus, $r = \frac{n^2 h^2 \epsilon_0}{\pi m_e e^2}$ donc :

$$E_m(n) = -\frac{e^4 m_e}{8\epsilon_0^2 h^2} \frac{1}{n^2} = -\frac{A}{n^2} \quad \text{avec} \quad A = \frac{e^4 m_e}{8\epsilon_0^2 h^2} = 2,17 \cdot 10^{-18} \text{ J} = 13,6 \text{ eV}.$$

$$6. \Delta E = A \left(\frac{1}{n_2^2} - \frac{1}{n_1^2} \right) = h\nu = h\frac{c}{\lambda} \quad \text{donc} \quad \frac{1}{\lambda} = \frac{A}{hc} \left(\frac{1}{n_2^2} - \frac{1}{n_1^2} \right) = R_H \left(\frac{1}{n_2^2} - \frac{1}{n_1^2} \right),$$

$$\text{avec } R_H = \frac{A}{hc} = \frac{e^4 m_e}{8\epsilon_0^2 ch^3} = 1,09 \cdot 10^7 \text{ m}^{-1}.$$

7. Les longueurs d'onde visibles sont comprises entre 400 et 800 nm environ.

Pour la série de Lyman ($n_2 = 1$), on a : pour $n_1 = 2$, $\lambda = 122 \text{ nm}$; pour $n_1 = 3$, $\lambda = 103 \text{ nm}$ et pour toutes les valeurs supérieures de n_1 les longueurs d'onde seront plus faibles. On n'a donc aucune raie dans le domaine visible, elles sont toutes dans le domaine ultra-violet.

Pour la série de Balmer ($n_2 = 2$), on a : pour $n_1 = 3$, $\lambda = 660 \text{ nm}$; pour $n_1 = 4$, $\lambda = 489 \text{ nm}$; pour $n_1 = 5$, $\lambda = 436 \text{ nm}$; pour $n_1 = 6$, $\lambda = 412 \text{ nm}$; pour $n_1 = 7$, $\lambda = 399 \text{ nm}$ et pour toutes les valeurs supérieures de n_1 les longueurs d'onde seront plus faibles. On a donc quatre raies (pour n_1 égal à 3, 4, 5 et 6) dans le domaine visible.

Ces raies et leurs longueurs d'ondes prédites par le modèle correspondent à celles que l'on mesure expérimentalement à l'aide d'un spectroscopie.

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

20.6 Caractéristique d'un météorite

1. On étudie le mouvement du météorite dans le référentiel héliocentrique supposé galiléen. Ce météorite est soumis à la seule force d'attraction gravitationnelle du Soleil qui est une force newtonienne donc centrale et conservative en $\frac{1}{r^2}$. On réalise le schéma suivant :

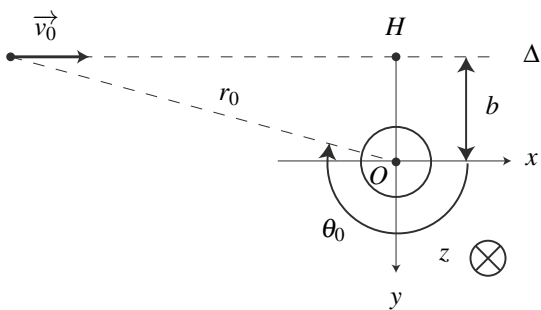


Figure 20.12

2. On applique le théorème de l'énergie mécanique au météorite soumis à l'attraction gravitationnelle du Soleil qui est une force conservative. On obtient : $\frac{dE_m}{dt} = 0 \Rightarrow E_m = \text{constante}$.
L'énergie mécanique du météorite vaut : $E_m = E_c + E_p = \frac{1}{2}mv^2 - m\frac{\mathcal{G}M_S}{r}$, où r est la distance séparant le météorite du centre de la Terre. A l'instant initial, cette relation s'écrit :

$$E_m = \frac{1}{2}mv_0^2 - m\frac{\mathcal{G}M_S}{r_0} \simeq \frac{1}{2}mv_0^2.$$

si son énergie potentielle initiale est négligeable devant son énergie cinétique initiale.

3. On applique le théorème du moment cinétique par rapport à O au météorite soumis à l'attraction gravitationnelle du Soleil qui est une force centrale. On obtient :

$$\frac{d\vec{L}_O}{dt} = \vec{0} \Rightarrow \vec{L}_O \text{ est une constante du mouvement.}$$

A l'aide du schéma, on détermine $\vec{L}_O = m\vec{OM} \wedge \vec{v}$ à l'instant initial : $\vec{L}_O = m\vec{OM}_0 \wedge \vec{v}_0 = m(\vec{OH} + \vec{HM}_0) \wedge \vec{v}_0 = m\vec{OM}_0 \wedge \vec{v}_0 = mbv_0\vec{u}_z$, où le point H et le vecteur $\vec{u}_z$ sont définis sur la figure 20.12. Le moment cinétique est alors :

$$\vec{L}_O = mbv_0\vec{u}_z.$$

Il faut faire attention au sens de $\vec{u}_z$: $(\vec{OH}, \vec{v}_0, \vec{u}_z)$ doit être un trièdre direct.

4. Le moment cinétique en O étant conservé, le mouvement est situé dans le plan perpendiculaire à $\vec{L}_O$ c'est-à-dire ici dans le plan (xOy) et vérifie la loi des aires.

5. On définit le système de coordonnées polaire sur la figure 20.12. Dans ces conditions, le mouvement de M étant situé dans le plan xOy , on utilise les relations du cours de cinématique du point :

$$\overrightarrow{OM} = r\vec{u}_r \quad ; \quad \vec{v} = \dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta.$$

On en déduit $\vec{L}_O = m\overrightarrow{OM} \wedge \vec{v} = mr^2\dot{\theta}\vec{u}_z$. On pose $\mathcal{C} = r^2\dot{\theta} = bv_0$ la constante des aires.

6. On écrit ensuite l'énergie mécanique massique à l'instant t :

$$E_m = E_c + E_p = \frac{1}{2}mv^2 - m\frac{\mathcal{G}M_S}{r} = \frac{1}{2}m(\dot{r}^2 + r^2\dot{\theta}^2) - m\frac{\mathcal{G}M_S}{r} = \frac{1}{2}m\dot{r}^2 + E_{p,eff}(r),$$

avec $E_{p,eff}(r) = \frac{1}{2}mr^2\dot{\theta}^2 - m\frac{\mathcal{G}M_S}{r}$. Comme $\mathcal{C} = r^2\dot{\theta} = bv_0$, on a $\dot{\theta}^2 = \left(\frac{bv_0}{r^2}\right)^2$ puis :

$$E_{p,eff}(r) = \frac{1}{2}m\frac{b^2v_0^2}{r^2} - m\frac{\mathcal{G}M_S}{r}.$$

7. On trace alors l'allure de $E_{p,eff}(r)$ qui est la même que dans le cours figure 20.4. L'énergie mécanique étant positive, le météorite est dans un état de diffusion et sa trajectoire n'est pas bornée.

8. Sa distance au Soleil sera minimale lorsque $E_m = E_{p,eff}(r)$ soit lorsque :

$$\frac{1}{2}mv_0^2 = \frac{1}{2}m\frac{b^2v_0^2}{r^2} - m\frac{\mathcal{G}M_S}{r} \quad \Rightarrow \quad r^2 + r\frac{2\mathcal{G}M_S}{v_0^2} - b^2 = 0.$$

On reconnaît un polynôme du deuxième ordre en r dont on ne retient que la racine positive :

$$r_{\min} = b \left(-\frac{\mathcal{G}M_S}{bv_0^2} + \sqrt{\left(\frac{\mathcal{G}M_S}{bv_0^2}\right)^2 + 1} \right).$$

9. Si $r_{\min}$ est inférieure au rayon du Soleil, le météorite ira s'écraser sur la surface de Soleil. En pratique, s'il s'en approche de trop près, il se brise sous l'action de la chaleur dégagée par le Soleil ou des forces de marée exercée par le Soleil.

20.7 Changement d'orbite d'un satellite

1. On étudie le satellite dans le référentiel géocentrique supposé galiléen. Il est soumis à la force $\vec{f} = -\frac{GM_T m}{r^2}\vec{u}_r$ qui est une force centrale. Par conséquent, le moment de cette force par rapport au centre de la Terre est nul et le moment cinétique est constant. Le mouvement a donc lieu dans le plan contenant le centre de la Terre.

Comme le satellite est géostationnaire, il suit la Terre dans sa rotation sur elle-même donc son mouvement a lieu dans un plan perpendiculaire à l'axe de rotation de la Terre sur elle-même. Le seul plan à la fois perpendiculaire à l'axe de rotation de la Terre sur elle-même et passant par son centre est le plan équatorial. C'est donc dans ce plan qu'a lieu le mouvement d'un satellite géostationnaire.

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

2. Le mouvement est plan pour chaque orbite ainsi que pour la trajectoire de transfert. Les conditions initiales d'une orbite sont les conditions finales de la trajectoire précédente donc les trois orbites évoluent dans un même plan qui est forcément le plan équatorial puisque le mouvement final est géostationnaire.

3. On projette le principe fondamental de la dynamique appliqué au satellite sur $\vec{u}_r$ soit $-m \frac{v^2}{r} = -\frac{GM_T m}{r^2}$ donc $v = \sqrt{\frac{GM_T}{r}}$. Dans le cas d'une orbite basse à l'altitude $z \ll R_T$,

on a $r \simeq R_T$ et $v_B = \sqrt{\frac{GM_T}{R_T}}$.

4. L'application numérique donne $v_B = 7,91.10^3 \text{ m.s}^{-1} = 28,5.10^3 \text{ km.h}^{-1}$.

5. On applique la troisième loi de Kepler soit $\frac{T^2}{R_G^3} = \frac{4\pi^2}{GM_T}$. On en déduit $R_G = \sqrt[3]{\frac{GM_T T^2}{4\pi^2}}$.

6. L'application numérique donne $R_G = 42,2.10^3 \text{ km}$.

7. Pour obtenir l'altitude, il suffit de retrancher au rayon qui vient d'être trouvé le rayon terrestre soit $h = R_G - R_T = 35,8.10^3 \text{ km}$.

8. On a la même relation qu'à la question 3 avec $r = R_G$ soit

$$v_G = \sqrt{\frac{GM_T}{R_G}} = 3,07.10^3 \text{ m.s}^{-1} = 11,1.10^3 \text{ km.h}^{-1}.$$

9. Les points A et P correspondent respectivement à l'apogée et au périégée. En ces points, la vitesse est orthoradiale car on a un extremum de la distance au centre de la Terre autrement dit r est maximal ou $\dot{r} = 0$, cette dernière relation impliquant que la composante radiale de la vitesse est nulle. On en déduit qu'en ces points $v = v_\theta = r\dot{\theta}$ et en utilisant la constante des aires, on en déduit $C = r^2\dot{\theta} = r_A v_A = r_P v_P$.

10. L'énergie mécanique Em est conservée. En exprimant l'énergie mécanique à l'apogée et au périégée, on pourra écrire l'égalité des deux expressions obtenues. A l'apogée, on a $Em = \frac{1}{2}mv_A^2 - \frac{GM_T m}{r_A}$ et au périégée $Em = \frac{1}{2}mv_P^2 - \frac{GM_T m}{r_P}$. En multipliant la première relation par r_A^2 et la seconde par r_P^2 avant de les soustraire, on obtient $(r_A^2 - r_P^2)Em = \frac{1}{2}m(r_A^2 v_A^2 - r_P^2 v_P^2) - GM_T m(r_A - r_P)$ avec la conservation de l'énergie mécanique. La relation établie à la question précédente permet d'obtenir $r_A^2 v_A^2 - r_P^2 v_P^2 = 0$ et en factorisant $r_A^2 - r_P^2 = (r_A - r_P)(r_A + r_P)$, on en déduit $(r_A - r_P)(r_A + r_P)Em = -GM_T m(r_A - r_P)$. Or $r_A + r_P = 2a$ donc en simplifiant par $r_A - r_P \neq 0$, on a l'expression demandée $Em = -\frac{GM_T m}{2a}$.

11. On utilise le fait que $2a = R_G + R_T$ pour déterminer la valeur numérique de l'énergie mécanique $Em = -1,23.10^{10} \text{ J}$.

12. L'énergie mécanique s'écrit dans le cas général :

$$Em = \frac{1}{2}mv^2 - \frac{GM_T m}{r} = -\frac{GM_T m}{R_G + R_T}$$

par conservation de l'énergie mécanique. On en déduit :

$$v = \sqrt{2GM_T \left(\frac{1}{r} - \frac{1}{R_G + R_T} \right)}.$$

13. Au périégée P , il faut imposer une variation de vitesse :

$$\Delta v_P = v_P - v_B = \sqrt{2GM_T \left(\frac{1}{R_T} - \frac{1}{R_T + R_G} \right)} - \sqrt{\frac{GM_T}{R_T}}$$

soit numériquement $\Delta v_P = 2,50.10^3 \text{ m.s}^{-1} = 8,96.10^3 \text{ km.h}^{-1}$.

14. La variation d'énergie mécanique s'écrit :

$$\Delta E_{mP} = E_m - E_{mB} = -GM_T m \left(\frac{1}{R_G + R_T} - \frac{1}{2R_T} \right) = -3,46.10^{10} \text{ J.}$$

15. De même, à l'apogée, on impose une variation de vitesse :

$$\Delta v_A = v_G - v_A = \sqrt{\frac{GM_T}{R_T}} - \sqrt{2GM_T \left(\frac{1}{R_G} - \frac{1}{R_T + R_G} \right)}$$

soit numériquement $\Delta v_A = 1,50.10^3 \text{ m.s}^{-1} = 5,4.10^3 \text{ km.h}^{-1}$.

16. La variation d'énergie mécanique s'écrit

$$\Delta E_{mA} = E_{mG} - E_m = -GM_T m \left(\frac{1}{2R_G} - \frac{1}{R_G + R_T} \right) = -5,23.10^{10} \text{ J.}$$

17. La troisième loi de Kepler s'écrit $\frac{T^2}{a^3} = \frac{4\pi^2}{GM_T}$. Or $a = \frac{R_G + R_T}{2}$, on en déduit la période

$T = \pi \sqrt{\frac{(R_G + R_T)^3}{2GM_T}}$ de la trajectoire de transfert. Le transfert a une durée égale à la moitié de la période soit par l'application numérique 5 h 14 min.

CHAPITRE 20 – MOUVEMENT DANS UN CHAMP DE FORCE CENTRALE. CHAMPS NEWTONIENS.

Quatrième partie

Thermodynamique

Trouver la bibliothèque des livres ici : <https://1024terabox.com/s/1mg0wxl22G8NjM8MA2RzgFw>

Système thermodynamique à l'équilibre

21

La thermodynamique est une science née au *XIX*ème siècle qui étudie les propriétés de la matière à l'échelle macroscopique. Son champ d'application est extrêmement vaste : moteurs et centrales électriques thermiques, dispositifs réfrigérateurs destinés à produire du froid, etc. Ce chapitre présentera des modèles thermodynamiques simples pour un corps pur dans l'état solide, liquide ou gaz. On envisagera aussi le cas où deux de ces états coexistent.

1 Descriptions microscopique et macroscopique de la matière

1.1 Les phases solide, liquide et gaz

a) Aspect macroscopique

La matière existe principalement dans trois **états** bien connus : solide, liquide et gaz. Ces états, appelés aussi **phases**, sont caractérisés d'après l'expérience commune, de la manière suivante :

- un solide a une forme propre et un volume propre invariables ;
- un liquide n'a pas de forme propre (il épouse la forme d'un récipient) mais il a un volume propre invariable ;
- un gaz n'a ni volume propre, ni forme propre (il occupe tout le volume qui lui est offert).

Les phases solide et liquide ont des masses volumiques du même ordre de grandeur, 1000 fois supérieure à l'ordre de grandeur de la masse volumique d'un gaz (par exemple, la masse volumique de l'eau vaut $1.10^3 \text{ kg}\cdot\text{m}^{-3}$ et la masse volumique de l'air dans les conditions normales de température et de pression vaut $1,29 \text{ kg}\cdot\text{m}^{-3}$). Pour cette raison ces deux phases sont appelées **phases condensées**.

Les phases liquide et gaz peuvent s'écouler ; elles sont appelées **phases fluides**. La phase gaz est parfois appelée vapeur.

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

b) Aspect microscopique

La matière est constituée de particules microscopiques qui peuvent être des atomes (cas des gaz rares tels que le néon Ne, des métaux), des molécules (cas du diazote N_2 , de l'eau H_2O) ou des ions (cas d'un sel comme le chlorure de sodium NaCl). À l'échelle de ces particules, appelée échelle microscopique, les états solide, liquide et gaz se différencient par leur structure (voir figure 21.1).

Dans un solide les particules occupent des positions d'équilibre bien définies et régulièrement disposées dans l'espace. Le solide présente un ordre moléculaire à longue portée.

Dans un liquide les particules occupent des positions aléatoires. La distance moyenne entre particules est, comme dans le solide, de l'ordre de la taille des particules. Il existe un ordre moléculaire à courte portée seulement.

Dans un gaz, la distance moyenne entre particules est bien plus grande que leur taille. Il n'y a pas d'ordre moléculaire.

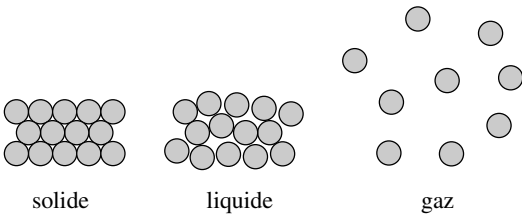


Figure 21.1 – Aspect microscopique des états solide, liquide et gaz.

1.2 L'agitation thermique

Les particules microscopiques sont constamment en mouvement, même lorsque la matière est immobile à l'échelle macroscopique. On parle d'**agitation thermique** à cause du caractère désordonné de ces mouvements.

Le physicien français Jean Perrin, au début du $XX^{\text{ème}}$ siècle, mit le premier en évidence l'agitation thermique en étudiant le *mouvement brownien*, mouvement de particules de taille de l'ordre de $1\ \mu m$ en suspension dans un liquide. Ces particules ont des trajectoires constituées de segments de droite successifs dont les directions et les longueurs sont aléatoires (voir figure 21.2). Les changements continuels de direction sont dus aux chocs de la particule brownienne avec les molécules du liquide qui sont en mouvement d'agitation thermique.



Figure 21.2 – Simulation d'une trajectoire brownienne.

Dans un solide, les particules microscopiques vibrent autour de leur position d'équilibre.
Dans un liquide ou un gaz les particules se déplacent en s'entrechoquant continuellement.
Leur mouvement est analogue à un mouvement brownien.

1.3 Échelles microscopique, mésoscopique et macroscopique

À notre échelle, tout échantillon de matière contient un très grand nombre de particules élémentaires (atomes ou molécules), nombre dont l'ordre de grandeur est donné par le nombre d'Avogadro $\mathcal{N}_A = 6,02 \cdot 10^{23} \text{ mol}^{-1}$.

On distingue trois échelles de longueurs :

- L'**échelle macroscopique** est notre échelle, son ordre de grandeur est 1 m. À cette échelle, la matière paraît continue.
- L'**échelle microscopique** est celle des particules élémentaires du système. Son ordre de grandeur est 10^{-10} m. À cette échelle, la matière est discontinue.
- L'**échelle mésoscopique** est une échelle intermédiaire, à la fois très petite devant l'échelle macroscopique et très grande devant l'échelle microscopique. Un volume de taille mésoscopique contient un très grand nombre de particules. À cette échelle la matière apparaît encore comme continue.

1.4 Le point de vue de la thermodynamique

Une description théorique complète d'un échantillon de matière à l'échelle microscopique supposerait la détermination des positions au cours de temps des N molécules de cet échantillon, soit de $3N$ coordonnées en fonction du temps avec N de l'ordre de 10^{23} . Ceci est impossible, même avec les ordinateurs les plus puissants (les méthodes de la *dynamique moléculaire* permettent de réaliser la simulation d'échantillon dont la taille, dépendant de la puissance de calcul disponible, a pu atteindre $N \sim 10^{12}$ atomes).

En thermodynamique on se place à l'échelle macroscopique (ou éventuellement mésoscopique). À cette échelle, on ne perçoit pas les constituants microscopiques individuellement mais uniquement des effets moyens dus à un très grand nombre de particules. Ainsi, on ne voit aucun mouvement provenant de l'agitation thermique parce qu'il y a autant de particules ayant une vitesse $\vec{v}$ donnée que de particules ayant la vitesse $-\vec{v}$.

Le passage d'une description au niveau microscopique à une description au niveau macroscopique est illustré par l'exemple suivant.

Exemple

On peut imaginer le système de simulation représenté sur la figure 21.3 : un cylindre contenant des billes d'acier, fermé à sa partie inférieure par un piston mobile, et à sa partie supérieure par une membrane permettant d'enregistrer les chocs des billes venant la percuter. Le piston a un mouvement oscillatoire.

S'il y a peu de billes ou si le piston va lentement, peu de billes viendront percuter la membrane pendant une seconde et on pourra enregistrer chaque choc indépendamment (voir figure 21.4).

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

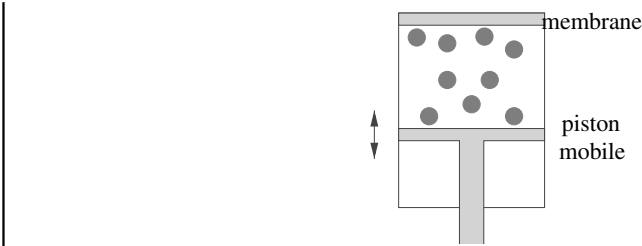


Figure 21.3 – Simulation : billes dans un cylindre.

En revanche, s'il y a beaucoup de billes et que l'oscillation est d'amplitude et de fréquence suffisantes, on ne pourra distinguer chaque choc car un nombre important de billes viendra percuter la membrane par seconde, on ne pourra mesurer que la force moyenne s'exerçant sur la paroi (voir figure 21.5).

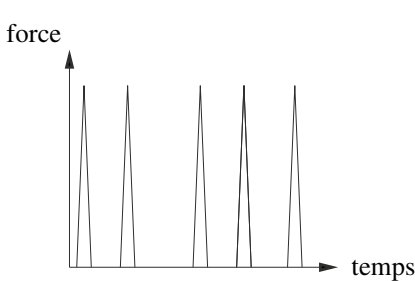


Figure 21.4 – Force exercée dans le cas d'une seule bille.

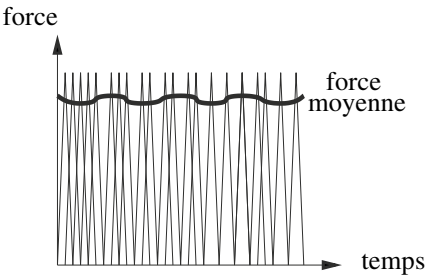


Figure 21.5 – Force exercée dans le cas de nombreuses billes.

2 Système thermodynamique, variables d'état

2.1 Système thermodynamique

On appelle **système thermodynamique** tout système constitué d'un très grand nombre de particules microscopiques.

Un échantillon de matière de taille macroscopique est un système thermodynamique. Un échantillon de matière de dimensions mésoscopiques est aussi un système thermodynamique. On peut prendre pour système un objet complexe comme, par exemple, un appareil de climatisation.

Choisir un système (on dit aussi « isoler » un système) revient à partager par l'esprit le monde en deux : d'une part le système choisi et d'autre part le reste de l'univers que l'on dénomme **extérieur**.

La surface fermée qui délimite le système est appelée **surface de contrôle**. Tout ce qui est à l'intérieur de la surface de contrôle fait partie du système et tout ce qui est à l'extérieur de la surface de contrôle ne fait pas partie du système.

On s'intéresse ensuite aux différents échanges (matière, énergie...) entre le système et l'extérieur.

Un **système fermé** est un système qui n'échange pas de matière avec l'extérieur. Dans le cas contraire on parle de **système ouvert**.

Dans le cas d'un système fermé, rien, ni matière, ni énergie, ne traverse la surface de contrôle.

Exemple
On prend pour système l'air contenu dans la chambre à air d'un pneu. La surface de contrôle est la surface intérieure de la chambre à air. C'est un système ouvert pendant qu'on gonfle le pneu. C'est un système fermé quand la voiture roule. En cas de crevaisson, c'est un système ouvert !

2.2 Variables d'état

a) Définition

Les **variables d'état** sont les grandeurs macroscopiques permettant de définir l'état d'un système thermodynamique.

Ce sont les grandeurs physiques que l'on peut définir et mesurer pour le système. Parmi ces grandeurs, il y a d'abord les grandeurs mécaniques, masse du système m , volume du système V , vitesse macroscopique qui sont définissable indépendamment de la nature moléculaire de la matière. La thermodynamique utilise deux variables d'états, pression P et la température T qui n'existent pas en mécanique, et sont issues de phénomènes à l'échelle microscopique.

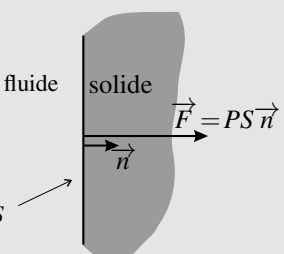
b) Pression et température

La **pression** mesure à l'échelle macroscopique l'effet des chocs des molécules en mouvement sur la paroi d'un récipient ou la membrane sensible d'un capteur de pression. Par leurs chocs, les molécules exercent une force sur la paroi. Cette force fluctue et on ne mesure, à l'échelle macroscopique, que sa valeur moyenne (comme dans le cas des billes ci-dessus). Cette force est la force de pression. Elle est orthogonale à la surface, dirigée vers l'intérieur de la paroi et de norme proportionnelle à la surface. On définit la pression ainsi :

Un fluide en contact avec une paroi solide plane de surface S exerce sur celle-ci une force :

$$\vec{F} = PS\vec{n},$$

où P est la **pression** du fluide et $\vec{n}$ le vecteur orthogonal à la paroi dirigée vers l'intérieur de celle-ci.
La pression se mesure en pascal, de symbole Pa : surface S
 $1 \text{ Pa} = 1 \text{ N}\cdot\text{m}^{-2}$.



CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

Remarque

Le pascal est une unité assez petite. On utilise souvent le bar : $1 \text{ bar} = 1.10^5 \text{ Pa}$. Mais il ne faut pas oublier que l'unité du Système International est le pascal.

La **température** est une grandeur mesurable à l'échelle macroscopique qui correspond à l'énergie cinétique qu'ont les particules microscopiques constituant la matière dans leur mouvement désordonné d'agitation thermique. La mesure de la température est fondée sur le postulat que deux systèmes en contact et en équilibre thermodynamique (voir ci-dessous) ont la même température. La température absolue T se définit à partir de la propriété des gaz d'avoir un comportement de gaz parfait lorsque la pression tend vers 0 (voir ci-dessous). Elle se mesure en kelvin, de symbole K. On utilise couramment les degrés Celsius, de symbole °C. Le lien entre la température T en kelvin et la température θ en degré Celsius est :

$$T = 273,15 + \theta.$$

c) Variables d'état d'un échantillon de corps pur

Pour donner la « taille » d'un échantillon de corps pur, on peut préciser sa masse. Il est parfois plus intéressant de considérer sa **quantité de matière** n . Ces deux variables sont liées par la relation :

$$n = \frac{m}{M},$$

où M est la **masse molaire**. Finalement :

Les variables d'état d'un échantillon de corps pur dans une seule phase sont : la température T , la pression P , le volume V et la quantité de matière n ou la masse m .

Dans le cas d'un système fermé, les variables d'état n et m sont fixées et ne peuvent être modifiées.

d) Variables extensives

Si l'on mélange deux volumes d'eau V identiques à la même température T , on obtient un volume $2V$ d'eau à la température T . Le volume et la température sont des variables de natures différentes : le volume est une variable extensive et la température une variable intensive.

Une **variable extensive** est une variable X qui dépend de la taille du système.

Si l'on réunit deux systèmes Σ_1 et Σ_2 identiques caractérisés par les valeurs $X_{\Sigma_1} = X_{\Sigma_2}$, la valeur de X pour le système $\Sigma_1 + \Sigma_2$ est :

$$X_{\Sigma_1 + \Sigma_2} = X_{\Sigma_1} + X_{\Sigma_2} = 2X_{\Sigma_1}.$$

Le volume V , le nombre de particules microscopiques N , la quantité de matière $n = \frac{N}{N_A}$ qui s'exprime en moles (unité du Système International de symbole mol) et la masse m sont des variables d'état extensives.

e) Variables intensives

Une **variable intensive** est une variable Y qui ne dépend pas de la taille du système.
Si l'on réunit deux systèmes Σ_1 et Σ_2 identiques caractérisés par les valeurs $Y_{\Sigma_1} = Y_{\Sigma_2}$, la valeur de Y pour le système $\Sigma_1 + \Sigma_2$ est :

$$Y_{\Sigma_1 + \Sigma_2} = Y_{\Sigma_1} = Y_{\Sigma_2}.$$

La pression P , la température T sont deux variables intensives.

f) Grandeurs extensives d'un corps pur monophasé

Dans ce paragraphe on suppose que le système est un échantillon de corps pur dans une seule phase. Toute grandeur extensive X relative à Σ peut s'exprimer en utilisant :

- la quantité de matière n dans le système et la **grandeur molaire** X_m par la formule :

$$X = nX_m;$$

- la masse m dans le système et la **grandeur massique** x par la formule :

$$X = mx.$$

Remarque

Les grandeurs molaires X_m et massiques x sont des grandeurs *intensives* car elles sont le quotient de deux grandeurs extensives.

On définit ainsi le **volume molaire** :

$$V_m = \frac{V}{n},$$

et le **volume massique** :

$$v = \frac{V}{m} = \frac{1}{\rho},$$

où $\rho = \frac{m}{V}$ est la masse volumique.

Les grandeurs molaires et massiques d'usage courant sont rassemblées dans le tableau 21.1.

grandeur extensive	notation	grandeur molaire	grandeur massique
volume	V	V_m	$v = \frac{1}{\rho}$
masse	m	M	1
quantité de matière	n	1	$\frac{1}{M}$
énergie interne	U	U_m	u
capacité thermique	C_V	$C_{V,m}$	c_V

Tableau 21.1 – Grandeurs extensives X , valeur molaires X_m et valeur massique x .
L'énergie interne et la capacité thermique sont définies dans le paragraphe 5.

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

3 Équilibre thermodynamique

3.1 Définition

Tout système thermodynamique abandonné à lui-même tend vers un état d'équilibre dans lequel ses variables d'état ne changent plus.

Un **système thermodynamique à l'équilibre** est un système dont les variables d'état sont toutes définies et constantes dans le temps.

3.2 Équilibre thermodynamique local

Il arrive qu'un système Σ soumis à une contrainte externe atteigne un état stationnaire (c'est-à-dire invariable dans le temps) mais dans lequel un paramètre d'état intensif Y n'est pas homogène (sa valeur dépend du point du système où on le mesure). Dans ce cas on ne peut pas attribuer une valeur à Y pour le système Σ et donc Σ n'est pas un système à l'équilibre.

Cependant on peut diviser Σ en volumes mésoscopiques pour lesquels Y a une valeur bien déterminée : il suffit que la taille des volumes mésoscopiques soit très inférieure à la distance caractéristique de variation de Y . Ceci permet de définir $Y(M)$, valeur locale de Y en un point M du système.

On fait de plus l'hypothèse d'**équilibre thermodynamique local** selon laquelle ces volumes mésoscopiques sont quasiment des systèmes thermodynamiques à l'équilibre.

Le système hors d'équilibre Σ peut être dès lors être décrit comme une réunion de systèmes thermodynamiques à l'équilibre.

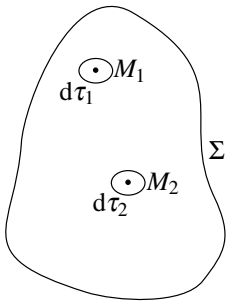


Figure 21.6 – $Y(M_1) \neq Y(M_2)$
donc Σ n'est pas un système à l'équilibre, mais $d\tau_1$ et $d\tau_2$ sont des systèmes à l'équilibre.

3.3 Conditions d'équilibre

L'équilibre thermodynamique d'un système implique que soient réalisées les trois conditions d'équilibre suivantes :

- la condition d'**équilibre mécanique**,
- la condition d'**équilibre thermique**,
- la condition d'**équilibre de diffusion**.

a) Équilibre mécanique

La condition d'équilibre mécanique suppose l'absence de tout mouvement macroscopique de matière dans le système.

Elle impose aussi que la somme des forces appliquées soit nulle sur toutes les parties mobiles dans le système et à la frontière du système.

Par exemple, les forces appliquées sur une paroi mobile de surface S séparant deux gaz de

pressions P_1 et P_2 (voir figure 21.7) sont :

- la force de pression exercée par le gaz à la pression P_1 qui s'écrit : $P_1 S \vec{u}_{1 \rightarrow 2}$ où $\vec{u}_{1 \rightarrow 2}$ est un vecteur unitaire dirigé du côté 1 du piston vers le côté 2 ;
- la force de pression exercée par le gaz à la pression P_2 qui s'écrit $-P_2 S \vec{u}_{1 \rightarrow 2}$;
- d'autres forces éventuelles (force de frottement, poids, ...) dont la résultante suivant le vecteur $\vec{u}_{1 \rightarrow 2}$ s'écrit $F_{1 \rightarrow 2} \vec{u}_{1 \rightarrow 2}$.

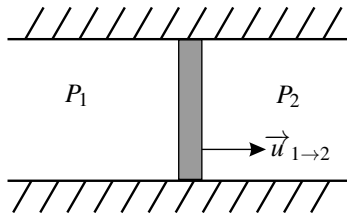


Figure 21.7 – Paroi mobile séparant deux gaz.

La condition d'équilibre mécanique de la paroi s'écrit :

$$0 = P_1 S - P_2 S + F_{1 \rightarrow 2} \quad \text{d'où} \quad P_1 = P_2 - \frac{F_{1 \rightarrow 2}}{S}.$$

Si la paroi se déplace sans frottement et si aucune force autre que les forces de pression ne s'y applique $F_{1 \rightarrow 2} = 0$, et dans ce cas la condition d'équilibre est :

$$P_1 = P_2.$$

b) Équilibre thermique

La condition d'équilibre thermique impose l'égalité de température dans tout le système. Elle impose aussi l'égalité entre la température du système et la température du milieu extérieur autour du système (qui doit être uniforme elle aussi).

Remarque

On introduira dans le chapitre suivant la notion de paroi adiabatique. Quand une telle paroi sépare deux systèmes, on peut considérer un état d'équilibre transitoire dans lequel les températures des deux systèmes ne sont pas égales.

c) Équilibre de diffusion

Cette condition intervient lorsque le système est un échantillon de corps pur *diphasé* (c'est-à-dire sous deux phases physiques différentes). Elle impose une relation entre la pression P et la température T dans le système. Ceci sera précisé dans le paragraphe 7.

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

4 Équation d'état

4.1 Définition

On appelle **équation d'état** une relation vérifiée par les variables d'état du système à l'équilibre.

Dans le cas d'un échantillon de matière sous une seule phase, il s'agit d'une relation liant la température T , la pression P , le volume V et la quantité de matière n , qui peut s'écrire symboliquement :

$$f(T, P, V, n) = 0.$$



L'équation d'état est un renseignement très utile quand on cherche à déterminer l'ensemble des variables d'état du système dans un état d'équilibre inconnu.

4.2 Équation d'état d'un gaz parfait

a) Le modèle du gaz parfait

Le **gaz parfait** est gaz théorique idéal composé de molécules ponctuelles qui n'ont aucune interaction entre elles.

Ce modèle correspond à peu près à la réalité si les molécules du gaz sont très éloignées les unes des autres, c'est-à-dire si le volume molaire du gaz V_m tend vers l'infini.

L'équation d'état du **gaz parfait** s'écrit :

$$PV = nRT, \quad (21.1)$$

où $R = 8,314 \text{ J} \cdot \text{K}^{-1} \cdot \text{mol}^{-1}$ est **constante des gaz parfaits**.



Il est important de noter que dans l'équation (21.1), les différentes grandeurs sont exprimées dans les unités du Système International :

- P en pascal,
- V en mètre cube,
- T en kelvin,
- n en mole.

D'après l'équation du gaz parfait le volume molaire V_m d'un gaz parfait, pour une température T et une pression P , est :

$$V_m = \frac{V}{n} = \frac{RT}{P},$$

et le volume massique :

$$v = \frac{V}{m} = \frac{nRT}{mP} = \frac{RT}{MP},$$

où M est la masse molaire du gaz parfait.

Exemple

Bien que les gaz réels ne se comportent en général pas comme le gaz parfait (voir ci-dessous), cette formule permet de trouver le bon ordre de grandeur du volume molaire

d'un gaz. Par exemple, dans les conditions normales de température et pression (CNTP), soit pour $T = 273,15\text{ K}$ et $P = 1,0133 \cdot 10^5\text{ Pa}$, on calcule :

$$V_m = \frac{8,3145 \times 273,15}{1,0133 \cdot 10^5} = 22,41 \cdot 10^{-3}\text{ m}^3 = 22,41\text{ L}.$$

Dans des conditions de température et de pression les gaz parfaits ont tous le même volume molaire, mais leur volume massique dépend de la masse molaire. Pour de l'air, de masse molaire $M_{\text{air}} = 29,0 \cdot 10^{-3}\text{ kg} \cdot \text{mol}^{-1}$, le volume massique dans les conditions normales est :

$$v_{\text{air}} = \frac{V_m}{M_{\text{air}}} = \frac{22,41}{29,0 \cdot 10^{-3}} = 0,773\text{ m}^3 \cdot \text{kg}^{-1}$$

et la masse volumique :

$$\rho_{\text{air}} = \frac{1}{v_{\text{air}}} = 1,29\text{ kg} \cdot \text{m}^{-3},$$

valeur qui a été donnée au début de ce chapitre.

b) Validité du modèle du gaz parfait

Comme il a été dit plus haut, le modèle du gaz parfait doit s'appliquer si le volume molaire V_m tend vers l'infini, soit, d'après l'équation d'état du gaz parfait, si la pression P tend vers 0. C'est bien ce que l'on observe expérimentalement. La figure 21.8 permet de comparer le comportement d'un gaz réel à celui d'un gaz parfait. Elle représente des courbes isothermes dans le diagramme d'Amagat et dans le diagramme de Clapeyron.

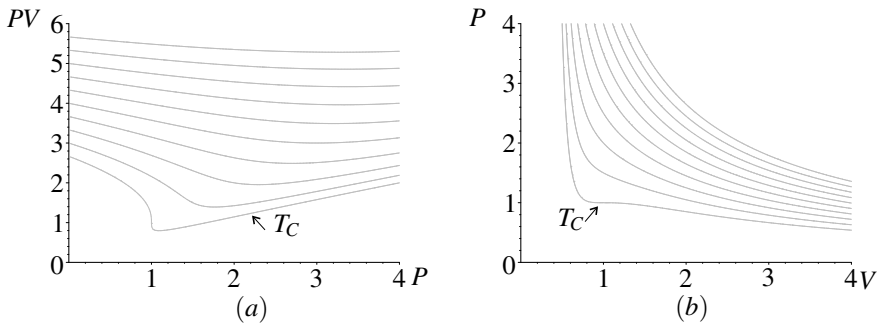


Figure 21.8 – Isothermes d'un gaz réel : (a) dans le diagramme d'Amagat, (b) dans le diagramme de Clapeyron. Les isothermes ont été tracées par un logiciel avec une équation d'état plus réaliste que l'équation du gaz parfait. Elles correspondent aux mêmes températures sur les deux diagrammes : la plus basse est la température critique T_C (définie plus loin dans ce chapitre) et les températures sont régulièrement espacées de $\frac{1}{8}T_C$. Les pressions et volumes sont dans des unités adaptées.

Le diagramme **diagramme d'Amagat** est un diagramme (PV, P) ce qui signifie que l'axe des ordonnées correspond au produit PV et l'axe des abscisses à la pression P . En thermo-

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

dynamique, il est d'usage de désigner les diagrammes par le couple (ordonnée, abscisse), soit dans l'ordre inverse de l'usage des mathématiciens. Les courbes représentées sont des **isothermes**. Sur la figure elles ont été tracées par ordinateur mais elles peuvent être obtenues expérimentalement. Pour cela on maintient un échantillon de gaz à température constante T_0 , on fait varier par étapes le volume V , on mesure la pression P et on porte le point d'abscisse P et d'ordonnée PV sur le diagramme. On construit ainsi l'isotherme T_0 point par point.

Les isothermes expérimentales pour la plupart des gaz ont l'allure qui est visible sur la figure 21.8. Ces gaz se comportent-ils comme des gaz parfaits ? La réponse est clairement non puisque, pour un gaz parfait, l'isotherme T_0 a pour équation $PV = nRT_0 = \text{constante}$, donc c'est une droite horizontale. Les isothermes d'un gaz réel dans le diagramme d'Amagat ne sont pas des droites horizontales et cela montre que les gaz réels ne se comportent pas comme des gaz parfaits. Cependant on constate expérimentalement que le point de départ des isothermes, sur l'axe des ordonnées, a une ordonnée proportionnelle à T_0 , c'est-à-dire que si P tend vers 0, le rapport $\frac{PV}{T_0}$ tend vers une constante, conformément à l'équation d'état du gaz parfait.

Le diagramme **diagramme de Clapeyron** est un diagramme (P, V) ce qui signifie que l'axe des ordonnées correspond à la pression P et l'axe des abscisses au volume V . Pour tracer expérimentalement une isotherme dans ce diagramme, il faut maintenir un échantillon de gaz à température constante T_0 , faire varier par étapes son volume V et mesurer sa pression P , pour construire point par point l'isotherme T_0 .

Pour un gaz parfait, l'isotherme T_0 dans le diagramme de Clapeyron a pour équation $PV = nRT_0 = \text{constante}$, c'est donc une hyperbole. Sur la figure 21.8, les isothermes correspondant aux plus grandes températures ont l'allure d'hyperboles mais le diagramme d'Amagat montre bien que le gaz ne se comporte pas comme un gaz parfait.

Les **gaz réels** se comportent pas comme des gaz parfaits. Mais dans la limite où la pression P tend vers 0, ils suivent approximativement la loi des gaz parfaits.

4.3 Équation d'état d'une phase condensée idéale

On appelle phase condensée un solide ou un liquide. Une phase condensée se distingue fortement d'un gaz par les caractères suivants :

- elle a une masse volumique nettement plus importante (d'où son nom de phase condensée) qui est typiquement 1000 fois plus grande que la masse volumique d'un gaz ;
- son volume diminue très peu lorsque la pression augmente : une phase condensée est très peu compressible ;
- son volume augmente très peu lorsque la température augmente : une phase condensée est très peu dilatable.

La modèle le plus simple est la **phase condensée indilatable et incompressible** dont le volume est une constante, indépendante de la température et de la pression. Son équation d'état s'écrit :

$$V = nV_{m,0}.$$

où $V_{m,0}$ est le volume molaire qui a une valeur constante, indépendante de la température et de la pression.

Exemple

Il faut connaître l'ordre de grandeur du volume molaire d'une phase condensée. On peut prendre comme point de repère l'exemple de l'eau dont la masse volumique dans les conditions normales est bien connue : $\rho_{\text{eau}} = 1,0 \cdot 10^3 \text{ kg} \cdot \text{m}^{-3}$.

La masse molaire de l'eau étant $M_{\text{eau}} = 18 \cdot 10^{-3} \text{ kg} \cdot \text{mol}^{-1}$, on peut en déduire :

$$V_{m,\text{eau}} = \frac{M_{\text{eau}}}{\rho_{\text{eau}}} = \frac{18 \cdot 10^{-3}}{1,0 \cdot 10^3} = 1,8 \cdot 10^{-5} \text{ m}^3 \cdot \text{mol}^{-1}.$$

Il est environ 1000 fois plus petit que le volume molaire du gaz parfait dans les conditions normales de température et de pression que l'on a calculé plus haut.

5 Énergie interne, capacité thermique à volume constant

5.1 L'énergie interne U

a) Définition

Pour définir l'énergie interne d'un système thermodynamique Σ on se place dans un référentiel où Σ est *macroscopiquement au repos*, c'est-à-dire qu'il n'y a pas de mouvement de matière à l'échelle macroscopique.

L'**énergie interne** U du système thermodynamique Σ est la valeur moyenne de l'énergie totale des particules microscopiques de Σ . Elle comprend :

- l'énergie cinétique des particules microscopiques,
- l'énergie potentielle d'interaction de ces particules.

L'énergie interne se mesure en joules, de symbole J.

U dépend des variables d'état T, P, V, n du système Σ . C'est une **fonction d'état** du système :

$$U = U(T, P, V, n).$$

b) Propriétés

L'énergie interne est une fonction d'état **extensive** : pour un échantillon de corps pur monophasé, elle est proportionnelle à la quantité de matière. Elle est aussi **additive**, ce qui s'écrit :

$$U_{\Sigma_1 + \Sigma_2} = U_{\Sigma_1} + U_{\Sigma_2},$$

où $\Sigma_1 + \Sigma_2$ désigne la réunion des systèmes Σ_1 et Σ_2 . Cette propriété sera utilisée pour calculer l'énergie interne d'un système complexe. Les propriétés d'extensivité et d'additivité sont liées (en particulier l'additivité implique l'extensivité) et très souvent confondues. On pourra considérer que les deux adjectifs sont synonymes.

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

5.2 La capacité thermique à volume constant C_V

On appelle **capacité thermique à volume constant** d'un système *fermé* Σ la grandeur C_V telle que la variation dU de l'énergie interne de Σ lorsque la température varie de dT , *le volume restant constant*, est :

$$dU = C_V dT.$$

C_V se mesure en joules par kelvin : $J \cdot K^{-1}$.

La capacité thermique à volume constant dépend *a priori* de la température. Ainsi, la variation d'énergie interne du système dans une transformation isochore (c'est-à-dire dans laquelle le volume du système reste constant) où la température passe de la valeur T_i à la valeur T_f est donnée par :

$$\Delta U = \int_{T_i}^{T_f} C_V(T) dT. \tag{21.2}$$

La capacité thermique à volume constant est une grandeur **extensive** et **additive**.

Pour un échantillon de corps pur, dont la « taille » est donnée par la quantité de matière n elle est se calcule par :

$$C_V = nC_{Vm},$$

où C_{Vm} est la **capacité thermique molaire à volume constant** qui s'exprime en $J \cdot K^{-1} \cdot mol^{-1}$.

Si la « taille » est donnée par la masse m , on la calcule par :

$$C_V = mc_V,$$

où c_V est la **capacité thermique massique à volume constant** qui s'exprime en $J \cdot K^{-1} \cdot kg^{-1}$.

5.3 Cas d'un gaz parfait

a) Cas du gaz parfait monoatomique

L'expression **gaz parfait monoatomique** désigne un gaz dont les molécules sont réduites à des atomes (cas des gaz rares comme l'hélium et le néon). On admet que l'énergie interne d'un gaz parfait monoatomique est donnée par :

$$U = \frac{3}{2}nRT. \tag{21.3}$$

Si T varie de dT , U varie de $dU = \frac{3}{2}nRdT$ donc la capacité thermique à volume constant est :

$$C_V = \frac{3}{2}nR. \tag{21.4}$$

La capacité thermique à volume constant molaire et la capacité thermique à volume constant massique sont donc, notant M la masse molaire du gaz :

$$C_{Vm} = \frac{3}{2}R \quad \text{et} \quad c_V = \frac{3}{2} \frac{R}{M}.$$

Numériquement : $C_{Vm} = 12,5 \text{ J} \cdot \text{K}^{-1} \cdot \text{mol}^{-1}$. La capacité thermique molaire ne dépend pas de la masse molaire du gaz parfait, mais la capacité thermique massique en dépend. Dans le cas, par exemple, de l'hélium de masse molaire $M_{\text{He}} = 4,00 \cdot 10^{-3} \text{ kg} \cdot \text{mol}^{-1}$, $c_v = 3,13 \cdot 10^3 \text{ J} \cdot \text{K}^{-1} \cdot \text{kg}^{-1}$. D'après l'équation (21.3), l'énergie interne molaire du gaz parfait monoatomique est

$$U_m = \frac{3}{2}RT.$$

Elle ne dépend que de la température. Ce résultat est connu sous le nom de **première loi de Joule** et il est vérifié par tous les gaz parfaits.

b) Première loi de Joule

Les gaz parfaits vérifient la **première loi de Joule** : leur énergie interne molaire ne dépend que de la température, ce que l'on peut écrire :

$$U_m = U_m(T).$$

La capacité thermique molaire à volume constant du gaz parfait est donc :

$$C_{Vm} = \frac{dU_m}{dT}.$$

Pour un gaz parfait non monoatomique, C_{Vm} dépend *a priori* de la température, contrairement au cas du gaz parfait monoatomique. Cependant, la plupart du temps, on travaille sur un domaine de températures dans lequel C_{Vm} est une constante. On a alors :

$$U_m(T) = C_{Vm}T + \text{constante},$$

et si la température varie de T_i à T_f la variation de l'énergie interne molaire, différence entre l'énergie interne molaire dans l'état final $U_{m,f}$ et dans l'état initial $U_{m,i}$ est :

$$\Delta U_m = U_{m,f} - U_{m,i} = C_{Vm}(T_f - T_i). \quad (21.5)$$

5.4 Cas d'une phase condensée incompressible

a) Énergie interne d'une phase condensée incompressible

Comme il a été dit plus haut, les phases condensées (solide et liquide) sont quasiment incompressibles. Ainsi, le volume molaire V_m est une constante V_{m0} . L'énergie interne molaire U_m peut dépendre *a priori* de la température T et de la pression P . On admettra qu'une variation de pression est sans influence, non seulement sur le volume, mais aussi sur les autres propriétés thermodynamiques de la phase condensée incompressible, parmi lesquelles son énergie interne.

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

L'énergie interne d'une phase condensée (solide ou liquide) incompressible ne dépend que de la température :

$$U_m = U_m(T).$$

La capacité thermique molaire à volume constant est donc :

$$C_{Vm} = \frac{dU_m}{dT}.$$

On utilise souvent pour les phases condensées l'énergie interne massique $u(T)$ et la capacité thermique à volume constant massique : $c_V = \frac{du}{dT}$.

b) Capacité thermique à volume constant

La capacité thermique à volume constant n'a pas la même valeur que pour un gaz parfait, elle peut être beaucoup plus grande. Par exemple, pour l'eau : $c_{V,\text{eau}} = 4,18 \cdot 10^3 \text{ J} \cdot \text{K}^{-1} \cdot \text{kg}^{-1}$, soit $C_{Vm,\text{eau}} = 75,2 \text{ J} \cdot \text{K}^{-1} \cdot \text{mol}^{-1}$. Toutefois l'eau est un liquide de capacité thermique particulièrement élevée.

La capacité thermique molaire à volume constant de nombreux corps purs à l'état solide aux températures usuelles vérifie $C_{Vm} \simeq 3R$ soit $C_{Vm} \approx 25 \cdot 10^3 \text{ J} \cdot \text{K}^{-1} \cdot \text{mol}^{-1}$. Ce résultat est la loi empirique de Dulong et Petit.

c) Variation de l'énergie interne molaire entre deux températures

La capacité thermique à volume constant d'une phase condensée dépend en toute rigueur de la température. Cependant, dans un domaine de température *limité*, on pourra la considérer comme quasiment constante et écrire la variation d'énergie interne molaire entre deux températures T_1 et T_2 de la manière suivante :

$$\Delta U_m = U_m(T_2) - U_m(T_1) = C_{Vm}(T_2 - T_1).$$

6 Vitesse quadratique moyenne (PTSI)

6.1 Définition

Dans ce paragraphe on s'intéresse à un gaz du point de vue microscopique. Les particules microscopiques sont animés de mouvement désordonnée et incessants. Peut-on calculer l'ordre de grandeur de leur vitesse ?

On appelle **vitesse quadratique moyenne** u la racine carrée de la moyenne du carré de la vitesse $\vec{v}$ des particules microscopiques (atomes ou molécules) du gaz :

$$u = \sqrt{\langle v^2 \rangle}, \text{ en notant } v^2 = \|\vec{v}\|^2.$$

S'il y a N atomes dans le gaz et que l'atome numéro i a pour vitesse $\vec{v}_i$, alors :

$$u^2 = \frac{1}{N} \sum_{i=1}^N v_i^2.$$

6.2 Expression en fonction de la température

Dans le cas d'un gaz monoatomique, on connaît l'expression de l'énergie interne :

$$U = \frac{3}{2} nRT.$$

Par définition, l'énergie interne est la somme de l'énergie cinétique des atomes et de leur énergie potentielle d'interaction. Pour un gaz parfait les particules n'ont pas d'énergie potentielle d'interaction.

L'énergie cinétique d'un atome, de masse m^* , et de vitesse $\vec{v}_i$ est $E_c = \frac{1}{2} m^* v_i^2$. L'énergie interne du gaz est donc :

$$U = \sum_{i=1}^N \frac{1}{2} m^* v_i^2 = \frac{1}{2} m^* \sum_{i=1}^N v_i^2 = \frac{1}{2} N m^* u^2.$$

D'autre part, $M = \mathcal{N}_A m^*$, où $\mathcal{N}_A$ est le nombre d'Avogadro, et $N = n \mathcal{N}_A$. Il vient donc :

$$U = \frac{1}{2} n M u^2.$$

En rapprochant cette expression de l'énergie interne de l'expression rappelée en début de calcul on trouve :

$$u = \sqrt{\frac{3RT}{M}}.$$

(21.6)

Cette relation est aussi valable aussi dans le cas d'un gaz de molécules polyatomiques.

6.3 Ordre de grandeur de la vitesse quadratique moyenne

On peut calculer u dans le cas de deux gaz pour se rendre compte de l'ordre de grandeur important de cette vitesse à une température de 300 K :

- pour le dihydrogène H_2 , $m_{H_2}^* = 2 \times 1,67 \cdot 10^{-27}$ kg (deux fois la masse d'un proton) et :

$$u_{H_2} = \sqrt{\frac{3 \times 1,38 \cdot 10^{-23} \times 300}{2 \times 1,67 \cdot 10^{-27}}} = 1,9 \cdot 10^3 \text{ m} \cdot \text{s}^{-1},$$

- pour le dioxygène O_2 , $m_{O_2}^* \simeq 32 \times 1,67 \cdot 10^{-27}$ kg et :

$$u_{O_2} \simeq \frac{1}{\sqrt{16}} u_{H_2} = 4,7 \cdot 10^2 \text{ m} \cdot \text{s}^{-1}.$$

On peut être impressionné par des valeurs aussi élevées. Les vitesses des molécules peuvent être très importantes.

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

7 Corps pur diphasé en équilibre

Ce paragraphe est consacré aux **systèmes diphasés**, c'est-à-dire qui contiennent un corps pur simultanément dans deux phases physiques différentes.

7.1 Changements d'état physique

Lorsque la matière évolue d'un état physique (solide, liquide ou gaz) à un autre, on dit qu'il y a **changement d'état**. Les noms des différents changements d'état sont sur la figure 21.9. Les changements d'état ont une très grande importance. Ils interviennent dans des phénomènes naturels comme le cycle de l'eau qui fait alterner précipitation et évaporation. Dans la vie courante, certains appareils utilisent ces transformations tels les réfrigérateurs, congélateurs, climatiseurs, pompes à chaleur (voir chapitre 25).

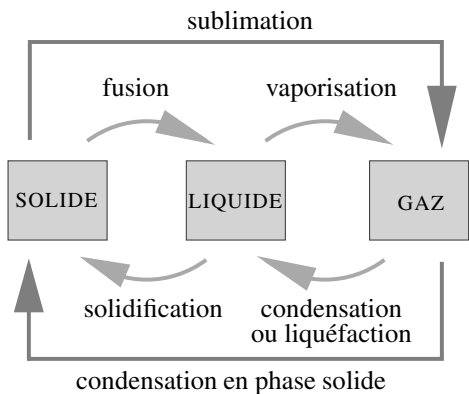


Figure 21.9 – Le vocabulaire des changements d'état.

Remarque

Pour certains corps, il existe plusieurs phases solides appelées *variétés allotropiques*. Pour l'eau, par exemple, on connaît plus de 11 types de glace. Le changement d'état entre deux variétés allotropiques d'un solide s'appelle *transition allotropique*.

7.2 Diagramme de phases (P,T)

Pour chaque corps pur on établit expérimentalement des **diagrammes de phases** qui indiquent sous quelle phase physique ce corps se présente suivant les valeurs de certains paramètres d'état. Le diagramme le plus simple est le diagramme (P,T) sur lequel on a la pression en ordonnée et la température en abscisse.

a) Deux exemples

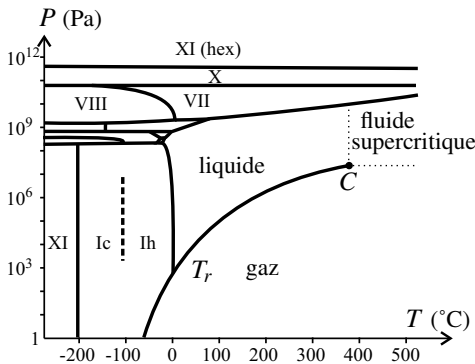


Figure 21.10 – Diagramme de phases (P, T) de l'eau. Chaque chiffre romain correspond à une variété allotropique de la glace.

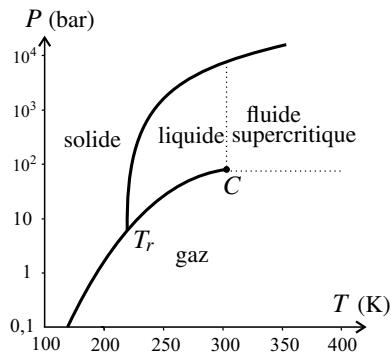


Figure 21.11 – Diagramme de phases (P, T) du dioxyde de carbone CO_2

Les figures 21.10 et 21.11 montrent les diagrammes de phases (P, T) de l'eau et du dioxyde de carbone avec une échelle logarithmique pour la pression afin de couvrir une très large gamme de valeurs. Les figures 21.12 et 21.13 montrent une partie des mêmes diagrammes avec une échelle linéaire pour la pression (et une gamme de pressions réduite).

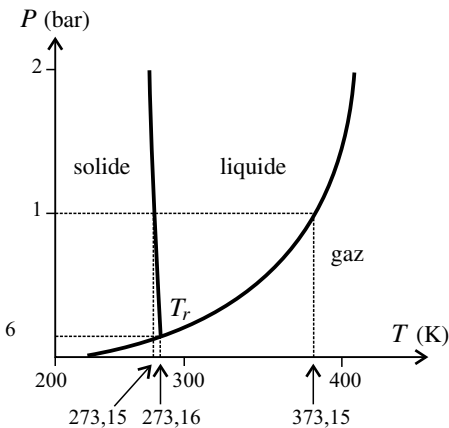


Figure 21.12 – Diagramme de phases de l'eau.

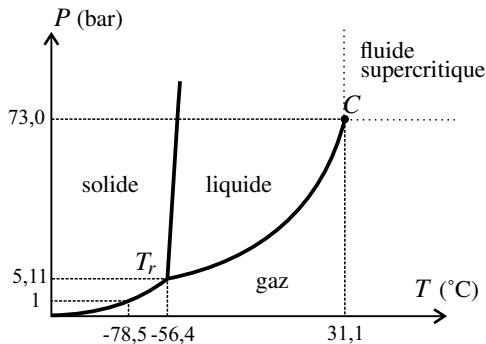


Figure 21.13 – Diagramme de phases du dioxyde de carbone.

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

b) Interprétation du diagramme

Identification des zones On peut retrouver facilement sur un diagramme simple comme ceux des figures 21.10 et 21.11 l'attribution des trois zones aux phases solide, liquide et gaz :

- la phase stable pour les plus hautes pressions et les plus basses températures est la phase solide,
- le liquide est stable pour des pressions plus élevées que le gaz.

Condition d'équilibre entre deux phases Deux phases différentes ne coexistent à l'équilibre que sur la courbe du diagramme (P, T) qui sépare leurs zones de stabilité. Ceci impose une relation entre la pression P et la température T appelée la condition d'équilibre de diffusion :

Un corps ne peut exister à l'équilibre simultanément dans deux phases différentes I et II que si la pression et la température vérifient :

$$P = P_{I-II}(T)$$

où $P_{I-II}(T)$ est la **pression d'équilibre** à la température T pour les phases I et II de ce corps pur.

La courbe séparant les zones de stabilité des phases I et II dans le diagramme de phases (P, T) est la représentation graphique de la pression d'équilibre P_{I-II} en fonction de T .

Exemple

Sur le diagramme de l'eau figure 21.12 on retrouve des valeurs connues :

- $P_{L-G}(373, 15\text{K}) = 1 \text{ bar}$, ce qui correspond au fait que l'eau bout à 100°C ;
- $P_{S-L}(273, 15\text{K}) = 1 \text{ bar}$, ce qui correspond au fait que l'eau gèle à 0°C .

Sur le diagramme du dioxyde de carbone figure 21.13, on lit que $P_{S-G}(-78, 5^\circ\text{C}) = 1 \text{ bar}$. Ainsi, sous une pression de 1 bar le gaz se transforme en solide (neige carbonique) à la température de $-78, 5^\circ\text{C}$.

Point triple Il existe sur le diagramme un point, appelé **point triple**, où les trois phases solide, liquide et gaz coexistent à l'équilibre. C'est le point noté T_r sur le diagramme, dont les coordonnées (P_r, T_r) sont telles que :

$$P_r = P_{S-L}(T_r) = P_{L-G}(T_r) = P_{S-G}(T_r).$$

Les coordonnées du point triple sont caractéristiques du corps pur.

Exemple

Les coordonnées du point triple de l'eau sont : $T_{r,H_2O} = 273,16 \text{ K}$ et $P_{r,H_2O} = 611 \text{ Pa}$.

Les coordonnées du point triple du dioxyde de carbone sont : $T_{r,CO_2} = 56,4 \text{ K}$ et $P_{r,CO_2} = 5,11 \text{ bar}$.

Point critique On peut constater sur les figures 21.10, 21.11 et 21.13 que la courbe d'équilibre entre le liquide et le gaz s'interrompt en un point noté C . Ce point est appelé **point**

critique. Les coordonnées P_C et T_C de ce point sont appelées pression critique et température critique.

L'état physique du corps pur pour une température supérieure à sa température critique et une pression supérieure sa pression critique est appelé **fluide supercritique**. Cet état fluide ne peut être qualifié ni de gazeux, ni de liquide.

Exemple

Les coordonnées du point critique de l'eau sont : $T_{C,H_2O} = 647^\circ\text{C}$ et $P_{C,H_2O} = 218 \text{ bar}$.

Expérience

La courbe d'équilibre solide-liquide de l'eau présente une pente négative, ce qui est exceptionnel : les courbes d'équilibre entre phases de natures différentes sont dans la quasi-totalité des cas des courbes croissantes.

L'expérience représentée ci-contre met en évidence cette particularité. On suspend à un fil de fer posé sur un bloc de glace deux masses de 1 kg. On constate que le fil de fer s'enfonce dans la glace qui se referme après son passage.



L'explication est que la pression sous le fil est très supérieure à la pression atmosphérique donc supérieure à la pression d'équilibre P_{S-L} entre l'eau solide et l'eau liquide à la température du glaçon (cette pression est proche de la pression atmosphérique, puisque le glaçon est à une température proche de 0°C).

Or dans le cas de l'eau, c'est le liquide qui est stable pour $P > P_{S-L}$ (voir figure 21.12) donc la glace fond sous le fil. Le fil s'enfonce et au-dessus du fil la pression étant égale à la pression atmosphérique la glace se reforme. Quand le fil a traversé la glace, le morceau de glace est intact, en un seul morceau.

7.3 Variables d'état d'un système diphasé

On considère un système Σ en équilibre thermodynamique constitué par un corps pur simultanément sous deux phases différentes I et II . Quelles sont les variables d'état nécessaires pour décrire ce système à l'équilibre ?

Si l'on connaît la température T de Σ on connaît sa pression puisqu'elle est nécessairement égale à $P_{I,II}(T)$, pression d'équilibre entre les phases I et II . De même si l'on connaît la pression P de Σ on connaît sa température T puisqu'elle est telle que : $P_{I,II}(T) = P$.

Il faut un paramètre qui donne la proportion des phases I et II dans le système Σ . On peut préciser les masses m_I et m_{II} respectives des deux phases, la masse du système étant $m = m_I + m_{II}$. On peut aussi donner les quantités de matière n_I et n_{II} respectives des deux phases, la quantité de matière du système étant $n = n_I + n_{II}$.

On utilise souvent les **titres massiques** w_I et w_{II} respectifs des phases I et II dans le système

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

définis par :

$$w_I = \frac{m_I}{m_I + m_{II}} \quad \text{et} \quad w_{II} = \frac{m_{II}}{m_I + m_{II}},$$

ou les **titres molaires** :

$$x_I = \frac{n_I}{n_I + n_{II}} \quad \text{et} \quad x_{II} = \frac{n_{II}}{n_I + n_{II}}.$$

Ces titres vérifient les relations : $w_I + w_{II} = 1$ et $x_I + x_{II} = 1$. De plus, M étant la masse molaire du corps pur considéré, on a : $m_I = n_I M$ et $m_{II} = n_{II} M$, de sorte que :

$$w_I = x_I \quad \text{et} \quad w_{II} = x_{II}.$$

L'état d'un échantillon de corps pur diphasé comportant les phases I et II est entièrement déterminé par les variables d'état suivantes :

- température T ,
- masse m ou quantité de matière n ,
- titre massique (ou molaire) x_{II} de la phase II .

Par exemple, si l'on connaît (T, m, x_{II}) on en déduit :

- la pression, qui est nécessairement :

$$P = P_{I-II}(T),$$

- les masses des deux phases :

$$m_I = (1 - x_{II})m \quad \text{et} \quad m_{II} = x_{II}m,$$

- les volumes des deux phases : $V_I = m_I v_I = (1 - x_{II})m v_I$ et $V_{II} = m_{II} v_{II} = x_{II}m v_{II}$, en notant v_I et v_{II} les volumes massiques respectifs des deux phases, et le volume du système :

$$V = V_I + V_{II} = m(1 - x_{II})v_I + m x_{II} v_{II},$$

- le volume massique global du système :

$$v = \frac{V}{m} = (1 - x_{II})v_I + x_{II} v_{II}.$$



Le volume massique ou le volume molaire d'une phase condensée sera toujours une donnée. Dans le cas d'une phase gaz en équilibre avec une phase condensée I , soit c'est une donnée, soit on le calcule, en faisant l'hypothèse que c'est un gaz parfait, par l'une des formules :

$$v_G = \frac{RT}{MP_{I-G}(T)} \quad \text{ou} \quad V_{m,G} = \frac{RT}{P_{I-G}(T)}.$$

8 Étude de l'équilibre liquide-gaz

Dans ce paragraphe on s'intéresse à l'équilibre entre un corps pur sous les formes liquide et gaz. Dans ce contexte, le gaz est fréquemment désigné par le mot « vapeur ».

8.1 Pression de vapeur saturante

La pression d'équilibre entre le liquide et le gaz à la température T , $P_{L-G}(T)$, est appelée **pression de vapeur saturante**. Dans la suite on la notera $P_{\text{sat}}(T)$, selon l'usage le plus répandu.

La courbe représentative de $P_{\text{sat}}(T)$ en fonction de T commence au point triple et se termine au point critique.

C'est, dans le plan (P, T) , le lieu des points où il y a équilibre entre le liquide et la vapeur. Elle sépare le domaine où le corps pur est liquide du domaine où il est gazeux.

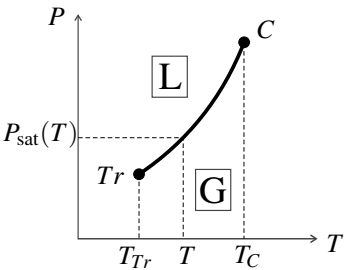


Figure 21.14 – Courbe $P_{\text{sat}}(T)$.

Pour un système contenant un corps pur, à la température T et la pression P :

- si $P < P_{\text{sat}}(T)$, le système à l'équilibre ne contient que de la vapeur, qualifiée de « vapeur sèche » ;
- si $P = P_{\text{sat}}(T)$, le système à l'équilibre contient à la fois du liquide et de la vapeur ;
- si $P > P_{\text{sat}}(T)$, le système à l'équilibre ne contient que du liquide.

On ne peut donc pas observer de vapeur sèche pour une pression supérieure à la pression de vapeur saturante.

Exemple

On introduit dans une enceinte de volume V une masse $m = 100$ g d'eau liquide. L'enceinte est maintenue à la température $T = 423$ K, température à laquelle la pression de vapeur saturante de l'eau est $P_{\text{sat}} = 4,76$ bar. On veut déterminer l'état d'équilibre atteint par l'eau pour $V = V_1 = 50$ L, puis pour $V = V_2 = 1$ L.

On fait l'hypothèse que la vapeur d'eau se comporte comme un gaz parfait. Le volume massique de l'eau liquide à la température de l'expérience est $v_L = 1,09 \cdot 10^{-3} \text{ m}^3 \cdot \text{kg}^{-1}$. La masse molaire de l'eau est $M_{\text{eau}} = 18,0 \cdot 10^{-3} \text{ kg} \cdot \text{mol}^{-1}$.

On se sait pas *a priori* si dans le système à l'état final l'eau est sous forme liquide uniquement, à la fois sous forme liquide et vapeur ou bien sous forme vapeur uniquement. On commence par supposer qu'il n'y a que de la vapeur. La pression est alors donnée par la formule des gaz parfaits :

$$P = \frac{mRT}{M_{\text{eau}}V}.$$

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

- Pour $V = V_1$ cette formule donne : $P_1 = \frac{0,100 \times 8,314}{0,018 \cdot 10^{-3}} = 3,9 \cdot 10^5 \text{ Pa} = 3,9 \text{ bar}$. Cette valeur étant inférieure à P_{sat} , l'hypothèse d'une vapeur sèche est confirmée.
- Pour $V = V_2$ cette formule donne : $P_2 = 195 \text{ bar}$, valeur supérieure à P_{sat} , ce qui est impossible. L'hypothèse d'une vapeur sèche est donc fausse. Le système est diphasé dans l'état final et $P_2 = P_{\text{sat}}$.

Pour connaître complètement l'état d'équilibre du système dans le deuxième cas, il faut calculer le titre massique en vapeur x_G .

La masse de liquide dans le système est $m_L = (1 - x_G)m$, et son volume $V_L = m_L v_L = (1 - x_G)mv_L$.

Le volume de la vapeur est ainsi $V_G = V_2 - V_L = V_2 - (1 - x_G)mv_L$. La masse de vapeur est $m_G = x_G m$, donc l'équation des gaz parfaits s'écrit :

$$P_{\text{sat}} V_G = \frac{m_G R T}{M} \quad \text{soit} \quad P_{\text{sat}} (V - (1 - x_G)mv_L) = \frac{x_G m R T}{M_{\text{eau}}}.$$

On en tire :

$$x_G = \frac{V_2 - mv_L}{\frac{m R T}{M_{\text{eau}} P_{\text{sat}}} - mv_L} = 0,22.$$

Conclusion : pour $V = V_2$ le système à l'équilibre est composé de $m_G = 22 \text{ g}$ de vapeur d'eau et $m_L = 78 \text{ g}$ d'eau liquide.

8.2 Variation de P_{sat} avec T

Comme on le voit sur la figure 21.14, P_{sat} est une fonction croissante de la température.

Pour l'eau, il existe formules semi-empiriques ou approchées donnant la pression de P_{sat} en fonction de T :

- la formule de Dupré, qui est valable sur un large domaine de température :

$$\ln P_{\text{sat}} = A - \frac{B}{T} - C \ln T \quad \text{où } A, B \text{ et } C \text{ sont des constantes ;}$$

- la formule de Duperray ci-dessous, dans laquelle θ est la température en $^{\circ}\text{C}$ et qui est valable seulement autour de 100°C :

$$P_{\text{sat}} = P_0 \left(\frac{\theta}{\theta_0} \right)^4 \quad \text{avec} \quad P_0 = 1 \text{ bar et } \theta_0 = 100^{\circ}\text{C}.$$

8.3 Température d'ébullition

On appelle **température d'ébullition** sous la pression P la température $T_{\text{eb}}(P)$ telle que :

$$P_{\text{sat}}(T_{\text{eb}}(P)) = P.$$

La température d'ébullition est une fonction croissante de la pression.

Exemple

La température d'ébullition de l'eau à la pression $P_0 = 1,0$ bar est $\theta_{eb,0} = 100^\circ\text{C}$. En altitude, la pression atmosphérique est notablement plus faible que P_0 . La température d'ébullition est donc inférieure à celle observée au niveau de la mer. À 5 000 m la pression n'est plus que $P = 0,7$ bar en moyenne et, d'après la formule de Dupré, sous cette pression la température d'ébullition est : $\theta_{eb} = \theta_{eb,0} \left(\frac{P}{P_0} \right)^{\frac{1}{4}} \simeq 90^\circ\text{C}$. Ceci n'est pas sans conséquence : en altitude, il faut plus de temps pour cuire un aliment dans l'eau bouillante car la température de l'eau est plus faible !

8.4 Diagramme de Clapeyron

Sur la figure 21.8 page 765, on a observé les isothermes d'un gaz réel dans le diagramme de Clapeyron pour des températures T supérieures à la température critique T_C . On s'intéresse maintenant aux isothermes pour des températures inférieures à T_C .

En dessous de la température critique on peut avoir deux phases fluides, liquide et vapeur. La figure 21.15 montre des isothermes obtenues expérimentalement pour l'hexafluorure de soufre SF_6 . On observe sur ces isotherme des paliers horizontaux correspondant à l'existence simultanée, dans la cellule expérimentale, de liquide et de vapeur. De telles isothermes portent le nom d'**isothermes d'Andrews**. L'ordonnée du palier est la pression de vapeur saturante $P_{\text{sat}}(T_0)$ à la température T_0 de l'isotherme.

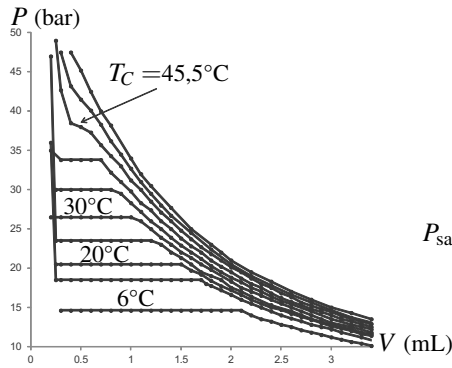


Figure 21.15 – Diagramme de Clapeyron expérimental de SF_6

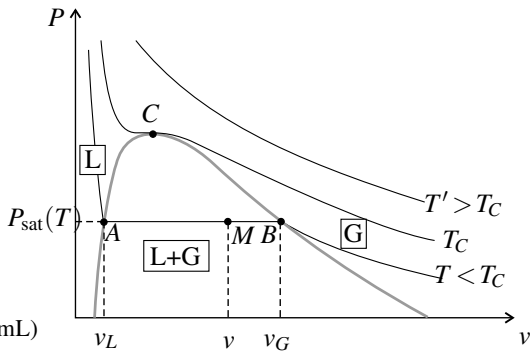


Figure 21.16 – Diagramme de Clapeyron ($P > P_r$ et $T > T_r$).

La figure 21.16 représente l'allure globale du diagramme. Les points A et B aux extrémités de ces paliers forment une courbe en forme de cloche (en gris sur la figure 21.16), appelée **courbe de saturation**, qui délimite le domaine d'existence simultanée de liquide et de gaz (domaine L+G sur la figure). La partie de la courbe de saturation située à gauche de C, de forte pente positive, s'appelle la **courbe d'ébullition**. La partie située à droite de C de pente

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

négative moins importante est la **courbe de rosée**.

De l'autre côté de la courbe de saturation par rapport au domaine L+G on trouve :

- du côté des petits volumes massiques le domaine où le liquide existe seul (domaine L sur la figure),
- du côté des grands volumes massiques le domaine où le gaz existe seul (domaine G sur la figure).

Remarque

On constate sur le diagramme expérimental de la figure 21.15 que les isothermes sont beaucoup plus pentues dans la zone du liquide que dans la zone du gaz. Ceci montre qu'il est beaucoup plus difficile de comprimer un liquide qu'un gaz.

8.5 Composition du mélange liquide-gaz

On considère, à la température T et à la pression $P_{\text{sat}}(T)$, un échantillon de corps pur contenant du liquide et du gaz. Ce système est représenté par le point M sur le diagramme de Clapeyron (voir figure 21.16). L'abscisse de M est le volume massique v de l'échantillon.

Pour connaître l'état du système il faut connaître les titres molaires ou massiques (ici identiques) x_L et x_G respectifs des phases liquide et gaz. On va voir que la position du point M sur le palier de liquéfaction AB donne ce renseignement.

On peut d'abord remarquer que si M est en A , c'est-à-dire sur la frontière avec le domaine du liquide seul, $x_L = 1$ et $x_G = 0$. On a du liquide en équilibre avec une quantité infinitésimale de vapeur, ce que l'on appelle du **liquide saturant seul**. Le volume massique du liquide saturant seul, noté v_L , se lit à l'abscisse du point A (voir figure 21.16).

De même, si M est en B , $x_L = 0$ et $x_G = 1$. On a de la vapeur en équilibre avec une quantité infinitésimale de liquide, ce que l'on appelle de la **vapeur saturante seule**. Le volume massique de la vapeur saturante seule, noté v_G , se lit à l'abscisse du point B (voir figure 21.16).

Le liquide contenu dans le système a pour volume massique v_L et le gaz a pour volume massique v_G . Ainsi, le volume total du système est : $V = m_L v_L + m_G v_G = m x_L v_L + m x_G v_G$. On a aussi : $V = m v$. Il vient donc : $v = x_L v_L + x_G v_G$ soit :

$$v = (1 - x_G) v_L + x_G v_G.$$

On en déduit que :

$$x_G = \frac{v - v_L}{v_G - v_L} = \frac{AM}{AB} \quad \text{et} \quad x_L = 1 - x_G = \frac{v_G - v}{v_G - v_L} = \frac{MB}{AB}.$$

Ainsi, plus le point M est près de B (respectivement de A) plus il y a de gaz (respectivement de liquide) dans le système.

Exemple

On reprend l'exemple de la page 777 mais on ne fait plus l'hypothèse que la vapeur est un gaz parfait. On donne la valeur tabulée du volume massique de la vapeur saturante à la température $T = 423 \text{ K}$: $v_G = 0,393 \text{ m}^3 \cdot \text{kg}^{-1}$. Le volume massique de l'eau liquide

à cette température est $v_L = 1,09 \cdot 10^{-3} \text{ m}^3 \cdot \text{kg}^{-1}$.

On cherche la nature de l'état final quand on met dans une enceinte vide, maintenue à 423 K, de volume variable V une masse $m = 100 \text{ g}$ d'eau liquide. Dans un premier cas $V = V_1 = 50 \text{ L}$, dans un deuxième cas $V = V_2 = 1 \text{ L}$.

Pour répondre à la question on calcule le volume massique du système :

- dans le premier cas : $v_1 = \frac{V_1}{m} = \frac{50 \cdot 10^{-3}}{0,100} = 0,50 \text{ kg} \cdot \text{m}^{-3}$,
- dans le deuxième cas : $v_2 = \frac{V_2}{m} = \frac{1,0 \cdot 10^{-3}}{0,100} = 0,010 \text{ kg} \cdot \text{m}^{-3}$.

On a : $v_L < v_2 < v_G < v_1$. Donc dans le premier cas le système à l'équilibre ne comporte que de la vapeur et dans le deuxième cas il comporte les deux phases. On obtient le titre massique de la phase vapeur en appliquant la formule trouvée ci-dessus :

$$x_G = \frac{v_2 - v_L}{v_G - v_L} = \frac{0,010 - 0,001}{0,393 - 0,001} = 0,23.$$

Remarque : en supposant que la vapeur se comporte comme un gaz parfait, on calcule, à 423 K le volume massique :

$$v_{G,\text{gaz parfait}} = \frac{RT}{M_{\text{eau}} P_{\text{sat}}} = \frac{8,314 \times 423}{0,018 \times 4,76 \cdot 10^5} = 0,410 \text{ m}^3 \cdot \text{kg}^{-1},$$

valeur différant de 4% de la valeur tabulée.

8.6 Point critique

Sur la figure 21.16 les courbes d'ébullition et de rosée se rejoignent au **point critique C** qui est le sommet de la courbe de saturation. L'abscisse de ce point dans le diagramme de Clapeyron est le **volume massique critique** v_C . À la température T_C , les volumes massiques v_L et v_G se rejoignent, il n'y a plus de différence entre les propriétés du liquide et de la vapeur.

Expérience

L'expérience des *tubes de Natterer* met en évidence le comportement du fluide au voisinage du point critique. Il faut choisir un fluide dont la température critique n'est pas trop élevée, comme le fréon. On dispose de trois tubes numérotés de (1) à (3), de même volume constant mais contenant des masses différentes d'un même corps : $m_1 > m_2 > m_3$ soit $v_1 < v_2 < v_3$. La masse m_2 est choisie de manière à ce que le volume massique v_2 soit égal au volume massique critique v_C .

Initialement ces tubes sont à la température T_0 comprise entre T_{Tr} et T_C (voir figure 21.17). Les points représentatifs apparaissent sur le diagramme de la figure 21.18. D'après la position des points de départ, il y a plus de liquide que de gaz dans le tube (1), à peu près autant de liquide et de gaz dans le tube (2) et plus de gaz dans le tube (3).

On chauffe doucement les trois tubes qui évoluent à volume massique constant,

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

puisque la masse et le volume de chaque tube sont constants. Les évolutions sont représentées par des droites verticales.

- Tube (1) : le point représentatif se rapproche de la courbe d'ébullition. La quantité de gaz diminue donc le ménisque de séparation liquide/gaz monte jusqu'à disparition complète du gaz.
- Tube (2) critique : le point représentatif reste à peu près à même distance des deux courbes de saturation donc le ménisque ne bouge pas. Lorsque la température atteint T_C , le système passe dans l'état de fluide supercritique, il n'y plus de différence entre liquide et gaz et le ménisque disparaît brusquement.
- Tube (3) : le point représentatif se rapproche de la courbe de rosée. La quantité de liquide diminue, donc le ménisque descend jusqu'à disparition complète du liquide.

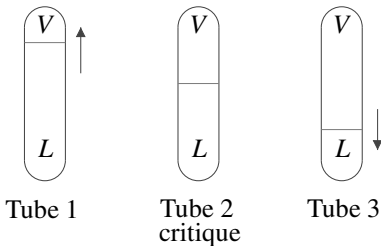


Figure 21.17 – Tubes de Natterer.

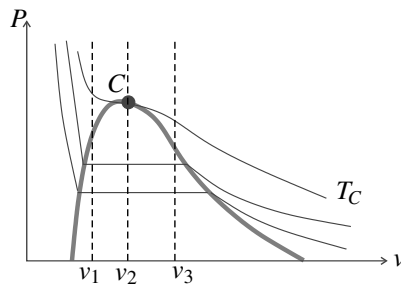


Figure 21.18 – Évolution des tubes de Natterer.

En fait, l'expérience est plus spectaculaire lorsqu'ayant réalisé ce qui précède, on laisse les tubes se refroidir à l'air libre, c'est-à-dire lentement. En ce qui concerne les tubes (1) et (3), il y a peu de changement si ce n'est l'évolution inverse des ménisques. En revanche dans le tube (2) on assiste au phénomène d'*opalescence critique* lorsque le point représentatif passe le point critique. Alors que pour une température supérieure à T_C le corps est translucide et qu'il n'y a pas de ménisque, au passage du point critique le corps devient « opaque » puis, juste après, lorsqu'il redevient translucide, le ménisque réapparaît exactement où il était auparavant.



Il faut prendre garde avant de dire qu'un liquide a une masse volumique très supérieure à celle de sa vapeur. Cela n'est vrai que si la température est très inférieure à la température critique.

8.7 Le stockage des fluides

Le point critique joue un rôle important dans le choix des conditions de stockage des fluides. Quand le fluide a sa température critique T_C inférieure à la température ambiante, la « bouteille de gaz » contient un fluide supercritique, sous pression élevée, pour réduire l'encombrement. C'est le cas avec N_2 (de température critique $\theta_{C,N_2} = -147^\circ\text{C}$), ou H_2 (de température critique $\theta_{C,H_2} = -240^\circ\text{C}$).

D'autres fluides tels NH_3 (de température critique $\theta_{C,\text{NH}_3} = 132^\circ\text{C}$) ou Cl_2 (de température critique $\theta_{C,\text{Cl}_2} = 144^\circ\text{C}$) ont leur température critique supérieure à la température de stockage. On a intérêt à stocker ces fluides sous pression élevée et sous forme de liquide en équilibre avec sa vapeur pour minimiser l'encombrement. Cependant, en cas d'échauffement accidentel de la bouteille, le fluide stocké évolue à volume massique constant, donc le long d'une verticale sur le diagramme de Clapeyron (voir figure 21.19), d'une isotherme T à une isotherme $T' > T$.

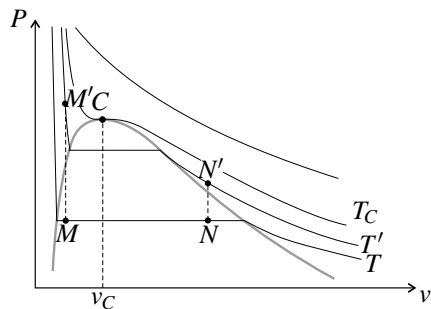


Figure 21.19 – Stockage d'un fluide.

Si le point du diagramme correspondant aux conditions de stockage se trouve à gauche du point critique (point M), cet échauffement peut amener le système dans le domaine du liquide où les isothermes sont très pentues, provoquant une augmentation considérable de la pression. Si le point de stockage se trouve à droite du point critique (point N), l'échauffement conduit le système dans la zone du gaz provoquant une augmentation de pression bien plus faible.

Il est très important que le volume massique du fluide stocké soit supérieur au volume massique critique v_C .

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

SYNTHÈSE

SAVOIRS

- ordre de grandeur du nombre d'Avogadro
- définition des échelles microscopique, mésoscopique et macroscopique
- équation d'état des gaz parfaits
- équation d'état d'une phase condensée incompressible et indilatable
- ordres de grandeur des volumes massiques et molaires
- première loi de Joule pour le gaz parfait
- $U_m = U_m(T)$ pour une phase condensée incompressible et indilatable
- vitesse quadratique moyenne (PTSI)
- diagramme d'Amagat
- diagramme de Clapeyron
- diagramme de phases (P, T)
- condition d'équilibre de diffusion entre deux phases
- condition d'équilibre d'un liquide avec un mélange gazeux
- problématique du stockage des fluides

SAVOIR-FAIRE

- reconnaître un système fermé
- déterminer le caractère extensif ou intensif d'une grandeur
- appliquer la condition d'équilibre mécanique
- appliquer la condition d'équilibre thermique
- comparer le comportement d'un gaz réel au modèle du gaz parfait dans le diagramme d'Amagat ou de Clapeyron
- analyser un diagramme de phases (P, T)
- analyser le diagramme (P, v) pour l'équilibre liquide-gaz
- calculer le titre massique en vapeur

MOTS-CLÉS

- | | | |
|---|---|-------------------------------------|
| • phase condensée | • température | • changement d'état |
| • phase fluide | • gaz parfait | • point triple |
| • macroscopique, mésosco-
pique, microscopique | • énergie interne | • point critique |
| • agitation thermique | • capacité thermique à vo-
lume constant | • pression de vapeur satu-
rante |
| • vitesse quadratique
moyenne | • diagramme d'Amagat | • stockage des fluides |
| • pression | • diagramme de Clapeyron | |
| | • incompressible | |

S'ENTRAÎNER

21.1 Pression des pneumatiques (★)

En hiver, par une température extérieure de -10°C , un automobiliste règle la pression de ses pneus à $p_1 = 2\text{ atm}$, pression préconisée par le constructeur. Cette valeur est affichée sur un manomètre qui mesure l'écart entre la pression des pneumatiques et la pression atmosphérique. On rappelle que $1\text{ atm} = 1,013.10^5\text{ Pa}$.

- 1. Quelle serait l'indication p_2 du manomètre en été à 30°C ? On suppose que le volume des pneus ne varie pas et qu'il n'y aucune fuite au niveau de ce dernier.
- 2. Calculer la variation relative de pression due au changement de température. Conclure

21.2 Existence d'une atmosphère à la surface des planètes (PTSI) (★)

1. Calculer numériquement la vitesse de libération (Cf. cours de mécanique sur les mouvements newtoniens) et la vitesse quadratique moyenne pour le dihydrogène et pour le diazote à la surface des quatre planètes telluriques (Mercure, Vénus, Terre, Mars), pour une température de 300 K . Conclure quant à la présence possible d'une atmosphère à cette température. On donne :

Planète	Diamètre équatorial en km	Rapport de la masse à celle de la Terre
Mercure	4878	0,055
Vénus	12 104	0,815
Terre	12 756	1
Mars	6794	0,107

Masse de la Terre : $M_T = 5,976\text{ }10^{24}\text{ kg}$; $G = 6,67.10^{-11}\text{ m}^3\text{kg}^{-1}\text{s}^{-2}$; masse des molécules : $m_{H_2} = 3,3.10^{-27}\text{ kg}$ et $m_{N_2} = 4,65.10^{-26}\text{ kg}$; constante Boltzmann : $k_B = 1,38.10^{-23}\text{ J.K}^{-1}$.

2. Quel devrait être l'ordre de grandeur de la température pour que les molécules de diazote échappent à l'attraction terrestre ?

21.3 Effusion d'un gaz (PTSI) (★★)

On considère deux compartiments de volume V_1 et V_2 , l'ensemble est maintenu à la température T . Entre les deux compartiments, un petit trou de section s a été pratiqué. Initialement on a N_a particules d'un gaz parfait dans V_1 . On note N_1 et N_2 les nombres de particules dans les volumes V_1 et V_2 et on adopte le modèle suivant : les particules ont toutes le même module de vitesse v et leur vitesse est suivant l'une des directions $\vec{u}_x, \vec{u}_y, \vec{u}_z, -\vec{u}_x, -\vec{u}_y, -\vec{u}_z$.

- 1. Quel est le nombre $dN_{1\rightarrow 2}$ passant de V_1 à V_2 entre t et $t + dt$?
- 2. En déduire les équations différentielles vérifiées par N_1 et N_2 en fonction de N_1, N_2, s, v et $V = V_1 = V_2$.
- 3. Établir les expressions de N_1 et de N_2 en fonction du temps.
- 4. Définir un temps caractéristique τ .

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

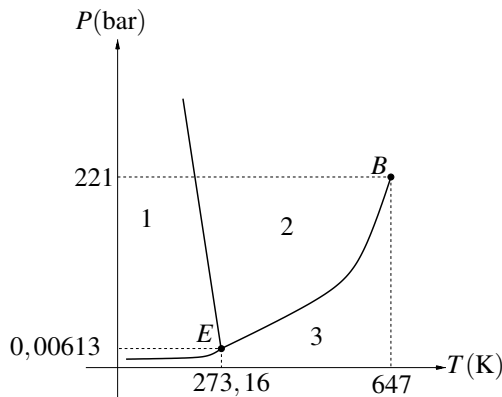
5. Comment varie-t-il en fonction de la masse du gaz si on admet que $v = \sqrt{\frac{3RT}{M}}$?
6. Quelle peut être l'application pratique de ce phénomène d'effusion gazeuse ?

21.4 Questions qualitatives sur les changements d'état (★)

1. Dans une centrale nucléaire à REP (réacteur à eau pressurisée) le circuit primaire contient de l'eau liquide dont la température varie entre 286°C en entrée de la cuve du réacteur et 323°C en sortie de cuve. En utilisant un document du cours, déterminer si la pression est de 1,55 bar, 15,5 bar ou 155 bar.
2. Pourquoi les patins glissent-ils sur la glace ?
3. On remplit à moitié une bouteille d'eau minérale en plastique avec de l'eau chaude, puis on la ferme bien. Que se passe-t-il quand la bouteille refroidit ? Pourquoi ?

21.5 Étude de quelques transformations d'un corps d'après CCP TSI (★)

Dans cet exercice, on s'intéresse à l'eau dont le diagramme des phases est donné par :



1. Reproduire ce diagramme et le compléter en donnant les domaines d'existence des différentes phases et en définissant les points caractéristiques.
2. Définir la pression de vapeur saturante et préciser de quel(s) paramètre(s) elle dépend.
3. Comment appelle-t-on le passage de la vapeur au liquide ?
4. Représenter le diagramme donnant la pression en fonction du volume pour la transformation correspondante. On définira les domaines et on tracera les courbes de rosée et d'ébullition.
5. Soit une enceinte cylindrique diathermane de volume initial V , ce volume pouvant être modifié en déplaçant sans frottement un piston. L'ensemble est maintenu sous la pression atmosphérique à la température $T = 373$ K. A cette température la pression de vapeur saturante vaut 1,0 bar. La vapeur d'eau sèche et saturante sera considérée comme un gaz parfait. On néglige le volume occupé par la phase liquide devant le volume occupé par la vapeur, ainsi le volume de la phase gazeuse est égal au volume total de l'enceinte. Le cylindre étant initialement vide, on introduit, piston bloqué, une masse m d'eau. Déterminer la masse maximale $m_{\max}$ d'eau qu'on peut introduire pour que l'eau soit entièrement sous forme vapeur. On

donnera sa valeur en fonction de R , T , V , P_S la pression de vapeur saturante et M_{eau} la masse molaire de l'eau.

6. On considère que la masse d'eau introduite est inférieure à m_{max} . Dans quel état se trouve l'eau ?

7. Pour obtenir l'équilibre entre les phases liquide et vapeur de l'eau, faut-il augmenter ou diminuer le volume ?

Déterminer le volume limite V_{lim} à partir duquel on a cet équilibre.

8. La masse m d'eau introduite est telle qu'on a l'équilibre entre les phases liquide et vapeur. Déterminer la fraction massique d'eau sous forme vapeur.

APPROFONDIR

21.6 Équilibre d'un piston (★)

Un cylindre vertical fermé aux deux bouts est séparé en deux compartiments égaux par un piston homogène, se déplaçant sans frottement. La masse du piston par unité de surface est $\sigma = 1360 \text{ kg.m}^{-2}$. Les deux compartiments contiennent un gaz parfait à la température $t_1 = 0^\circ\text{C}$. La pression qui règne dans le compartiment supérieur est égale à $P_H = 0,133 \text{ bar}$. L'intensité de la pesanteur est $g = 10 \text{ m.s}^{-2}$.

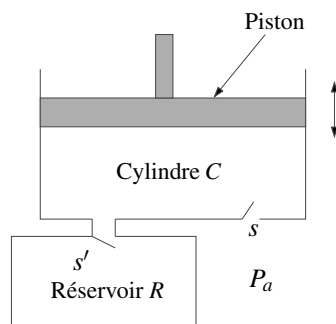
1. En écrivant que le piston est à l'équilibre, déterminer la pression, en bars, du gaz dans le compartiment du bas.

2. On porte les deux compartiments à $t_2 = 100^\circ\text{C}$. De combien se déplace le piston ?

21.7 Étude d'un compresseur (★★)

Un compresseur est constitué de la façon suivante : un piston se déplace dans un cylindre C qui communique par des soupapes s et s' respectivement avec l'atmosphère (pression P_a) et avec le réservoir R contenant l'air comprimé. Le réservoir R contient initialement de l'air considéré comme gaz parfait à la pression $P_0 \geq P_a$.

Le volume du réservoir R , canalisations comprises, est V . Le volume offert au gaz dans C varie entre un volume maximum V_M et un volume minimum V_m , volume nuisible résultant de la nécessité d'allouer un certain espace à la soupape s .



- La soupape s s'ouvre lorsque la pression atmosphérique P_a devient supérieure à la pression dans le cylindre C et se ferme pendant la descente du piston.
- La soupape s' s'ouvre lorsque la pression dans le réservoir devient supérieure à celle du gaz dans le cylindre C et se ferme pendant la montée du piston.

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

Au départ, le piston est dans sa position la plus haute ($V = V_M$), s' est fermée, s est ouverte et le volume V_m est rempli d'air à la pression P_a .

- 1. a. En supposant que le piston se déplace assez lentement pour que l'air reste à température constante, calculer le volume V'_1 pour lequel s' s'ouvre, en fonction de P_0 , P_a et V_M .
b. Calculer la pression P_1 dans le réservoir R après le premier aller et retour.
c. En écrivant une condition sur V'_1 , calculer la valeur P_{max} au-dessus de laquelle la pression ne peut pas monter dans le réservoir.
- 2. Calculer la pression P_n dans le réservoir R après n allers et retours du piston.
- 3. Donner la valeur limite de P_n quand $n \gg 1$. Comparer cette limite avec P_{max} .
- 4. Calculer P_1 et P_{max} avec $V = 5 \text{ L}$, $V_M = 0.25 \text{ L}$, $V_m = 10 \text{ cm}^3$, $P_0 = P_a = 1 \text{ bar}$.

21.8 Énergie interne d'un gaz - Gaz réel ou parfait ? (★★)

Les valeurs expérimentales de l'énergie interne massique de la vapeur d'eau sont les suivantes :

$T \text{ (K)}$	523	573	623	673
à $P = 10 \text{ bars}$	2711	2793	2874	2956
à $P = 20 \text{ bars}$	2683	2773	2859	2944

- 1. Tracer les courbes donnant l'énergie interne en fonction de la température.
- 2. A-t-on un gaz parfait ? Justifier.
- 3. Comparer la capacité thermique à volume constant à celle d'un gaz parfait monoatomique.

CORRIGÉS

21.1 Pression de pneumatiques

1. En considérant le gaz à l'intérieur comme parfait et le volume constant, on a $\frac{P_1}{T_1} = \frac{nR}{V_1} = \frac{nR}{V_2} = \frac{P_2}{T_2}$ avec $P_1 = P_0 + p_1$ et $P_2 = P_0 + p_2$. On en déduit $p_2 = \frac{T_2}{T_1} (P_0 + p_1) - P_0$ soit 2,5 atm.
2. $\Delta p = 0,50$ atm soit 20 % qui est supérieur à l'écart maximal toléré dont il faut changer la pression en fonction de la saison.

21.2 Existence d'une atmosphère à la surface des planètes (PTSI)

1. La vitesse de libération est obtenue pour une énergie mécanique nulle :

$$E_m = \frac{1}{2}mv_l^2 - \frac{GMm}{R} = 0 \Rightarrow v_l = 2\sqrt{\frac{MG}{R}}$$

M étant la masse de la planète et R son rayon.

On rappelle que la vitesse quadratique moyenne est $u = \sqrt{\frac{3k_B T}{m}}$, m étant la masse d'une particule.

Planète	v_l en km.s^{-1}
Mercure	4,24
Vénus	10,4
Terre	11,2
Mars	5,01

Les vitesses quadratiques moyennes sont : $u_{H_2} = 1,94 \text{ km.s}^{-1}$ et $u_{N_2} = 0,52 \text{ km.s}^{-1}$. Il faut se souvenir que la vitesse quadratique moyenne est une moyenne, donc de nombreuses particules sont plus rapides. D'après les données précédentes, la vitesse quadratique est beaucoup plus proche des vitesses de libération pour Mercure et Mars que pour la Terre et Vénus. Les particules s'échappent plus facilement de l'attraction de Mars et Mercure, ce qui peut expliquer qu'il n'y ait pas d'atmosphère.

2. Pour que N_2 échappe à l'attraction terrestre, il faut que $u_{N_2} = v_l$, soit $T = 1,4.10^5 \text{ K}$.

21.3 Effusion d'un gaz (PTSI)

1. Les particules de V_1 passant dans V_2 entre t et $t + dt$ sont celles contenues dans un cylindre de base s et de hauteur $v dt \vec{u}_x$ compte tenu des hypothèses très simplifiées considérées. Son volume est $v s dt$ et le nombre de ces particules $N_1 \frac{v s dt}{6V}$, le facteur 6 est lié au fait que seulement un sixième des particules ont une vitesse $\vec{v} u_x$. On en déduit $dN_{1 \rightarrow 2} = \frac{N_1 v s}{6V} dt$.

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

2. De même, on a $dN_{2 \rightarrow 1} = \frac{N_2 v s}{6V} dt$. On en déduit

$$dN_1 = -dN_2 = -dN_{1 \rightarrow 2} + dN_{2 \rightarrow 1} = \frac{vs}{6V} (N_2 - N_1) dt$$

3. Comme $N_1 + N_2 = N_a$, on obtient $dN_1 = \frac{vs}{6V} (N_a - 2N_1) dt$. On en déduit l'équation différentielle $\frac{dN_1}{dt} + \frac{vs}{3V} N_1 = \frac{vs}{6V} N_a$.

La résolution de cette équation différentielle du premier ordre à coefficients constants donne $N_1 = \frac{N_a}{2} + C \exp\left(-\frac{vs t}{3V}\right)$. Or à $t = 0, 0$ s, $N_1 = N_a = \frac{N_a}{2} + C$ donc $C = \frac{N_a}{2}$. Finalement $N_1 = \frac{N_a}{2} \left(1 + \exp\left(-\frac{t}{\tau}\right)\right)$ avec $\tau = \frac{3V}{vs}$ et $N_2 = \frac{N_a}{2} \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$.

4. Avec $v = \sqrt{\frac{3RT}{M}}$, on en déduit $\tau = \frac{V}{s} \sqrt{\frac{3M}{RT}}$ donc le temps caractéristique est proportionnel à $\sqrt{M}$.

5. La diffusion est proportionnelle à $\sqrt{M}$, donc si on a deux types de particules de masses différentes, on peut enrichir les mélanges en bloquant le système avant la fin.

21.4 Questions qualitatives sur les changements d'état

1. D'après le diagramme d'équilibre de l'eau représenté sur la figure 21.10, page 773, pour que l'eau soit liquide à une température voisine de 300°C , il faut que la pression soit de l'ordre de 1.10^7 Pa. La réponse est donc 155 bar.

2. Sous les patins à glace la pression est plus forte que la pression atmosphérique. Elle est augmentée du rapport du poids du patineur divisé par la surface de contact des patins. À cause de la pente négative de la courbe d'équilibre liquide-solide dans le diagramme (P, T) de l'eau (voir figure 21.12, page 773), une augmentation de pression peut faire passer le point représentant l'état de l'eau de la zone du solide à la zone du liquide. Ainsi la glace fond sous les patins.

3. Quand elle refroidit, la bouteille est écrasée par la pression atmosphérique parce que sa pression interne diminue et devient fortement inférieure à la pression atmosphérique. Il y a deux raisons pour cela :

- d'après la loi des gaz parfaits, la pression diminue si la température diminue ;
- la quantité d'eau dans la phase gazeuse à l'intérieur de la bouteille diminue parce qu'elle est fixée par la condition d'équilibre : $P_{\text{eau}} = \frac{n_{\text{eau}} RT}{V} = P_{\text{sat}}(T)$ et que $P_{\text{sat}}(T)$ diminue fortement avec la température.

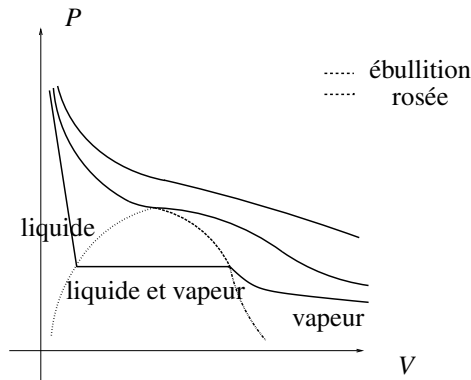
21.5 Étude de quelques transformations d'un corps d'après CCP TSI

1. Le domaine (1) est le domaine de stabilité de la phase solide, (2) celui de la phase liquide et (3) celui de la phase gazeuse. Le point (B) est le point critique au-delà duquel on ne distingue plus les phases liquide et vapeur : on a une phase fluide.

Le point (E) est le point triple, en ce point coexistent les trois phases.

2. La pression de vapeur saturante est la pression d'équilibre entre la phase gazeuse et la phase liquide à une température donnée ; elle ne dépend que de la température.

- 3. Le passage de la vapeur au liquide s'appelle liquéfaction.
- 4. Le diagramme demandé est le suivant :



- 5. Les vapeurs étant considérées comme des gaz parfaits, on peut écrire $P_s V = n_{\max} RT$. On en déduit $m_{\max} = \frac{M_{\text{eau}} P_s V}{RT}$.
- 6. Si $m < m_{\max}$, on a de la vapeur car cela correspond à $P < P_s$.
- 7. On n'a que du gaz si $P < P_s$ soit $V > V_{\text{lim}} = \frac{mRT}{M_{\text{eau}} P_s}$.
- 8. Le titre de vapeur est $x_{\text{vap}} = \frac{m_{\max}}{m} = \frac{M_{\text{eau}} P_s V}{mRT}$.

21.6 Équilibre d'un piston

On utilisera un indice B (respectivement H) pour les variables relatives au gaz du bas (respectivement au gaz du haut).

- 1. Les forces qui s'exercent sur le piston sont : son poids $m \vec{g}$, la force de pression exercée par le gaz du bas $\vec{F}_{p,B}$ dirigée vers le haut, et la force de pression exercée par le gaz du haut $\vec{F}_{p,H}$ dirigée vers le bas. On projette le principe fondamental de la dynamique à l'équilibre sur un axe vertical ascendant : $-mg + P_B S - P_H S = 0$, où S est la section du tube. La masse par unité de surface σ étant $\frac{m}{S}$, on obtient la relation cherchée :

$$P_B = P_H + \sigma g$$

L'application numérique donne :

$$P_B = \left(\frac{10}{76} 1,013 \cdot 10^5 + 13600 \right) \text{ Pa} = 0,269 \text{ bar.}$$

- 2. On cherche à déterminer la pression finale dans le compartiment du haut, ce qui permettra de déterminer le volume final de ce compartiment et donc le déplacement du piston. Le volume total, noté V, se conserve.
Pour le système gaz du haut (nombre de moles n_H) :

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

$$\begin{array}{c|c} \text{Etat initial} & \begin{array}{l} P_H = P_1 \\ V_H = \frac{V}{2} \\ T_H = 273 \text{ K} \end{array} \\ \hline \text{Etat final} & \begin{array}{l} P'_H \\ V'_H \\ T'_H = 373 \text{ K} \end{array} \end{array}$$

Pour le système gaz du bas (nombre de moles n_B) :

$$\begin{array}{c|c} \text{Etat initial} & \begin{array}{l} P_B = P_1 + \sigma g \\ V_B = \frac{V}{2} \\ T_B = 273 \text{ K} \end{array} \\ \hline \text{Etat final} & \begin{array}{l} P'_B \\ V'_B \\ T'_B = 373 \text{ K} \end{array} \end{array}$$

L'équilibre du piston entraîne : $P'_B = P'_H + \sigma g$.

L'équation du gaz parfait donne, pour le gaz H :

$$n_H R = \frac{P_1 V_H}{T_1} = \frac{P'_H V'_H}{T_2} \quad \text{donc} \quad V'_H = \frac{T_2}{T_1} P_1 V_H$$

et pour le gaz B :

$$n_B R = \frac{(P_1 + \sigma g) V_B}{T_1} = \frac{(P'_H + \sigma g) V'_B}{T_2} \quad \text{donc} \quad V'_B = \frac{T_2}{T_1} \frac{P_1 + \sigma g}{P'_H + \sigma g} V_B$$

La conservation du volume entraîne : $V = V_H + V_B = V'_H + V'_B$ soit $2V_H = V'_H + V'_B$. On remplace V'_H et V'_B par les expressions obtenues plus haut :

$$2V_H = \frac{T_2 P_1}{T_1 P'_H} V_H + \frac{T_2}{T_1} \frac{P_1 + \sigma g}{P'_H + \sigma g} V_H$$

La pression P'_H vérifie l'équation du second degré :

$$2 \frac{T_1}{T_2} P_H'^2 + P'_H \left(2 \frac{T_1}{T_2} \sigma g - 2P_1 - \sigma g \right) - P_1 \sigma g = 0$$

La seule solution positive est $P'_H = 1,99 \cdot 10^4 \text{ Pa}$.

Pour calculer le rapport $\frac{h'}{h}$ des hauteurs du piston dans l'état final et dans l'état initial, il suffit d'écrire :

$$\frac{h'}{h} = \frac{V'_H}{V_H} = \frac{P_1 T_2}{P'_H T_1} = 0,9$$

Le piston monte de donc de $0,1h$.

21.7 Étude d'un compresseur

1. a. Initialement le piston est en haut. Lorsqu'il commence sa descente, la soupape s se ferme. Le système considéré est donc le gaz dans le cylindre de pression P , jusqu'à l'ouverture de s' pour V'_1 et $P = P_0$.

$$\begin{array}{c|c} \text{Etat initial} & \begin{array}{l} P = P_a \\ V = V_M \\ T \\ n_C \end{array} \\ \hline \text{Etat final} & \begin{array}{l} P = P_0 \\ V = V'_1 \\ T \\ n_C \end{array} \end{array}$$

On applique l'équation du gaz parfait pour trouver V_1' :

$$V_1' = \frac{P_a V_M}{P_0}$$

b. La soupape s' étant ouverte, le système est constitué du gaz dans le cylindre et du gaz dans le réservoir. L'état final est obtenu lorsque le piston est au fond du cylindre.

Etat initial	$\left \begin{array}{c} P_0 \\ V_1' + V \\ T \\ n_C + n_R \end{array} \right $	Etat final	$\left \begin{array}{c} P_1 \\ V_m + V \\ T \\ n_C + n_R \end{array} \right $
--------------	---	------------	--

On applique de nouveau l'équation du gaz parfait :

$$(n_C + n_R)RT = P_1(V + V_m) = P_0(V + V_1')$$

En utilisant la relation $P_0 V_1' = P_a V_M$, on obtient :

$$P_1 = P_a \frac{V_M}{V + V_m} + P_0 \frac{V}{V + V_m} \tag{21.7}$$

c. Pour que la soupape s' s'ouvre, il faut que le volume V_1' soit atteint avant que le piston ne soit complètement descendu, donc que : $V_1' \geq V_m$. On en déduit : $P_0 \leq P_{\max}$ avec $P_{\max} = P_a \frac{V_M}{V_m}$.

2. Pour les allers et retours suivants le raisonnement est semblable. Pour obtenir la pression P_2 , il suffit de remplacer P_0 par P_1 , dans l'expression (21.7) :

$$P_2 = P_a \frac{V_M}{V + V_m} + P_1 \frac{V}{V + V_m}$$

On remplace P_1 par son expression :

$$P_2 = P_a \left(\frac{V_M}{V + V_m} + \left(\frac{V_M}{V + V_m} \right)^2 \right) + P_0 \left(\frac{V}{V + V_m} \right)^2$$

On voit apparaître la formule de récurrence suivante :

$$P_n = P_a \sum_{i=1}^n \left(\frac{V_M}{V + V_m} \right)^i + P_0 \left(\frac{V}{V + V_m} \right)^n = P_a \frac{1 - \left(\frac{V_M}{V + V_m} \right)^n}{1 - \left(\frac{V_M}{V + V_m} \right)} + P_0 \left(\frac{V}{V + V_m} \right)^n$$

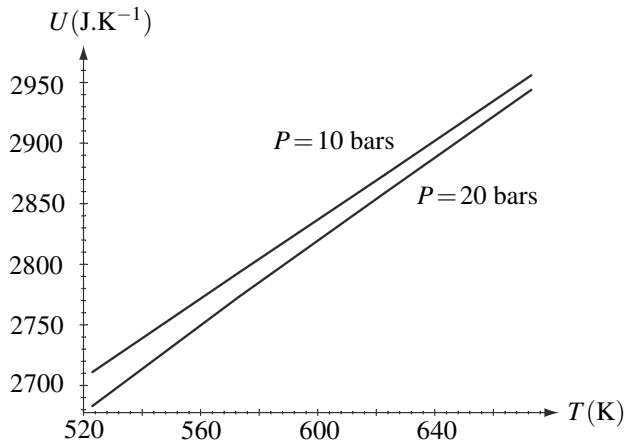
La limite pour n très grand donne : $\lim_{n \rightarrow \infty} P_n = P_a \frac{1}{1 - \frac{V_M}{V + V_m}} = P_{\max}$.

3. Applications numériques : $P_1 = 1,05$ bar et $P_{\max} = 25$ bar.

CHAPITRE 21 – SYSTÈME THERMODYNAMIQUE À L'ÉQUILIBRE

21.8 Énergie interne d'un gaz - Gaz réel ou parfait ?

1. On obtient :



2. La variation d'énergie interne est proportionnelle à celle de la température du fait qu'on a une droite mais les deux courbes ne sont pas confondues donc l'énergie interne dépend de la pression. Le gaz n'est pas parfait : il faut tenir compte des interactions entre molécules.

3. Le calcul des pentes donne 28 J.mol^{-1} pour $P = 10 \text{ bars}$ et 29 J.mol^{-1} pour $P = 20 \text{ bars}$. Ce sont les valeurs des capacités thermiques. Dans le cas d'un gaz parfait monoatomique, on aurait $\frac{3}{2}R = 12,5 \text{ J.mol}^{-1}$. On a donc une pente plus élevée du fait des vibrations et des rotations.

Énergie échangée par un système au cours d'une transformation

22

La thermodynamique des systèmes à l'équilibre s'intéresse à des transformations d'un système thermodynamique d'un état d'équilibre à un autre. Dans ce chapitre on va définir différents types de transformations. On va ensuite s'intéresser aux deux types d'échanges d'énergie que le système peut avoir avec l'extérieur au cours d'une transformation : le travail et le transfert thermique.

1 Transformation thermodynamique

1.1 Transformation, état initial, état final

On appelle **transformation thermodynamique** le passage d'un système thermodynamique d'un état d'équilibre, appelé **état initial**, à un nouvel état d'équilibre, appelé **état final**.

Lorsqu'on étudie une transformation thermodynamique il faut toujours bien préciser le système Σ considéré. Celui-ci devra toujours être un **système fermé**.

Pour provoquer la transformation d'un système Σ il faut imposer à Σ une modification d'une de ses variables d'état ou bien changer les conditions extérieures. On met ainsi le système hors d'équilibre et il évolue vers un nouvel état d'équilibre.

On connaît toujours l'état d'équilibre initial. Comment détermine-t-on l'état d'équilibre final ?

On obtient des renseignements sur les variables d'état finales en appliquant :

- la condition d'équilibre mécanique,
- la condition d'équilibre thermique (sauf si la transformation est trop rapide pour qu'il s'établisse, voir paragraphe 3),

CHAPITRE 22 – ÉNERGIE ÉCHANGÉE PAR UN SYSTÈME AU COURS D’UNE TRANSFORMATION

- la condition d’équilibre de diffusion dans le cas d’un système diphasé,
- les équations d’état des différentes phases existant dans le système.

Ces différents renseignements ne sont pas toujours suffisants. Il faut parfois utiliser d’autres informations concernant la transformation. Celle-ci peut avoir différentes propriétés que l’on va définir maintenant.

1.2 Différents types de transformations

a) Transformation isochore

Une transformation est **isochore** quand le volume du système est constant au cours de la transformation.

En notant V_i , V_f et V le volume du système respectivement dans l’état initial, dans l’état final et dans un état intermédiaire quelconque au cours de la transformation, on a : $V_i = V = V_f$.

Un système enfermé dans un récipient rigide indéformable subit des transformations obligatoirement isochores. Lorsque la transformation est isochore on connaît *a priori* le volume dans l’état final.

Exemple

Soit un échantillon de gaz dans l’état initial (T_i, P_i, V_i) enfermé dans un récipient indéformable. On peut le faire passer dans un état final $(T_f, P_f, V_f = V_i)$ de manière isochore en le plaçant dans un milieu extérieur à la température T_0 . La transformation est terminée quand l’équilibre thermique entre le gaz et l’extérieur est établi donc quand : $T_f = T_0$.

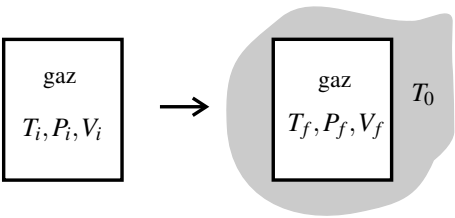


Figure 22.1 – Exemple de transformations isochore d’un échantillon de gaz. Le récipient indéformable impose un volume constant au gaz.

Si le gaz est parfait, on peut calculer la pression de l’état final P_f par la loi des gaz parfaits : $\frac{P_i V_i}{T_i} = \frac{P_f V_f}{T_f} = nR$, en notant n la quantité de gaz, d’où : $P_f = P_i \frac{T_f}{T_i} = P_i \frac{T_0}{T_i}$.

b) Transformation isobare

Une transformation est **isobare** quand la pression du système est définie tout au long de la transformation et garde une valeur constante.

En notant P_i , P_f et P la pression du système respectivement dans l'état initial, dans l'état final et dans un état intermédiaire quelconque au cours de la transformation, on a : $P_i = P = P_f$.

En pratique une transformation isobare est une transformation assez lente, dans laquelle on impose de l'extérieur sa pression au système.

Exemple

Soit un échantillon de gaz dans l'état initial (T_i, P_i, V_i) enfermé dans un récipient fermé par un piston sur lequel une force constante $\vec{F}$ s'exerce. On peut le faire passer dans un état final $(T_f, P_f = P_i, V_f)$ de manière isobare en le plaçant dans un milieu extérieur à la température T_0 . La transformation est terminée quand l'équilibre thermique entre le gaz et l'extérieur est établi donc quand : $T_f = T_0$.

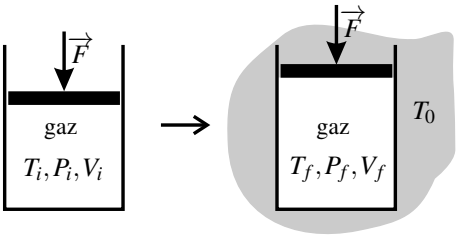


Figure 22.2 – Exemple de transformations isobare d'un échantillon de gaz.

L'équilibre mécanique du piston de surface S impose au gaz une pression

constante, $P_i = P = P_f = \frac{F}{S}$. Dans le cas représenté $T_f = T_0 > T_i$ donc $V_f > V_i$.

Si le gaz est parfait, on peut calculer le volume de l'état final V_f par l'équation d'état des gaz parfaits : $\frac{P_i V_i}{T_i} = \frac{P_f V_f}{T_f} = nR$, en notant n la quantité de gaz, d'où : $V_f = V_i \frac{T_f}{T_i} = V_i \frac{T_0}{T_i}$.

c) Transformation monobare

Une transformation **monobare** est une transformation au cours de laquelle la pression exercée par le milieu extérieur sur les parois mobiles du système garde une valeur P_0 constante :

$$P_{\text{ext}} = P_0.$$

Une transformation isobare est obligatoirement monobare si le système a une paroi mobile.

Exemple

La transformation représentée sur la figure 22.2 est monobare avec : $P_{\text{ext}} = \frac{F}{S}$.

Toute transformation d'un système en contact direct (ou bien par une paroi mobile) avec l'atmosphère, dont la pression P_{atm} est constante, subit une transformations monobare.

CHAPITRE 22 – ÉNERGIE ÉCHANGÉE PAR UN SYSTÈME AU COURS D'UNE TRANSFORMATION

d) Transformation isotherme

Une transformation est **isotherme** quand la température du système est définie tout au long de la transformation et garde une valeur constante.

En notant T_i , T_f et T la température du système respectivement dans l'état initial, dans l'état final et dans un état intermédiaire au cours de la transformation, on a : $T_i = T = T_f$.

Ces conditions sont très contraignantes et difficilement réalisables en pratique. La transformation isotherme est une transformation théorique idéale. On reviendra sur ce point dans le paragraphe 3.

e) Transformation monotherme

Une transformation **monotherme** est une transformation au cours de laquelle le milieu extérieur avec lequel le système échange de l'énergie par transfert thermique (voir paragraphe 3) a une température T_0 constante :

$$T_{\text{ext}} = T_0.$$

Une transformation isotherme est obligatoirement monotherme s'il y a un transfert thermique. Si le système est en contact thermique avec ce milieu extérieur dans l'état final, la condition d'équilibre thermique impose : $T_f = T_0$.

Exemple

Les transformations représentées sur les figures 22.1 et 22.2 sont monothermes. Au cours de ces transformations, la température du système n'est pas définie dans les états intermédiaires qui ne sont pas des états d'équilibre. Elle est définie seulement dans l'état initial et dans l'état final.

1.3 Influence du choix du système

Pour interpréter une expérience on a parfois le choix entre plusieurs systèmes possibles. Suivant le système choisi, la transformation n'a pas les mêmes propriétés.

Pour illustrer ceci on considère l'expérience représentée sur la figure 22.3 : une enceinte indéformable est séparée en deux compartiments par une cloison étanche et mobile. Dans l'état initial les deux compartiments contiennent des échantillons de gaz : les variables d'état initiales du gaz contenu dans le compartiment 1 sont (T_i, P_i, V_i, n) , pour le gaz contenu dans le compartiment 2 elles valent $(T_i, 2P_i, V_i, 2n)$ et une cale bloque la cloison mobile. On enlève la cale et on place l'enceinte dans un environnement à température T_0 . Quelles sont les variables d'état des gaz dans l'état d'équilibre final ?

Dans l'état final on doit avoir :

- équilibre thermique avec l'extérieur, donc les températures dans les deux compartiments sont : $T_{f1} = T_{f2} = T_0$,
- équilibre mécanique de la cloison mobile, donc des pressions égales dans les deux compartiments : $P_{f1} = P_{f2}$.

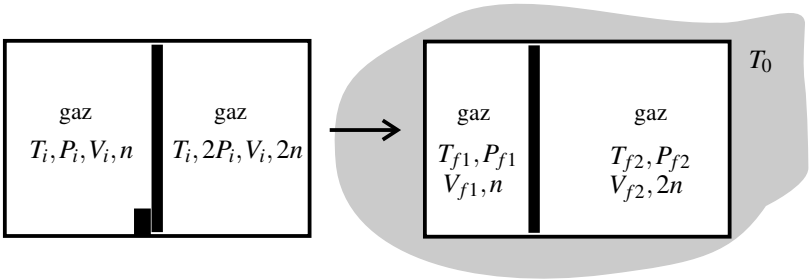


Figure 22.3 – Transformation d'un système composé.

Par ailleurs, le volume total de l'enceinte étant invariable, $V_{f1} + V_{f2} = 2V_i$.

Enfin l'équation d'état des gaz parfaits donne : $P_{f1}V_{f1} = nRT_{f1}$ et $P_{f2}V_{f2} = 2nRT_{f2}$. Étant donné que $T_{f1} = T_{f2}$ et $P_{f1} = P_{f2}$, ceci entraîne que : $V_{f2} = 2V_{f1}$.

On en déduit $V_{f1} = \frac{2}{3}V_i$ et $V_{f2} = \frac{4}{3}V_i$. Puis : $P_{f1} = P_{f2} = \frac{3}{2} \frac{nRT_0}{V_i}$.

Pour qualifier la transformation il faut préciser quel est le système considéré : soit le gaz contenu dans le compartiment 1 (système Σ_1), soit le gaz contenu dans le compartiment 2 (système Σ_2), soit tout ce qui se trouve à l'intérieur de l'enceinte (système Σ).

La transformation du système Σ est isochore et monotherme. Mais la transformation du système Σ_1 (ou Σ_2) n'a aucune propriété remarquable. Elle n'est pas monotherme car Σ_1 et Σ_2 peuvent *a priori* avoir un échange thermique à travers la paroi mobile qui les sépare (et ils n'ont pas des températures constantes). Toutefois, si cette paroi est suffisamment épaisse pour qu'on puisse négliger cet échange thermique (voir le paragraphe 3), les transformations de Σ_1 et Σ_2 sont monothermes.

2 Travail des forces de pression

Au cours d'une transformation, un système échange généralement de l'énergie avec l'extérieur. D'une manière générale :

Les échanges d'énergie d'un système sont toujours exprimés en valeur algébrique : ils sont positifs lorsque le système choisi reçoit de l'énergie et négatifs lorsqu'il en cède.

Dans ce paragraphe on s'intéresse à l'énergie reçue par un système grâce aux forces de pression. Cette énergie n'est autre que le **travail** (voir la partie *Mécanique*) de ces forces .

2.1 Expression générale du travail de la pression extérieure

a) Travail élémentaire des forces de pression dans le déplacement d'un piston

On considère dans ce paragraphe un fluide contenu dans un cylindre indéformable fermé par un piston mobile (voir figure 22.4). On prend pour système Σ l'ensemble { fluide + piston }.

CHAPITRE 22 – ÉNERGIE ÉCHANGÉE PAR UN SYSTÈME AU COURS D'UNE TRANSFORMATION

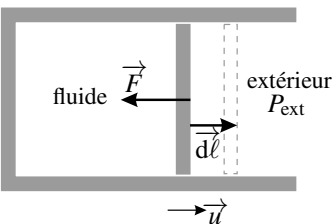


Figure 22.4 – Calcul du travail de la force de pression $\vec{F}$ s'exerçant sur le piston.

La pression P_{ext} régnant à l'extérieur du système applique au piston une force

$$\vec{F} = -P_{\text{ext}}S\vec{u},$$

où S est la surface du piston et $\vec{u}$ un vecteur unitaire perpendiculaire à la surface du piston dirigé de l'intérieur du système vers l'extérieur (voir figure 22.4). Pour un déplacement élémentaire du piston $d\vec{\ell} = d\ell\vec{u}$, le travail élémentaire de cette force s'écrit :

$$\delta W = \vec{F} \cdot d\vec{\ell} = (d\ell\vec{u}) \cdot (-P_{\text{ext}}S\vec{u}) = -P_{\text{ext}}Sd\ell.$$

Or $Sd\ell$ est le volume balayé par le piston dans son déplacement, et aussi la variation algébrique dV du volume V du système Σ (on se convainc facilement en observant la figure que $dV > 0$ pour $d\ell > 0$ et inversement). Ainsi :

$$\delta W = -P_{\text{ext}}dV. \tag{22.1}$$

⚠ Il est très important de bien respecter la notation. Le travail élémentaire se note δW avec un δ et non un d parce que *ce n'est pas la variation d'une fonction de l'état du système*. Pour la même raison le travail total sur une transformation se note W « sans rien » et surtout pas avec un Δ , la notation ΔX étant réservée à la variation d'une grandeur X fonction de l'état du système ($\Delta X = X_f - X_i$). On reviendra sur ce point au chapitre 23, page 804.

Remarque

Quel est le travail reçu par le système Σ' contenant uniquement le fluide ?

On ne peut répondre à cette question que dans le cas où le piston est « sans masse ». En effet, un système sans masse transmet les actions mécaniques. Ainsi, s'il est sans masse, le piston exerce sur le fluide une force égale à la force $\vec{F}$ qu'il subit de la part du milieu extérieur et le fluide reçoit donc de sa part exactement le travail δW donné par (22.1).

De plus, si le piston est sans masse, sa capacité thermique à volume constant est nulle (puisque $C_{V,\text{piston}} = m_{\text{piston}}c_{V,\text{piston}}$). Il « ne compte pas », ce qui fait que les systèmes Σ et Σ' sont équivalents.

Dans la pratique on se placera, sauf précision contraire, dans cette hypothèse et on considérera que le travail δW est reçu aussi bien par Σ' que par Σ .

b) Généralisation du résultat précédent

On considère de manière plus générale un système Σ soumis à une pression extérieure P_{ext} **uniforme** (c'est-à-dire identique en tous les points de la surface du système).

On suppose que la frontière du système se déforme de manière infinitésimale, passant de $\mathcal{S}$ à $\mathcal{S}'$: tout point M de $\mathcal{S}$ se déplace en un point M' de $\mathcal{S}'$ et on note $\vec{d\ell} = \overrightarrow{MM'}$ le petit déplacement de ce point (voir figure 22.5) . On note $\vec{dS}$ le vecteur surface élémentaire de $\mathcal{S}$ en M , vecteur dirigé vers l'extérieur du système.

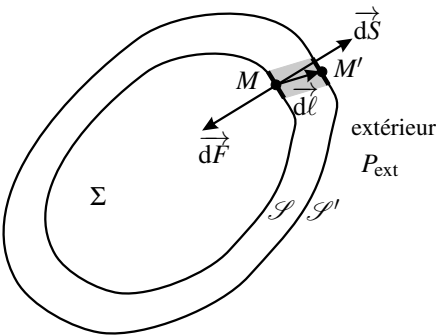


Figure 22.5 – Calcul du travail de la force de pression, dans le cas général. Le volume en gris est égal à $\vec{d\ell} \cdot \vec{dS}$.

La force exercée par l'extérieur sur l'élément de surface dS est $\vec{dF} = -P_{\text{ext}}\vec{dS}$ et son travail élémentaire dans ce déplacement est :

$$\delta^2 W = \vec{dF} \cdot \vec{d\ell} = -P_{\text{ext}}\vec{dS} \cdot \vec{d\ell}.$$



Le 2 en exposant dans la notation $\delta^2 W$ est là pour rappeler que ce travail est doublement élémentaire : parce que le vecteur surface $\vec{dS}$ est élémentaire et parce que le déplacement $\vec{d\ell}$ est élémentaire.

Ainsi le travail des forces de pression s'exerçant sur toute la surface du système est :

$$\delta W = \int_{M \in \mathcal{S}} \delta^2 W = -P_{\text{ext}} \int_{M \in \mathcal{S}} \vec{dS} \cdot \vec{d\ell} = -P_{\text{ext}} dV,$$

où dV est le volume compris entre $\mathcal{S}$ et $\mathcal{S}'$ soit la variation élémentaire de volume du système. On retrouve l'expression (22.1) qui est donc valable dans un cas plus général.

Dans ce calcul, pour sortir P_{ext} de l'intégrale, on a utilisé l'hypothèse selon laquelle cette pression est uniforme, donc identique en tous les points de la surface du système. Cette hypothèse sera supposée vérifiée dans toute la suite du chapitre sauf dans le paragraphe 2.2 ci-dessous.

CHAPITRE 22 – ÉNERGIE ÉCHANGÉE PAR UN SYSTÈME AU COURS D’UNE TRANSFORMATION

c) Travail des forces de pression reçu par un système dans une transformation

Lors d’une transformation d’un système entre un état initial i et un état final f , le travail des forces de la pression extérieure sur le système est :

$$W = \int_i^f \delta W \quad \text{soit, d'après (22.1) :} \quad W = - \int_i^f P_{\text{ext}} dV. \quad (22.2)$$

Il faut bien comprendre la signification physique du signe « $-$ » dans cette formule. Il s’agit du travail *reçu* par le système constitué par le gaz. Il est positif lorsque le volume du gaz diminue ($dV < 0$) : il faut fournir un travail pour comprimer un gaz dans un volume plus petit (on expérimente très clairement ce résultat lorsqu’on gonfle un pneu de vélo). Inversement, un gaz qui se détend en augmentant de volume ($dV > 0$) reçoit un travail négatif de l’extérieur, donc en fait fournit du travail.

2.2 Cas particulier d’un fluide en écoulement

Lorsque la condition d’application de la formule (22.1) n’est pas vérifiée il faut revenir à la définition du travail d’une force vue en mécanique. Dans ce paragraphe on étudie le cas important où le système est un volume de fluide en écoulement dans une conduite.

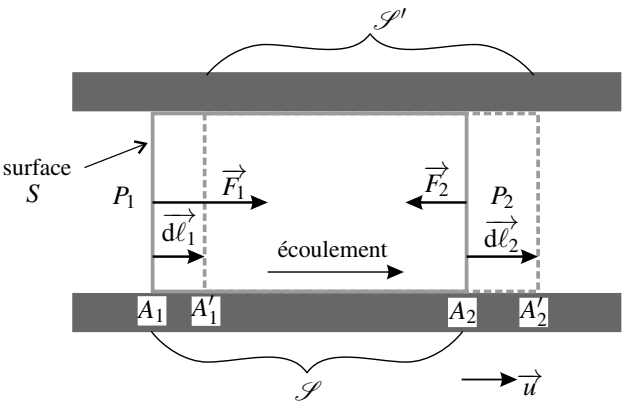


Figure 22.6 – Calcul du travail des forces de pression, dans le cas d’un écoulement dans une conduite.

Soit un fluide s’écoulant dans une conduite dont la section a une surface S . Dans ce fluide, on isole par l’esprit un système fermé Σ constitué par le fluide contenu dans la surface $\mathcal{S}$ comprise entre les sections A_1 et A_2 de la conduite à l’instant t (voir figure 22.6). A l’instant $t' = t + dt$ le système Σ (donc le même fluide) est contenu dans la surface $\mathcal{S}'$ comprise entre les sections A_1' et A_2' .

On appelle $\overrightarrow{d\ell_1} = \overrightarrow{A_1 A_1'}$ et $\overrightarrow{d\ell_2} = \overrightarrow{A_2 A_2'}$ les déplacements élémentaires entre t et $t + dt$ des deux sections délimitant le système (voir figure 22.6).

La pression en A_1 est égale à P_1 et elle est égale à P_2 en A_2 .

On note S la surface de la section de la conduite. La force de pression appliquée à Σ sur la section A_1 s'écrit : $\vec{F}_1 = P_1 S \vec{u}$, où $\vec{u}$ est le vecteur unitaire dans le sens de l'écoulement. Elle fournit dans le déplacement considéré le travail :

$$\delta W_1 = \vec{F}_1 \cdot d\vec{\ell}_1 = (P_1 S \vec{u}) \cdot (d\ell_1 \vec{u}) = P_1 S d\ell_1 = P_1 dV_1,$$

où dV_1 est le volume compris entre la section A_1 et A'_1 , volume balayé par la surface limitant le système. Ce travail est positif : le fluide en amont pousse le fluide de Σ .

La force de pression appliquée à Σ sur la section A_2 s'écrit : $\vec{F}_2 = -P_2 S \vec{u}$. Elle fournit dans le déplacement considéré le travail :

$$\delta W_2 = \vec{F}_2 \cdot d\vec{\ell}_2 = (-P_2 S \vec{u}) \cdot (d\ell_2 \vec{u}) = -P_2 S d\ell_2 = -P_2 dV_2,$$

où dV_2 est le volume compris entre la section A_2 et A'_2 , volume balayé par la surface limitant le système. Ce travail est négatif : le fluide en aval repousse le fluide de Σ .

Au total, le travail des forces de pressions est dans ce cas :

$$\delta W = P_1 dV_1 - P_2 dV_2. \quad (22.3)$$

On ne peut pas appliquer la formule (22.1) parce que la pression extérieure a ici deux valeurs différentes sur la surface du système.

Remarque

La variation de volume du système est $dV = dV_2 - dV_1$. Ainsi, si on avait $P_1 = P_2 = P_{\text{ext}}$, on aurait $\delta W = P_{\text{ext}}(dV_1 - dV_2) = -P_{\text{ext}}dV$ et on retrouverait bien la formule (22.1).

2.3 Travail des forces de pression dans deux cas particuliers

Dans toute la suite on suppose que la pression P_{ext} appliquée par l'extérieur sur le système est uniforme (identique sur toute la surface du système), ce qui permet d'appliquer la formule (22.1).

a) Cas d'une transformation isochore

Pour une transformation isochore, le volume ne variant pas $dV = 0$, donc le travail élémentaire est nul : $\delta W = -P_{\text{ext}}dV = 0$.

Au cours d'une transformation isochore le travail des forces de pression est nul.

b) Cas d'une transformation monobare

Dans le cas d'une transformation monobare, $P_{\text{ext}} = P_0$ où P_0 est une constante.

Le travail élémentaire des forces de pression s'écrit alors : $\delta W = -P_{\text{ext}}dV = -P_0dV$.

Sur la transformation complète entre l'état initial i et l'état final f , le volume varie entre le volume initial V_i et le volume final V_f , et le travail des forces de pression est :

$$W = \int_i^f \delta W = \int_i^f -P_0 dV = -P_0 \int_{V_i}^{V_f} dV = -P_0(V_f - V_i).$$

CHAPITRE 22 – ÉNERGIE ÉCHANGÉE PAR UN SYSTÈME AU COURS D'UNE TRANSFORMATION

Le travail des forces de pression au cours d'une transformation monobare telle que $P_{\text{ext}} = P_0$ est :

$$W = -P_0 \Delta V, \quad (22.4)$$

où $\Delta V = V_f - V_i$ est la variation de volume du système.

c) Cas d'une transformation isobare

Si la transformation est isobare et que le volume varie, elle est obligatoirement monobare et $P = P_{\text{ext}} = P_0$. Le résultat précédent s'applique donc :

Le travail des forces de pression au cours d'une transformation isobare est :

$$W = -P(V_f - V_i) = -P \Delta V. \quad (22.5)$$

2.4 Travail des forces de pression dans le cas d'une transformation mécaniquement réversible

a) Expression du travail

Une transformation **mécaniquement réversible** est une transformation au cours de laquelle la pression P du système est définie à chaque instant et toujours égale à la pression extérieure, soit :

$$P = P_{\text{ext}}.$$

Dans le cas d'une transformation mécaniquement réversible le travail des forces de pression s'exerçant sur le système s'écrit :

$$W = - \int_i^f P dV. \quad (22.6)$$

où P est la pression dans le système.

b) Interprétation géométrique

Le travail des forces de pression s'interprète géométriquement dans le diagramme de Clapeyron, diagramme (P, V) avec la pression du système P en ordonnée et son volume V en abscisse.

On représente la courbe suivie par le système dans sa transformation entre le point (P_i, V_i) représentant l'état initial et le point (P_f, V_f) représentant l'état final (voir figure 22.7). Sauf exception, cette courbe ne passe pas deux fois par la même valeur de V et c'est donc la représentation d'une fonction $P(V)$. Le travail $W = - \int_{V_i}^{V_f} P(V) dV$ est au signe près, d'après une propriété classique de l'intégrale, l'aire $\mathcal{A}$ comprise entre la courbe de la fonction $P(V)$ et l'axe des abscisses, aire grisée sur la figure 22.7(a).

La valeur absolue du travail des forces de pression est égale à l'aire $\mathcal{A}$ comprise entre la courbe représentant la transformation du système dans le diagramme de Clapeyron et l'axe des abscisses.

Le travail algébrique est :

- $W = +\mathcal{A}$ si $V_f < V_i$ (le gaz reçoit du travail si son volume diminue),
- $W = -\mathcal{A}$ si $V_f > V_i$ (le gaz cède du travail si son volume augmente).

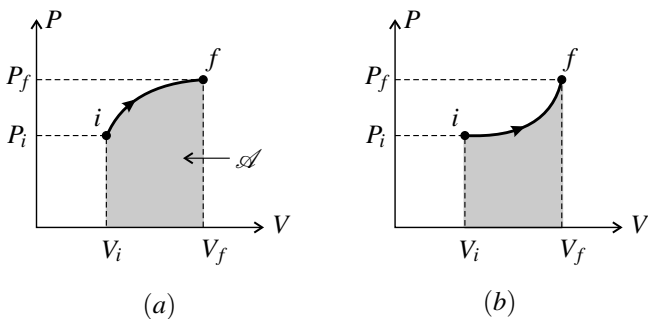


Figure 22.7 – Interprétation géométrique du travail des forces de pression.

D'autre part, les figures 22.7(a) et 22.7(b) montrent deux transformations différentes menant de l'état i à l'état f et représentées par deux courbes différentes. Les aires sous ces deux courbes sont différentes, donc le système ne reçoit pas le même travail dans les deux transformations qui vont pourtant du même état initial au même état final. On constate donc que :

Le travail de pression dépend de la transformation entre l'état initial et l'état final.

c) Travail reçu par le système au cours d'une évolution cyclique

On s'intéresse à une transformation cyclique du système au cours de laquelle il passe d'un état A à un état B , puis revient à l'état A par un autre chemin : $A \rightarrow B \rightarrow A$.

On suppose pour fixer les idées que $V_B > V_A$. Le travail reçu par le système est négatif lors de la transformation $A \rightarrow B$ soit $W_{A \rightarrow B} = -\mathcal{A}$ où $\mathcal{A}$ est l'aire sous la courbe suivie par le point représentatif du système dans le diagramme de Clapeyron (voir figure 22.8(a)). Le travail est positif lors de la transformation $B \rightarrow A$ soit $W_{B \rightarrow A} = +\mathcal{A}'$ en notant $\mathcal{A}'$ l'aire sous la courbe suivie (voir figure 22.8(b)). Le travail algébrique reçu par le système sur le cycle,

$$W_{\text{cycle}} = W_{A \rightarrow B} + W_{B \rightarrow A} = \mathcal{A}' - \mathcal{A},$$

n'est pas nul. Il est égal en valeur absolue à l'aire $\mathcal{A}_{\text{cycle}}$ entourée par le chemin du cycle dans le diagramme de Clapeyron (voir figure 22.8(c)).

Sur la figure 22.8(c) le cycle est décrit dans le sens des aiguilles d'une montre et $W < 0$. Il aurait été positif si le cycle avait été parcouru dans l'autre sens.

CHAPITRE 22 – ÉNERGIE ÉCHANGÉE PAR UN SYSTÈME AU COURS D'UNE TRANSFORMATION

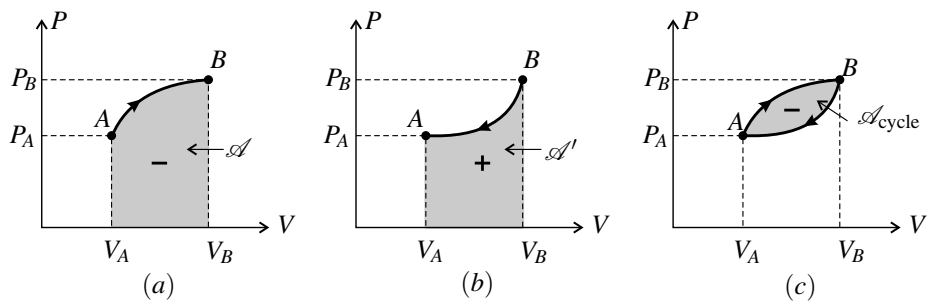


Figure 22.8 – Travail des forces de pression reçu par un système au cours d'un cycle.

Le travail W des forces de pression reçu par un système au cours d'un cycle est négatif lorsque le cycle est décrit dans le sens horaire dans le diagramme de Clapeyron. Dans ce cas le système fournit du travail. Un tel cycle est appelé **cycle moteur**.

Le travail W des forces de pression reçu par un système au cours d'un cycle est positif lorsque le cycle est décrit dans le sens trigonométrique dans le diagramme de Clapeyron. Dans ce cas le système reçoit du travail. Un tel cycle est appelé **cycle récepteur**.

Dans les deux cas, la valeur absolue du travail échangé par le système est égale à l'aire de la surface délimitée par le cycle : $|W| = \mathcal{A}_{\text{cycle}}$.

Remarque

Ces résultats montrent bien que le travail n'est pas la variation d'une fonction d'état X entre l'état initial et l'état final puisque :

- toute variation $\Delta X = X_f - X_i$ entre un état initial i et un état final f est indépendante du chemin suivi entre ces deux états, ce qui n'est pas le cas du travail ;
- lors d'un cycle, la variation de toute fonction d'état est nulle ($\Delta X = X_A - X_A = 0$), ce qui n'est pas le cas du travail.

d) Exemples de calculs pour un échantillon de gaz parfait

Dans ce paragraphe on calcule l'intégrale (22.6) pour un système constitué par une quantité de matière n de gaz parfait (plus éventuellement une paroi mobile). Les formules ne sont pas à retenir, mais il faut savoir refaire ces calculs.

Transformation isotherme mécaniquement réversible Dans le cas d'une transformation isotherme la température T du système reste constamment égale à $T_0 = T_i = T_f$. On peut donc écrire en appliquant l'équation d'état du gaz parfait :

$$PV = nRT = nRT_0 \quad \text{d'où} \quad P = \frac{nRT_0}{V}.$$

Le travail des forces de pression s'écrit alors :

$$W = - \int_i^f P \, dV = - \int_i^f \frac{nRT_0}{V} \, dV = -nRT_0 \int_{V_i}^{V_f} \frac{dV}{V} = -nRT_0 \ln \frac{V_f}{V_i}.$$

Par ailleurs, $P_i V_i = P_f V_f = nRT_0$ donc : $\frac{V_f}{V_i} = \frac{P_i}{P_f}$. Ainsi :

$$W = -nRT_0 \ln \frac{V_f}{V_i} = nRT_0 \ln \frac{P_f}{P_i}. \tag{22.7}$$

Transformation polytropique mécaniquement réversible Une transformation polytropique est une transformation au cours de laquelle la pression (définie à chaque instant) et le volume vérifient une relation de la forme $PV^k = \text{constante}$, où k est un exposant dépendant de la transformation. Comme une augmentation de volume s'accompagne naturellement d'une diminution de pression, l'exposant k est positif.

La pression P au cours de la transformation vérifie donc : $PV^k = P_i V_i^k$ soit $P = P_i \left(\frac{V_i}{V}\right)^k$. Le travail reçu par le système s'écrit ainsi :

$$W = - \int_i^f P \, dV = - \int_i^f P_i \left(\frac{V_i}{V}\right)^k \, dV = -P_i V_i^k \int_{V_i}^{V_f} \frac{dV}{V^k} = -P_i V_i^k \frac{V_f^{1-k} - V_i^{1-k}}{1-k}.$$

En utilisant $P_f V_f^k = P_i V_i^k$ et l'équation d'état du gaz parfait, on peut mettre ce résultat sous deux formes simples :

$$W = \frac{1}{k-1} (P_f V_f - P_i V_i) = \frac{nR}{k-1} (T_f - T_i). \tag{22.8}$$

3 Transfert thermique

3.1 Définition

Un système thermodynamique peut recevoir de l'énergie sans l'intervention d'une action mécanique mesurable à l'échelle macroscopique. Ce transfert d'énergie complémentaire du travail mécanique s'appelle **transfert thermique**.

La quantité d'énergie échangée entre un système Σ et l'extérieur par transfert thermique est notée Q . Elle est algébrique et, par convention, positive lorsque Σ reçoit de l'énergie. Le transfert thermique Q est parfois appelée **chaleur**. Il se mesure en joules.

3.2 Les trois modes de transfert thermique (PTSI)

Le transfert thermique s'opère entre deux systèmes en contact si leurs températures sont différentes. Le système dont la température est la plus élevée cède de l'énergie au système dont la température est la plus basse. Ce transfert d'énergie se passe à l'échelle microscopique

CHAPITRE 22 – ÉNERGIE ÉCHANGÉE PAR UN SYSTÈME AU COURS D’UNE TRANSFORMATION

et il n’est perceptible à l’échelle macroscopique que par la transformation des systèmes qu’il provoque : variation de température, changement d’état...

Il existe trois **modes de transfert thermique**, schématisés sur la figure 22.9 :

1. La **conduction thermique** est le mode de transfert thermique entre deux systèmes séparés par un milieu matériel immobile, par exemple une paroi solide. Le transfert d’énergie résulte des collisions entre les particules microscopiques constituant les systèmes et la paroi. Ces particules sont animées d’un mouvement d’agitation thermique quelle que soit la nature (solide, liquide ou gaz) des systèmes. Les particules du système ayant la température la plus élevée (« système chaud ») ont une énergie cinétique d’agitation thermique supérieure à celle du système ayant une température plus basse (« système froid »). Lors des chocs, les premières cèdent de l’énergie aux particules de la paroi et les deuxièmes reçoivent de l’énergie de la paroi.
2. La **convection thermique** met en jeu un fluide en mouvement. Le fluide passe d’un système à l’autre, reçoit de l’énergie du système chaud et cède de l’énergie au système froid.
3. Le **rayonnement thermique** met en jeu les ondes électromagnétiques émises par les particules microscopiques des systèmes à cause de leur mouvement d’agitation thermique (quelle que soit la nature des systèmes). Les photons émis par chacun des systèmes sont reçus par l’autre qui en absorbe une partie. Il y a ainsi transfert d’énergie dans les deux sens, mais du fait que le système chaud émet plus d’énergie que le système froid, le transfert d’énergie global se fait du système chaud vers le système froid.

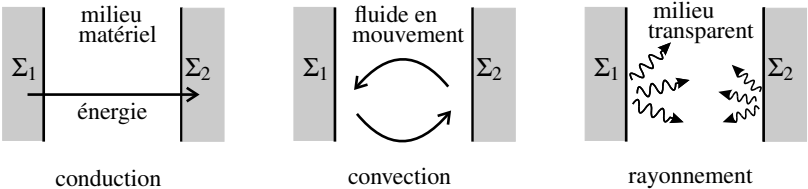


Figure 22.9 – Les trois modes de transfert thermique entre un système Σ_1 de température T_1 et un système Σ_2 de température $T_2 < T_1$.

Exemple

Le transfert thermique entre l’intérieur d’une habitation et l’extérieur est :

- conductif à travers un mur ou une fenêtre fermée ;
- convectif à travers une fenêtre ouverte,
- radiatif quand le rayonnement du soleil entre à travers une vitre.

La cuisson d’un plat dans un four électrique classique est due à un transfert thermique entre la résistance chauffante du four et le plat qui passe pour l’essentiel par le rayonnement et pour une petite part par la conduction thermique à travers l’air. Dans un four dit « à chaleur tournante » un ventilateur provoque un mouvement d’air à l’intérieur du four qui ajoute un transfert par convection très efficace.

3.3 Transformation adiabatique

a) Définition

Une transformation **adiabatique** est une transformation au cours de laquelle le système n'a aucun échange d'énergie par transfert thermique : $Q = 0$.

Dans le cas d'une transformation adiabatique, la température du système dans l'état final n'est pas déterminée par une condition d'équilibre thermique, puisque le système n'est en contact thermique avec aucun autre système.

b) Réalisation pratique

La définition théorique de la transformation adiabatique correspond à une idéalisation dont la réalité ne peut que s'approcher. Pour réaliser une transformation adiabatique, on entoure le système d'un matériau dit « isolant thermique », à travers lequel la conduction thermique est difficile.

Exemple

La vase de Dewar est un récipient à double paroi de verre entre lesquelles il y a de l'air sous très faible pression pour réduire fortement le transfert thermique conductif. De plus la paroi de verre est métallisée pour réduire le transfert thermique par rayonnement. Les calorimètres utilisés en travaux pratiques (voir chapitre suivant), bien que de fabrication moins coûteuse, sont conçus sur le même principe.

L'efficacité d'un tel dispositif est limitée dans le temps et le système finit toujours par être en équilibre thermique avec l'extérieur. Le rôle de l'isolation thermique est d'augmenter fortement le temps caractéristique d'établissement de l'équilibre thermique. Celui-ci peut facilement devenir très long.

Exemple

Le temps caractéristique d'établissement du transfert thermique conductif à travers une paroi d'épaisseur e est $\tau = \frac{e^2}{a}$ où a est la diffusivité thermique du matériau constituant la paroi. La diffusivité thermique du liège est $a = 3,0 \cdot 10^{-7} \text{ m}^2 \cdot \text{s}^{-1}$, donc pour $e = 2,0 \text{ cm}$ on trouve $\tau = 1,3 \cdot 10^3 \text{ s} \simeq 22 \text{ min}$.

Dans ce cas, l'équilibre mécanique du système s'établit beaucoup plus rapidement que l'équilibre thermique et on peut dire que le système, au moment où il atteint l'équilibre mécanique, a subi une transformation adiabatique.

Une transformation rapide peut être considérée comme adiabatique.

CHAPITRE 22 – ÉNERGIE ÉCHANGÉE PAR UN SYSTÈME AU COURS D'UNE TRANSFORMATION

Exemple

Le cylindre d'un moteur est en acier de diffusivité thermique $a = 1,5 \cdot 10^{-5} \text{ m}^2 \cdot \text{s}^{-1}$ et d'épaisseur $e \simeq 0,5 \text{ cm}$. Le temps caractéristique pour les échanges thermiques est $\tau = \frac{e^2}{a} \simeq 167 \text{ s}$. Pour un moteur ayant une vitesse de rotation de l'ordre de 1000 tours par minute, le cycle de transformation des gaz dure environ 0,06 s et les échanges thermiques n'ont pas le temps de se faire. On peut faire une modélisation adiabatique.

3.4 Notion de thermostat

a) Définition

Un **thermostat** est un système thermodynamique dont la température T_0 ne varie pas, même s'il échange de l'énergie (sous forme de transfert thermique ou de travail).

b) Intérêt pratique de la notion de thermostat

Soient deux systèmes monophasés Σ et Σ_0 échangeant de l'énergie par transfert thermique et de capacités thermiques à volume constant C_V et C_{V0} telles que $C_{V0} \gg C_V$. Lorsqu'un transfert thermique algébrique Q est fourni par Σ_0 et reçue par Σ , dans une transformation où les volumes des deux systèmes sont constants, cela induit des variations ΔT et ΔT_0 des températures respectives de Σ et Σ_0 telles que :

$$Q = \Delta U_{\Sigma} = C_V \Delta T \quad \text{et} \quad -Q = \Delta U_{\Sigma_0} = C_{V0} \Delta T_0,$$

d'après le premier principe, appliqué successivement aux deux systèmes. Il en résulte que :

$$|\Delta T_0| = \frac{C_V}{C_{V0}} |\Delta T| \ll |\Delta T|.$$

Ainsi le système Σ_0 peut être considéré comme un thermostat dans son interaction avec Σ .

Lorsque deux systèmes échangeant de l'énergie par transfert thermique ont des capacités thermiques d'ordres de grandeur très différents, on peut modéliser le système ayant la plus grande capacité thermique par un thermostat.

Exemple

Les centrales nucléaires sont construites à proximité d'un fleuve dont elles utilisent l'eau pour le refroidissement et qui peut être modélisé par un thermostat.

On peut aussi réaliser un thermostat de petite taille. Au laboratoire on utilise un mélange eau-glace. En effet un tel mélange est à température fixe, voisine de 273 K sous la pression atmosphérique. Quand ce système reçoit (resp. cède) un transfert thermique $Q > 0$, cela provoque la fusion d'une partie de la glace (resp. la solidification d'une partie de l'eau liquide) mais la température ne change pas. Bien sûr, il ne faut pas que la quantité d'énergie Q soit trop grande.

3.5 Retour sur les transformations monotherme et isotherme

On peut reformuler la définition de la transformation monotherme :

Une transformation est monotherme si le système échange de l'énergie par transfert thermique avec un et un seul thermostat.

Pour qu'une transformation soit **isotherme** il faut que la température du système ne varie pas. Or dans la plupart des cas (l'exception étant le mélange diphasé évoqué ci-dessus) tout apport d'énergie au système tend à faire varier sa température.

La réalisation d'une transformation isotherme nécessite donc un contrôle de la température que l'on obtient en mettant le système en contact avec un thermostat. Il faut que les échanges thermiques entre le système et le thermostat soient faciles. Ceux-ci doivent donc être séparés par une **paroi diathermane**, c'est-à-dire perméable à la chaleur. De plus l'évolution du système doit être suffisamment lente pour que les échanges thermiques aient le temps de s'établir et assurent le maintien de la température T du système à la même valeur que la température T_0 du thermostat. La plupart des transformations isothermes sont des transformations lentes, au cours desquelles le système est constamment en équilibre thermique avec un thermostat.

3.6 Choix d'un modèle : adiabatique ou isotherme ?

Les transformations isotherme et adiabatique sont deux transformations idéales aux caractères diamétralement opposés : la transformation adiabatique suppose des échanges thermiques nuls alors que la transformation isotherme n'est possible (dans la majeure partie des cas) que s'il y a des échanges thermiques avec un thermostat.

Une transformation réelle pourra se rapprocher de l'une ou l'autre des ces deux transformations limites et il importe de savoir choisir la bonne modélisation.

Si la transformation est très rapide ou si les parois délimitant le système sont très épaisses on pourra faire une modélisation adiabatique.
Si la transformation est lente et que le système est en contact avec un thermostat on pourra faire une modélisation isotherme.

Exemple

Un gaz est contenu dans un récipient fermé par un piston de surface S sur lequel on exerce une force $\vec{F}$. On impose ainsi, par la condition d'équilibre mécanique du piston, une pression $P = \frac{F}{S}$ au gaz. Dans l'état initial $\vec{F} = \vec{F}_i$ et dans l'état final $\vec{F} = \vec{F}_f$.

Dans une première expérience (figure 22.10) le récipient a des parois fines, diathermanes. On augmente lentement la force $\vec{F}$ provoquant une descente progressive du piston et laissant le gaz s'équilibrer thermiquement à chaque instant avec le milieu extérieur, de température T_0 . On peut modéliser la compression par une *transformation isotherme*.

CHAPITRE 22 – ÉNERGIE ÉCHANGÉE PAR UN SYSTÈME AU COURS D’UNE TRANSFORMATION

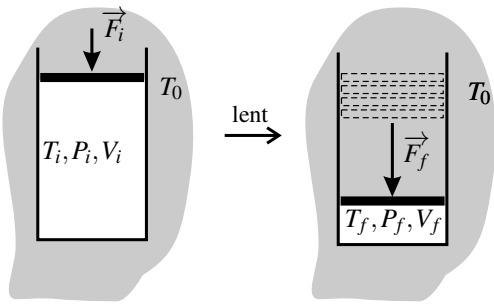


Figure 22.10 – Compression isotherme.

Dans ce cas, on peut trouver tous les paramètres d’état du gaz dans l’état final : $P_f = \frac{F_f}{S}$, $T_f = T_0$ et d’après l’équation d’état du gaz parfait : $V_f = V_i \frac{P_i}{P_f} \frac{T_0}{T_i} = V_i \frac{F_i}{F_f} \frac{T_0}{T_i}$.

Dans une autre expérience (figure 22.11) les parois du récipient sont épaisses. On augmente la force $\vec{F}$ trop rapidement pour que l’échange thermique entre le gaz et l’extérieur puisse se faire. On peut modéliser la compression comme une *transformation adiabatique*. Dans ce cas, l’équation d’état ne suffit pas pour trouver l’état final car on n’a pas de renseignement sur la température T_f .

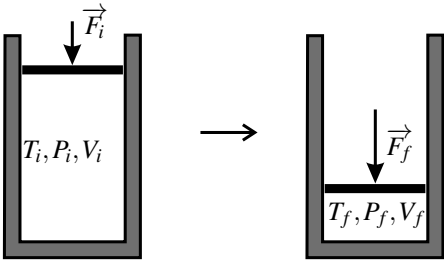


Figure 22.11 – Compression adiabatique.

Enfin, l’équilibre mécanique étant atteint beaucoup plus rapidement que l’équilibre thermique, on peut imaginer une transformation à la fois *adiabatique et mécaniquement réversible*, suffisamment rapide pour qu’il n’y ait pas de transfert thermique, mais suffisamment lente pour qu’il y ait équilibre mécanique à chaque instant. L’équation d’état ne suffit pas dans ce cas pour trouver l’état final.

SYNTHÈSE

SAVOIRS

- conditions d'équilibre d'un système thermodynamique
- différents types de transformations
- interprétation géométrique du travail des forces de pression
- définition du transfert thermique
- les trois modes de transfert thermique

SAVOIR-FAIRE

- définir un système
- exploiter les conditions d'équilibre pour trouver l'état d'équilibre final
- calculer le travail des forces de pression
- distinguer les trois modes de transfert thermique
- identifier un thermostat
- choisir un modèle limite entre isotherme et adiabatique

MOTS-CLÉS

- | | | |
|----------------------------|-----------------------|-------------------------|
| • transformation thermody- | • isobare | • conduction thermique |
| namique | • monobare | • convection thermique |
| • isochore | • adiabatique | • rayonnement thermique |
| • isotherme | • travail | |
| • monotherme | • transfert thermique | |

CHAPITRE 22 – ÉNERGIE ÉCHANGÉE PAR UN SYSTÈME AU COURS D'UNE TRANSFORMATION

S'ENTRAÎNER

22.1 Questions qualitatives (★)

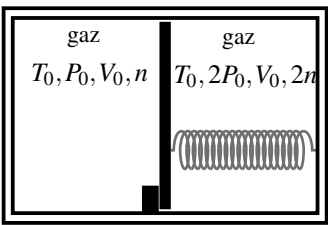
1. On dépose un glaçon sortant d'un congélateur dans une coupelle et on l'abandonne à l'air libre. Quel est l'état final ? Dans cette transformation le système constitué par le glaçon reçoit-il du travail ? du transfert thermique ?
2. En hiver, un ballon de baudruche initialement à l'équilibre dans un lieu chauffé est apporté à l'extérieur. Le système constitué par le ballon et l'air qu'il contient reçoit-il du travail ? du transfert thermique ?
3. Comment peut-on qualifier les transformations précédentes ?

22.2 Transfert thermique et cuisson (PTSI) (★)

1. Citer un procédé de cuisson dans lequel l'aliment reçoit du transfert thermique majoritairement :
 - par conduction,
 - par convection,
 - par rayonnement.
2. On cuit un œuf en le plongeant dans une casserole d'eau bouillante. Quel est le mode de transfert thermique dominant pour le système {œuf} ? pour le système {œuf + eau} ? pour le système {jaune d'œuf} ?

22.3 Recherche d'un état final (★)

Une enceinte indéformable aux parois calorifugées est séparée en deux compartiments par une cloison étanche de surface S , mobile, diathermane et reliée à un ressort de constante de raideur k . Les deux compartiments contiennent chacun un gaz parfait. Dans l'état initial, le gaz du compartiment 1 est dans l'état (T_0, P_0, V_0) , le gaz du compartiment 2 dans l'état $(T_0, 2P_0, V_0, 2n)$, une cale bloque la cloison mobile et le ressort est au repos. On enlève la cale et on laisse le système atteindre un état d'équilibre.



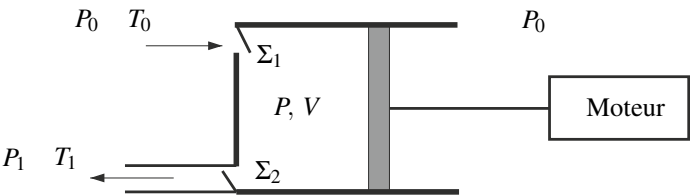
1. Décrire l'évolution du système.
2. Écrire cinq relations faisant intervenir certaines des six variables d'état : V_1 , V_2 (volumes finaux des deux compartiments), P_1 , P_2 (pressions finales dans les deux compartiments), T_1 , T_2 (températures finales dans les deux compartiments).

APPROFONDIR

22.4 Étude d'un compresseur, d'après ESIM 2000 (★★)

Le problème étudie le compresseur d'un moteur à air comprimé (celui d'un marteau-piqueur, par exemple). L'air est assimilé à un gaz parfait de masse molaire $M = 29 \text{ g.mol}^{-1}$, de capacité thermique massique à pression constante $c_p = 1,00 \text{ kJ.kg}^{-1}.\text{K}^{-1}$ et de rapport des capacités thermiques à pression et à volume constants $\gamma = 1,4$. La constante des gaz parfaits est $R = 8,314 \text{ J.K}^{-1}.\text{mol}^{-1}$.

L'air est aspiré dans les conditions atmosphériques, sous la pression $P_0 = 1 \text{ bar}$ et à la température $T_0 = 290 \text{ K}$, jusqu'au volume V_m , puis comprimé jusqu'à la pression P_1 , où il occupe le volume V_1 , et refoulé à la température T_1 dans un milieu où la pression est $P_1 = 6 \text{ bar}$. Bien que le mécanisme réel d'un compresseur soit différent, on suppose que celui-ci fonctionne comme une pompe à piston, qui se compose d'un cylindre, d'un piston coulissant entraîné par un moteur et de deux soupapes.



- La soupape d'entrée Σ_1 est ouverte si la pression P dans le corps de pompe est inférieure ou égale à la pression atmosphérique P_0 .
- La soupape de sortie Σ_2 est ouverte si P est supérieure à P_1 .
- Le volume V du corps de pompe est compris entre 0 et V_m .
- À chaque cycle (chaque aller et retour du piston), la pompe aspire et refoule une mole d'air.

- a. Tracer sur un diagramme de Watt (P en ordonnée, V en abscisse) l'allure de la courbe représentant un aller et un retour du piston. Indiquer le sens de parcours par une flèche.
 - b. Montrer que le travail de l'air situé à droite du piston est nul sur un aller-retour.
 - c. Montrer que le travail fourni par le moteur qui actionne le piston est égal à l'aire d'une surface sur le diagramme. On supposera que le mouvement est assez lent pour que l'évolution soit mécaniquement réversible.

2. Pendant la phase de compression, l'air suit une loi polytropique $PV^k = \text{cste}$; il sort du compresseur à la température $T_1 = 391 \text{ K}$. Trouver la valeur de k .
3. Exprimer le travail mécanique W_{moteur} fourni par le moteur pendant un aller-retour en fonction de R , n , k , T_1 et T_0 .
4. Le débit massique de l'air dans le compresseur est $D_m = 0,013 \text{ kg.s}^{-1}$. Calculer la puissance $\mathcal{P}_{\text{moteur}}$ fournie par le moteur.

CHAPITRE 22 – ÉNERGIE ÉCHANGÉE PAR UN SYSTÈME AU COURS D'UNE TRANSFORMATION

CORRIGÉS

22.1 Questions qualitatives

1. Dans l'état final le glaçon a fondu : il est devenu de l'eau liquide à la température de la pièce. Comme il est au départ à température plus basse que la pièce, il a reçu du transfert thermique. Le glaçon ayant un volume massique plus grand que l'eau liquide, le volume du système a diminué, donc il a reçu du travail des forces de pression atmosphérique.
2. Le système ballon, initialement à température plus élevée que l'air extérieur, cède du transfert thermique. Dans l'état final il a une température égale à la température de l'air atmosphérique qui est un thermostat. De ce fait, le gaz contenu dans le ballon a diminué de volume, donc le système ballon a reçu du travail des forces de pression atmosphérique.
3. L'atmosphère est un thermostat et un pressio-stat (système de pression constante). Les deux transformations sont monobares et monothermes. Comme elles sont lentes et qu'au départ les deux systèmes sont à la pression atmosphérique, elles sont aussi isobares.

22.2 Transfert thermique et cuisson (PTSI)

1. On peut citer :
 - pour la conduction : la cuisson dans une casserole (le système constitué par le contenu de la casserole reçoit un transfert thermique conductif traversant le fond de la casserole) ;
 - pour la convection : la cuisson vapeur et la cuisson au bain-marie dans lesquelles le système constitué par l'aliment reçoit le transfert thermique de la vapeur chaude ou de l'eau chaude en mouvement ;
 - pour le rayonnement : la cuisson dans un four traditionnel, la cuisson au barbecue.
2. Pour le système {œuf} c'est un transfert thermique convectif, pour le système {œuf + eau} un transfert thermique conductif et pour le système {jaune d'œuf} un transfert thermique conductif (quand l'œuf est dur).

Remarque

Si la casserole est posée sur une plaque halogène, le système constitué par la casserole et son contenu reçoit du transfert thermique essentiellement par rayonnement. Ainsi, le mode de transfert thermique mis en jeu dans une situation donnée dépend du système considéré.

22.3 Recherche d'un état final

1. La paroi mobile n'est pas à l'équilibre : la pression étant plus forte dans le compartiment 2, il est poussé vers la gauche.
2. On peut écrire les conditions suivantes :
 1. conservation du volume total : $V_1 + V_2 = 2V_0$,
 2. équation d'état pour le gaz du compartiment 1 : $\frac{P_1 V_1}{T_1} = \frac{P_0 V_0}{T_0}$;
 3. équation d'état pour le gaz du compartiment 2 : $\frac{P_2 V_2}{T_2} = \frac{2P_0 V_0}{T_0}$;

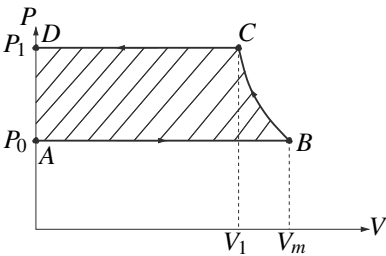
4. équilibre thermique de part et d'autre de la cloison diathermane : $T_1 = T_2$;
5. équilibre mécanique de la cloison mobile : $P_1 = P_2 - k \frac{V_2 - V_0}{S}$, en tenant compte de la force du ressort qui est allongé.



Il n'y a aucune raison pour que T_1 et T_2 soient égales à T_0 . L'enceinte étant calorifugée et indéformable, le système global est isolé et on obtient, en appliquant le premier principe de la thermodynamique (voir chapitre suivant), la relation : $T_1 = T_2 = T_0 - \frac{k}{6nC_V} \left(\frac{V_2 - V_0}{S} \right)^2$. La température finale est donc plus petite que T_0 .

22.4 Étude d'un compresseur

1. a. Le système est le gaz qui entre dans le corps de la pompe pendant la première phase. L'évolution étudiée comporte les trois étapes suivantes :
- admission du gaz de A à B à la pression P_0 ;
 - compression de B à C ;
 - refoulement de C à D à la pression P_1 .



- b. La pression extérieure est constante, donc le travail des forces de pression du gaz à droite du piston est :

$$W_{\text{droite}} = W_{\text{droite,AB}} + W_{\text{droite,BC}} + W_{\text{droite,CD}} = P_0(V_B - V_A + V_C - V_B + V_D - V_C) = 0$$

Effectivement, puisque le piston effectue un aller-retour, le volume balayé est le même à l'aller (travail résistant) et au retour (travail moteur).

- c. La force exercée par le moteur pour déplacer le piston est $F = (P - P_0)s$, où s est la section du piston, donc le travail fourni par le moteur est :

$$W_{\text{moteur}} = - \int_{ABCD} (P - P_0) dV = - \int_{ABCD} P dV - W_{\text{droite}} = - \int_{ABCD} P dV$$

W_{moteur} est donc l'aire du cycle. C'est ce qu'on appelle habituellement le travail de l'opérateur.

2. On étudie le gaz dans le corps de pompe entre B et C. La loi polytropique et l'équation du gaz parfait donnent :

$$\begin{cases} P_B V_B^k = P_C V_C^k \\ P_B V_m = nRT_B \quad \text{et} \quad P_C V_C = nRT_C \end{cases}$$

CHAPITRE 22 – ÉNERGIE ÉCHANGÉE PAR UN SYSTÈME AU COURS D'UNE TRANSFORMATION

En combinant ces expressions, on obtient :

$$\left(\frac{P_B}{P_C}\right) = \left(\frac{V_B}{V_C}\right)^k \Rightarrow \left(\frac{P_B}{P_C}\right)^{1-k} = \left(\frac{T_B}{T_C}\right)^{-k}$$

Comme on veut calculer k , on prend le logarithme de cette expression, ce qui donne :

$$k = \frac{\ln \frac{P_B}{P_C}}{-\ln \frac{T_B}{T_C} + \ln \frac{P_B}{P_C}} = \frac{\ln \frac{P_0}{P_1}}{-\ln \frac{T_0}{T_1} + \ln \frac{P_0}{P_1}} = 1,2$$

3. On a vu précédemment que $W_{\text{moteur}} = \int (-PdV)$. On décompose le calcul suivant les trois transformations élémentaires :

– Dans la transformation de A à B , $P = P_0$ le travail fourni par le moteur est :

$$W_{\text{moteur},AB} = -P_0 \int_0^{V_m} dV = -P_0 V_m = -nRT_0.$$

– Dans la transformation de B à C , il vaut :

$$W_{\text{moteur},BC} = \int_{V_m}^{V_1} (-PdV) = -P_0 V_m^k \int_{V_m}^{V_1} \frac{dV}{V^k} = \frac{P_0 V_m^k}{k-1} \left[V^{-k+1} \right]_{V_m}^{V_1},$$

ce qui peut s'écrire :

$$W_{\text{moteur},BC} = \frac{P_0 V_m}{k-1} \left(\left(\frac{V_m}{V_1} \right)^{k-1} - 1 \right).$$

Or :

$$\left(\frac{V_m}{V_1} \right)^{k-1} = \left(\frac{V_m}{V_1} \right)^k \frac{V_1}{V_m} = \frac{P_1 V_1}{P_0 V_m} = \frac{T_1}{T_0},$$

d'où finalement :

$$W_{\text{moteur},BC} = \frac{P_0 V_m}{k-1} \left(\frac{T_1}{T_0} - 1 \right).$$

– Dans la transformation de C à D , $P = P_1$ et le travail du moteur est :

$$W_{CD} = - \int_{V_1}^0 P_1 dV = P_1 V_1 = nRT_1.$$

Ainsi, finalement :

$$W_{\text{moteur}} = \frac{k}{k-1} nR(T_1 - T_0).$$

4. Par définition de la puissance : $W_{\text{moteur}} = P_{\text{moteur}} \Delta t$.

En $\Delta t = 1$ s, la quantité d'air passant dans le compresseur est : $n = \frac{D_m \Delta t}{M}$. Ainsi :

$$\mathcal{P}_{\text{moteur}} = \frac{k}{k-1} \frac{D_m}{M} R(T_1 - T_0) = 2,26 \text{ kW}.$$

Premier principe. Bilans d'énergie.

23

Le **premier principe de la thermodynamique** exprime la conservation de l'énergie. Ainsi, un système isolé, c'est-à-dire un système n'ayant aucun échange d'énergie avec l'extérieur, a une énergie constante. Les échanges d'énergie d'un système non isolé sont le travail mécanique et le transfert thermique, qui ont été étudiés au chapitre précédent. Faire le bilan d'énergie d'une transformation d'un système choisi consiste à calculer sa variation d'énergie, ainsi que les contributions respectives de ces deux types d'échange d'énergie à cette variation.

1 Le premier principe de la thermodynamique

1.1 Énergie d'un système

La notion d'énergie généralise au cas où il y a un mouvement de matière à l'échelle macroscopique la notion d'énergie interne, définie au chapitre 21.

L'**énergie** E d'un système thermodynamique Σ est la somme de :

- son énergie interne U ,
- son énergie cinétique macroscopique E_c dans le référentiel de l'étude,
- une éventuelle énergie potentielle d'interaction avec un système extérieur $E_{p,ext}$,

soit :

$$E = U + E_c + E_{p,ext}.$$

L'énergie cinétique macroscopique est l'énergie cinétique attribuée aux mouvements observables à l'échelle macroscopique. Pour la calculer on divise le système en volumes mésoscopiques $d\tau$ de masse dm . Chacun des volumes mésoscopiques est assimilable, du fait de sa très petite taille, à un point matériel animé d'une vitesse $\vec{v}$ et a donc une énergie cinétique égale à $\frac{1}{2}dm v^2$. L'énergie cinétique de Σ est la somme des énergies cinétiques de tous les volumes mésoscopiques qui le composent.

L'énergie potentielle d'interaction avec un système extérieur est dans la plupart des cas l'éner-

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

gie potentielle de pesanteur (interaction avec la Terre), qui est :

$$E_{p,\text{ext}} = mgz_G,$$

où m est la masse de Σ et z_G l'altitude du centre d'inertie du système, c'est-à-dire sa coordonnée sur une axe (Oz) vertical dirigé vers le haut. $E_{p,\text{ext}}$ est définie à une constante additive près correspondant au choix de l'origine O .

L'énergie E est définie à une constante additive près. Ceci n'est pas gênant parce que les expériences ne permettent de mesurer que des variations de E entre deux états.

Remarque

Par rapport à la définition qui a été donnée au chapitre 21, on ajoute parfois à l'énergie interne l'énergie de masse mc^2 où m est la masse de Σ et c la célérité de la lumière. Comme cette énergie ne varie pas dans les phénomènes courants, cela revient simplement à changer la constante additive.

1.2 Premier principe de la thermodynamique

Le **premier principe de la thermodynamique** s'applique à un système Σ donné pour une transformation thermodynamique donnée. On notera, pour la transformation considérée,

$$\Delta X = X_f - X_i$$

la variation d'une fonction d'état quelconque X du système, différence entre la valeur X_f dans l'état final et la valeur X_i dans l'état initial. Ainsi, ΔE est la variation d'énergie, ΔU la variation d'énergie interne, ΔE_c la variation d'énergie cinétique macroscopique, etc.

a) Énoncé

Au cours d'une transformation thermodynamique quelconque d'un système fermé Σ , la variation de l'énergie E de Σ est égale à l'énergie qu'il reçoit, somme du travail mécanique W et du transfert thermique Q , soit :

$$\Delta E = W + Q. \quad (23.1)$$

Le premier principe exprime la **conservation de l'énergie**.

Il est important de se rappeler que le travail W et le transfert thermique Q sont des **grandeurs algébriques**. Ces énergies sont, par convention, positives si elle sont effectivement reçues par le système Σ de l'extérieur et négatives si elles sont données par Σ à l'extérieur.

Les grandeurs ΔE , W et Q dépendent du système choisi et de la transformation considérée, qu'il faudra toujours préciser avec soin.

b) Cas d'un système isolé

Un **système isolé** est un système fermé n'échangeant pas d'énergie avec l'extérieur.

Un système isolé n'a pas d'interaction avec l'extérieur donc $E_{p,\text{ext}} = 0$. Son énergie s'écrit ainsi : $E = U + E_c$.

De plus, l'absence d'échange d'énergie se traduit par $Q = 0$ et $W = 0$ lors de toute transformation du système. Ainsi :

L'énergie $E = U + E_c$ d'un système isolé est constante.

Il peut y avoir à l'intérieur du système des transformation d'une forme d'énergie en une autre, mais l'énergie E du système, somme de toutes les formes d'énergie, est constante.

c) Forme usuelle du premier principe

La variation d'énergie du système peut s'écrire :

$$\begin{aligned}\Delta E &= (U_f + E_{c,f} + E_{p,\text{ext } f}) - (U_i + E_{c,i} + E_{p,\text{ext } i}) \\ &= (U_f - U_i) + (E_{c,f} - E_{c,i}) + (E_{p,\text{ext } f} - E_{p,\text{ext } i}) = \Delta U + \Delta E_c + \Delta E_{p,\text{ext}}.\end{aligned}$$

De plus, le programme limite les applications aux cas où l'énergie potentielle d'interaction avec l'extérieur est nulle (pas d'interaction) ou bien constante.

Dans le cadre du programme on suppose qu'il n'y a pas de variation d'énergie potentielle et le premier principe s'écrit :

$$\Delta U + \Delta E_c = W + Q. \tag{23.2}$$

d) Remarque importante

L'équation (23.2) est relative à un système et à une transformation de ce système qui se définit par : un état initial i , un état final f , et un chemin menant de l'état initial à l'état final.

Les termes $\Delta U = U_f - U_i$ et $\Delta E_c = E_{c,f} - E_{c,i}$ ne dépendent que de l'état initial i et de l'état final f et ne dépendent pas du chemin suivi.

En revanche, les termes Q et W dépendent du chemin suivi (ceci a été montré au chapitre précédent pour W).

Il y a donc une différence de statut entre les termes de l'équation, différence que la notation souligne par la présence ou absence d'un « Δ ».



C'est une erreur grave de mettre un « Δ » devant le W ou le Q .

1.3 Obtention de la valeur du transfert thermique

Dans le cadre de la thermodynamique des états d'équilibre étudiée dans le cours de première année, le premier principe est la seule loi permettant de calculer le transfert thermique Q reçu par un système au cours d'une transformation. Si on ne fait pas l'hypothèse d'une transformation adiabatique, on ne peut obtenir Q que par l'équation :

$$Q = \Delta U + \Delta E_c - W. \tag{23.3}$$

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

Remarque

La théorie des transferts thermiques qui sera vue en deuxième année donnera un moyen de calculer le transfert thermique Q directement.

1.4 Transfert thermique dans une transformation isochore sans travail autre que celui de la pression et sans variation d'énergie cinétique

Dans une transformation isochore le travail des forces de pression est nul (voir chapitre précédent). S'il n'y a pas d'autre force que les forces de pression, on a alors : $W = 0$. Si, de plus, l'énergie cinétique macroscopique du système dans le référentiel est soit nulle, soit constante $\Delta E_c = 0$. Le premier principe s'écrit alors :

$$Q = \Delta U. \quad (23.4)$$

En particulier, pour élever la température d'un système de $\Delta T = 1$ K par chauffage isochore, il faut fournir une énergie égale, en joules, à la valeur, en joules par kelvin, de sa capacité thermique à volume constant C_V , puisque dans ce cas :

$$Q = \Delta U = C_V \Delta T.$$

1.5 Exemples d'application du premier principe

a) Échauffement isochore d'un gaz (exemple du chapitre 22, page 796)

On prend pour système le gaz contenu dans le récipient. On est dans le cas du paragraphe 1.4 donc :

$$W = 0 \quad \text{et} \quad Q = \Delta U.$$

Par définition de la capacité thermique à volume constant C_V : $\Delta U = \int_{T_i}^{T_f} C_V dT$. Cette expression se simplifie si C_V est indépendante de T . Par exemple pour un gaz de capacité thermique à volume constant C_V indépendante de T , $C_V = nC_{Vm}$ où n est la quantité de matière dans le système et :

$$Q = \Delta U = nC_{Vm}(T_f - T_i) = nC_{Vm}(T_0 - T_i).$$

b) Échauffement isobare d'un gaz (exemple du chapitre 22, page 797)

On prend pour système l'ensemble {gaz contenu dans le récipient + piston}.

Le travail de la force $\vec{F}$ peut se calculer :

- à partir de la formule de la mécanique : le déplacement du piston est $\Delta \ell = \frac{V_f - V_i}{S}$, la force est dans le sens opposé au déplacement donc :

$$W = -F \Delta \ell = -\frac{F}{S}(V_f - V_i);$$

- en considérant que cette force correspond à une pression sur le piston égale à $P_{\text{ext}} = \frac{F}{S}$,

constante, et en appliquant la formule du travail de la force de pression dans le cas monobare :

$$W = -P_{\text{ext}}(V_f - V_i) = -\frac{F}{S}(V_f - V_i).$$

On retrouve, bien sûr, la même expression.

Il n'y a pas d'énergie cinétique macroscopique donc : $\Delta E = \Delta U$.

On en déduit, par application du premier principe :

$$Q = \Delta U - W = \Delta U + \frac{F}{S}(V_f - V_i).$$

On remarque que le transfert thermique nécessaire pour chauffer le gaz est plus grand que la variation d'énergie interne puisque $V_f > V_i$. Ceci est dû au fait que le gaz fournit du travail en poussant le piston.

On verra, dans la section 2, une manière directe de calculer le transfert thermique dans ce cas.

c) Échauffement d'un gaz par compression

Un cylindre fermé par un piston étanche contient de l'air à la température T_i . On déplace brutalement le piston sur une longueur ℓ en exerçant une force de norme F constante. D'autre part il s'exerce sur le piston qui se déplace une force de frottement de norme F_f constante. Quelle est la température finale T_f de l'air dans le cylindre ?

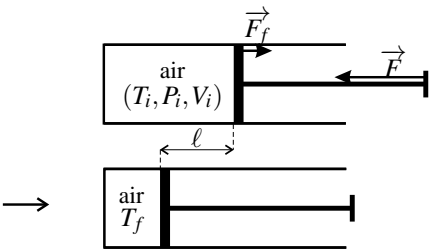


Figure 23.1 – Compression brutale d'un gaz

On va appliquer le premier principe au système constitué par l'air contenu dans le cylindre pour cette transformation. On fait l'hypothèse que l'air est un gaz parfait diatomique de capacité thermique à volume constant $C_{Vm} = \frac{5}{2}R$ indépendante de la température.

Il n'y a pas de variation d'énergie cinétique macroscopique donc $\Delta E_c = 0$. La variation d'énergie interne s'écrit :

$$\Delta U = nC_{Vm}(T_f - T_i) = \frac{5}{2}nR(T_f - T_i),$$

en notant n la quantité de matière dans le système.

La transformation étant brutale, on peut la considérer comme adiabatique, soit : $Q = 0$.

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

Le travail reçu par le gaz est le travail de la force $F - F_f$ exercée par le piston supposé sans masse, soit : $W = (F - F_f)\ell$, car les forces $\vec{F}$ et $\vec{F}_f$ sont de normes constantes.

Le premier principe s'écrit donc :

$$\Delta U = W \quad \text{soit} \quad \frac{5}{2}nR(T_f - T_i) = (F - F_f)\ell.$$

On en tire : $T_f = T_i + \frac{2}{5nR}(F - F_f)\ell$, expression que l'on peut réécrire, en utilisant l'équation d'état du gaz parfait $P_i V_i = nRT_i$, sous la forme :

$$T_f = T_i \left(1 + \frac{2}{5} \frac{(F - F_f)\ell}{P_i V_i} \right).$$

Cet exemple illustre le fait que la température peut augmenter sans qu'il y ait un transfert thermique, quand le système reçoit l'énergie correspondante sous forme de travail.

L'expérience est réalisable et spectaculaire : l'augmentation de température du gaz est suffisante pour provoquer l'ignition d'un petit morceau de papier placé à l'intérieur. Par exemple, si on utilise un tube de diamètre $d = 1$ cm dans lequel l'air occupe initialement une longueur

$L = 15$ cm, on a $V_i = \frac{\pi d^2}{4}L = 1,2 \cdot 10^{-5} \text{ m}^3$ et, avec $P_i = 1 \cdot 10^5$ Pa, on a : $P_i V_i = 1,2$ J. Si la force de frottement vaut $F_f = 5$ N et qu'on exerce la force $F = 45$ N en déplaçant le piston de $\ell = 10$ cm, alors $(F - F_f)\ell = (45 - 5) \times 0,10 = 4,0$ J. Dans ce cas la température passe de la valeur $T_i = 300$ K à $T_f = 300 \times \left(1 + \frac{2}{5} \frac{4,0}{1,2} \right) = 700$ K !

d) Transformation d'un système composé (voir figure 22.3 du chapitre 22, page 799)

On fait le bilan d'énergie du système Σ constitué par tout ce que l'enceinte contient.

Le système ayant une paroi indéformable, il ne reçoit pas de travail :

$$W = 0.$$

Il n'y a pas d'énergie cinétique macroscopique donc : $\Delta E = \Delta U$. Par additivité de l'énergie interne, la variation d'énergie interne de Σ est :

$$\Delta U = \Delta U_{\Sigma_1} + \Delta U_{\text{cloison}} + \Delta U_{\Sigma_2} = C_{V,\Sigma_1}(T_{f1} - T_i) + 0 + C_{V,\Sigma_2}(T_{f2} - T_i),$$

puisque la cloison sans masse a une capacité thermique à volume constant nulle. En supposant que les deux compartiments contiennent un même gaz parfait de capacité thermique à volume constant molaire C_{Vm} , il vient :

$$\Delta U = nC_{Vm}(T_0 - T_i) + 2nC_{Vm}(T_0 - T_i) = 3nC_{Vm}(T_0 - T_i).$$

On en déduit par application du premier principe le transfert thermique reçu par Σ :

$$Q = \Delta U = 3nC_{Vm}(T_0 - T_i).$$

Le bilan d'énergie du système Σ_1 ne peut être établi complètement. En effet, on sait calculer ΔU_{Σ_1} mais on ne sait pas calculer le travail W_{Σ_1} reçu par Σ_1 de la part de la cloison qui se déplace. Ce travail est positif puisque le volume de Σ_1 diminue. De même Σ_2 reçoit un travail W_{Σ_2} négatif. Si la cloison est sans masse, $W_{\Sigma_1} = -W_{\Sigma_2}$: les deux systèmes Σ_1 et Σ_2 échangent de l'énergie sous forme de travail. Ils en échangent aussi sous forme de transfert thermique si la cloison n'est pas adiabatique.

Ainsi, le choix du système Σ englobant Σ_1 et Σ_2 permet d'ignorer les échanges d'énergie entre ces deux systèmes que l'on ne sait pas modéliser.

e) Système mécanique avec frottements

On considère un solide A de masse m_A en translation à la vitesse $\vec{v}$ glissant avec frottement sur un solide fixe B (voir figure 23.2). On considère la transformation suivante : dans l'état initial la vitesse de A est $\vec{v}_i = \vec{v}_0$. Dans l'état final A est immobilisé $\vec{v}_f = \vec{0}$.

On va faire le bilan d'énergie du système $\{A + B\}$ sur cette transformation.

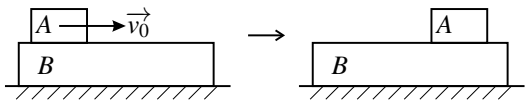


Figure 23.2 – Système mécanique avec frottements

La variation d'énergie cinétique du système est

$$\Delta E_c = E_{c,f} - E_{c,i} = \frac{1}{2} m_A v_f^2 - \frac{1}{2} m_A v_i^2 = 0 - \frac{1}{2} m_A v_0^2.$$

La variation d'énergie interne du système est, par additivité :

$$\Delta U = \Delta U_A + \Delta U_B.$$

Il est raisonnable de supposer la transformation adiabatique parce qu'elle est très rapide. De plus les corps A et B sont à la même température que l'air ambiant au début de l'expérience.

Par ailleurs les forces extérieures agissant sur le système $\{A + B\}$ sont les suivantes : le poids de A qui ne travaille pas puisque le mouvement de A est horizontal, le poids de B qui ne travaille pas puisque B est fixe, la force exercée par la table sur B qui ne travaille pas parce que B est fixe. Les forces de pression ne travaillent pas car la pression extérieure est uniforme et les volumes des solides ne varient pas.

Le premier principe pour le système $\{A + B\}$ dans cette transformation s'écrit donc :

$$\Delta U + \Delta E_c = 0 \quad \text{soit} \quad \Delta U_A + \Delta U_B = \frac{1}{2} m_A v_0^2.$$

Cette relation exprime le fait que l'énergie cinétique de A est transformée en énergie interne de A et de B . Localement, à l'endroit où il y a eu contact, la température des solides a augmenté.

Il est important de noter que le bilan d'énergie de $\{A + B\}$ ne comporte pas le travail de la force de frottement que B exerce sur A car ce travail représente un transfert d'énergie à l'intérieur du système.

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

2 La fonction d'état enthalpie

2.1 Définitions

a) Enthalpie d'un système thermodynamique

On appelle **enthalpie** d'un système thermodynamique la fonction d'état :

$$H = U + PV,$$

où U est l'énergie interne, P la pression et V le volume.

L'enthalpie se mesure en joules.

H est une fonction d'état **extensive**, puisque U et le produit PV sont des fonctions extensives (le produit de la variable intensive P par la variable extensive V est extensif). Elle est aussi **additive** :

$$H_{\Sigma_1 + \Sigma_2} = H_{\Sigma_1} + H_{\Sigma_2},$$

si on réunit deux systèmes Σ_1 et Σ_2 .

b) Capacité thermique à pression constante

On appelle **capacité thermique à pression constante** d'un système *fermé* Σ la grandeur C_P telle que la variation dH de l'enthalpie de Σ lorsque la température varie de dT , la *pression restant constante*, est :

$$dH = C_P dT.$$

C_P se mesure en joules par kelvin : $\text{J} \cdot \text{K}^{-1}$.

C_P dépend en général de la température T . La variation d'enthalpie du système dans une transformation isobare où la température passe de la valeur T_i à la valeur T_f est donc donnée par :

$$\Delta H = \int_{T_i}^{T_f} C_P(T) dT. \quad (23.5)$$

La capacité thermique à pression constante est une grandeur **extensive** et **additive** qui s'exprime en $\text{J} \cdot \text{K}^{-1}$.

Pour un échantillon de corps pur monophasé on définit la **capacité thermique massique à pression constante** c_P en $\text{J} \cdot \text{kg}^{-1} \cdot \text{K}^{-1}$ et la **capacité thermique molaire à pression constante** C_{Pm} en $\text{J} \cdot \text{mol}^{-1} \cdot \text{K}^{-1}$. On a alors :

$$C_P = mc_P = nC_{Pm},$$

en notant m la masse et n la quantité de matière du système. De plus, comme $m = nM$, où M est la masse molaire du corps pur :

$$C_{Pm} = Mc_P.$$

2.2 Premier principe pour une transformation monobare avec équilibre mécanique dans l'état initial et l'état final

On considère donc un système subissant une évolution monobare, sous la pression extérieure constante $P_{\text{ext}} = P_0$ constante, entre un état initial i caractérisé par les variables d'état (T_i, P_i, V_i) et un état final f caractérisé par les variables (T_f, P_f, V_f) . On suppose de plus qu'il y a équilibre mécanique entre le système et l'extérieur dans l'état final et l'état initial, soit :

$$P_i = P_f = P_0.$$

Le travail des forces de pression dans la transformation est donné par la formule (22.4), page 804 :

$$W_{\text{pression}} = -P_0(V_f - V_i) = -P_f V_f + P_i V_i = -\Delta(PV).$$

On appelle W_{autre} le travail des forces autres que les forces de pression : $W = W_{\text{pression}} + W_{\text{autre}}$. Le premier principe s'écrit :

$$\Delta U + \Delta E_c = W + Q \quad \text{soit} \quad \Delta U + \Delta E_c = -\Delta(PV) + W_{\text{autre}} + Q.$$

Or : $\Delta U + \Delta(PV) = \Delta(U + PV) = \Delta H$. Finalement :

Pour un système subissant une **transformation monobare**, avec **équilibre mécanique dans l'état initial et l'état final**, le premier principe peut s'écrire sous la forme :

$$\Delta H + \Delta E_c = W_{\text{autre}} + Q \tag{23.6}$$

où W_{autre} est le travail des forces autres que les forces de pression.

L'intérêt de cette formule est qu'elle évite le calcul du travail des forces de pression pour évaluer le transfert thermique.

Dans le cas où il n'y a pas d'autre travail que celui des forces de pression ($W_{\text{autre}} = 0$) et pas d'énergie cinétique ($\Delta E_c = 0$) elle s'écrit :

$$Q = \Delta H.$$

2.3 Transfert thermique dans une transformation isobare sans travail autre que celui de la pression et sans variation d'énergie cinétique

On considère une transformation isobare d'un système ($P_i = P = P_f$) dans laquelle il n'y a pas de variation d'énergie cinétique, ni de travail d'autres forces que les forces de pression.

Le travail des forces de pression dans la transformation est donné par la formule (22.5), page 804 : $W = -P(V_f - V_i) = P_i V_i - P_f V_f$.

Le premier principe permet d'exprimer le transfert thermique reçu par le système au cours de la transformation : $Q = \Delta U - W = U_f - U_i + P_f V_f - P_i V_i = H_f - H_i$, soit :

$$Q = \Delta H. \tag{23.7}$$

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

En particulier, pour élever la température d'un système de $\Delta T = 1 \text{ K}$ par chauffage isobare, il faut fournir une énergie égale, en joules, à la valeur, en joules par kelvin, de sa capacité thermique à pression constante C_P , puisque dans ce cas :

$$Q = \Delta H = C_P \Delta T.$$

Il faut rapprocher la formule (23.7) de la formule (23.4) et ne pas les confondre. L'encadré ci-dessous récapitule les expressions du transfert thermique à connaître.

Pour un système ne recevant de travail que des forces de pression et n'ayant pas d'énergie cinétique :

- $Q = \Delta U$ si la transformation est **isochore**,
- $Q = \Delta H$ si la transformation est **isobare** ou si elle est **monobare avec équilibre mécanique dans l'état initial et l'état final**.



Le fait que la valeur de Q coïncide dans certains cas avec la variation d'une fonction d'état ne doit pas faire oublier que Q n'est pas la variation d'une fonction d'état.

2.4 Enthalpie d'un gaz parfait

a) Gaz parfait quelconque

Dans le cas d'un gaz parfait, l'enthalpie molaire s'écrit :

$$H_m = U_m + PV_m = U_m + RT, \tag{23.8}$$

en notant U_m et V_m l'énergie interne et le volume molaires et en utilisant l'équation d'état du gaz parfait. D'après la première loi de Joule, U_m ne dépend que de T . Il en est clairement de même pour H_m . C'est la **deuxième loi de Joule** :

L'enthalpie molaire d'un gaz parfait ne dépend que de sa température, ce que l'on peut écrire :

$$H_m = H_m(T).$$

La capacité thermique à pression constante molaire du gaz parfait est :

$$C_{Pm} = \frac{dH_m}{dT}.$$

Dans un domaine de température où C_{Pm} est une constante on a :

$$H_m(T) = C_{Pm}T + \text{constante},$$

et, pour une transformation entre un état i et un état f :

$$\Delta H_m = C_{Pm}(T_f - T_i). \tag{23.9}$$

En dérivant l'équation (23.8) par rapport à T on obtient : $\frac{dH_m}{dT} = \frac{dU_m}{dT} + R$, soit :

$$C_{Pm} - C_{Vm} = R, \quad (23.10)$$

qui est la **relation de Mayer**. Ainsi, la capacité thermique à pression constante du gaz parfait est supérieure à sa capacité thermique à volume constant.

On donne souvent, pour un gaz parfait, le **rapport des capacités thermiques** :

$$\gamma = \frac{C_{Pm}}{C_{Vm}}. \quad (23.11)$$

Ce rapport est toujours plus grand que 1, d'après la relation de Mayer. On déduit facilement des relations (23.10) et (23.11) les expressions suivantes des deux capacités thermiques molaires :

$$C_{Vm} = \frac{R}{\gamma - 1} \quad \text{et} \quad C_{Pm} = \frac{R\gamma}{\gamma - 1}. \quad (23.12)$$



De ces relations découlent les relations suivantes concernant un échantillon de gaz de quantité de matière n :

$$C_P - C_V = nR, \quad C_V = \frac{nR}{\gamma - 1} \quad \text{et} \quad C_P = \frac{nR\gamma}{\gamma - 1},$$

et les relations suivantes concernant l'unité de masse de gaz :

$$c_P - c_V = \frac{R}{M}, \quad c_V = \frac{R}{M(\gamma - 1)} \quad \text{et} \quad c_P = \frac{R\gamma}{M(\gamma - 1)}.$$

b) Gaz parfait monoatomique

Dans le cas d'un gaz parfait monoatomique on a vu au chapitre 21 que :

$$C_{Vm} = \frac{3}{2}R.$$

On en déduit que :

$$C_{Pm} = \frac{5}{2}R \quad \text{et} \quad \gamma = \frac{5}{3}.$$

c) Gaz parfait diatomique

Dans le cas d'un gaz parfait diatomique, aux températures usuelles, on a vu au chapitre 21 que :

$$C_{Vm} = \frac{5}{2}R.$$

On en déduit que :

$$C_{Pm} = \frac{7}{2}R \quad \text{et} \quad \gamma = \frac{7}{5}.$$

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

2.5 Enthalpie d'une phase condensée indilatable et incompressible

Pour les solides et liquides incompressibles, on a vu au chapitre 21 que l'énergie interne molaire est fonction uniquement de la température $U_m(T)$. On ne peut pas dire la même chose de l'enthalpie molaire $H_m = U_m(T) + PV_m$, parce qu'elle dépend explicitement de la pression.

Cependant, la variation d'enthalpie molaire correspondant à une variation de pression ΔP , égale à $V_m \Delta P$, est souvent négligeable parce que le volume molaire d'une phase condensée est très inférieur à celui d'un gaz. Dans le cadre du programme, on considérera que H_m ne dépend que de T .

D'autre part, si la température varie de dT , la pression restant constante, le terme PV_m ne varie pas, puisque V_m est une constante V_{m0} , donc la variation de l'enthalpie molaire est $dH_m = dU_m = C_{Vm}dT$. Or par définition, $dH_m = C_{Pm}dT$, on a donc $C_{Pm} = C_{Vm}$.

Pour une phase condensée incompressible et indilatable, on peut faire l'approximation que l'enthalpie molaire est indépendante de la pression, soit :

$$H_m \simeq H_m(T).$$

De plus les capacités thermiques molaires à pression constante et à volume constant sont pratiquement égales :

$$C_{Pm} \simeq C_{Vm}.$$

La valeur commune aux deux capacités thermiques molaires est noté le plus souvent C_m .

Sur un domaine de température où C_m est une constante, on peut écrire :

$$U_m = C_m T + \text{constante} \quad \text{et} \quad H_m = C_m T + \text{constante}', \tag{23.13}$$

avec deux constantes différentes. Et pour une transformation au cours de laquelle la température varie de ΔT :

$$\Delta U_m = \Delta H_m = C_m \Delta T. \tag{23.14}$$

Pour les grandeurs massiques, on a de même :

$$u = cT + \text{constante}, \quad h = cT + \text{constante}' \quad \text{et} \quad \Delta u = \Delta h = c\Delta T,$$

où c est la capacité thermique massique.

Remarque

Le modèle de la phase condensée incompressible et indilatable est un modèle théorique limite. Pour les liquides et solides réels : $c_P \simeq c_V$.

Il faut connaître l'ordre de grandeur de la capacité thermique massique de l'eau liquide dont la valeur à température ambiante est : $c_{\text{eau}} = 4,18 \cdot 10^3 \text{ J} \cdot \text{K}^{-1} \cdot \text{kg}^{-1}$. C'est une valeur particulièrement élevée.

2.6 Enthalpie d'un système diphasé

a) Expression de l'enthalpie

On considère un système constitué par un corps pur dans deux phases différentes notées I et II (I et II représentant deux des lettres S pour solide, L pour liquide ou G pour gaz) et dont l'état est déterminé par les variables d'état suivantes : la quantité de matière totale n ou la masse m , la température T et le titre massique x_{II} de la phase II (voir chapitre 21, page 775).

Pour calculer l'enthalpie du système on utilise l'additivité de l'enthalpie :

$$H = n_I H_{m,I} + n_{II} H_{m,II} = n((1 - x_{II})H_{m,I} + x_{II}H_{m,II}),$$

ou encore :

$$H = m_I h_I + m_{II} h_{II} = m((1 - x_{II})h_I + x_{II}h_{II}).$$

Soit, en regroupant les termes en x_{II} :

$$H = n(H_{m,I} + x_{II}(H_{m,II} - H_{m,I})) = m(h_I + x_{II}(h_{II} - h_I)). \quad (23.15)$$

b) Enthalpies de changement d'état

On appelle **enthalpie molaire de changement d'état** $\Delta_{I-II}H_m$ la variation d'enthalpie au cours de la transformation d'une mole de corps pur de l'état I à l'état II en un point du plan (P, T) où les phases I et II coexistent, soit :

$$\Delta_{I-II}H_m = H_{m,II} - H_{m,I}.$$

$\Delta_{I-II}H_m$ se mesure en $\text{J}\cdot\text{mol}^{-1}$.

On appelle **enthalpie massique de changement d'état** $\Delta_{I-II}h$ la variation d'enthalpie au cours de la transformation d'un kilogramme de corps pur de l'état I à l'état II en un point du plan (P, T) où les phases I et II coexistent, soit :

$$\Delta_{I-II}h = h_{II} - h_I.$$

$\Delta_{I-II}h$ se mesure en $\text{J}\cdot\text{kg}^{-1}$.

Les enthalpies molaire et massique de changement d'état **ne dépendent que de la température** T puisque la pression est imposée par la condition d'équilibre de diffusion $P = P_{I-II}(T)$ nécessaire à la coexistence à l'équilibre des phases I et II .

Dans la pratique, l'habitude est de prendre toujours la phase II moins ordonnée que la phase I . De plus, l'indice « $I - II$ » est remplacé par les trois premières lettres du nom du changement d'état (voir figure 21.9, page 772) :

$$S - L = \text{fus}, \quad L - V = \text{vap} \quad \text{et} \quad S - G = \text{sub},$$

Par exemple l'enthalpie molaire de fusion est notée $\Delta_{\text{fus}}H_m$ et l'enthalpie massique de vaporisation est notée $\Delta_{\text{vap}}h$. L'enthalpie massique de liquéfaction n'a pas de notation propre, mais elle est égale à $-\Delta_{\text{vap}}h$.

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

L'enthalpie d'un échantillon de corps pur diphasé, comportant les phases *I* et *II*, peut se mettre sous les formes suivantes :

$$H = n(H_{m,I} + x_{II}\Delta_{I-II}H_m) = m(h_I + x_{II}\Delta_{I-II}h). \quad (23.16)$$

2.7 Variations d'enthalpie isobares

Le calcul de la variation d'enthalpie entre deux états de même pression est fondamental pour l'application du premier principe aux transformations isobares ou aux transformations monobares avec équilibre mécanique dans l'état final et dans l'état initial.

Pour le calcul de la variation d'une fonction d'état, on peut choisir le chemin menant de l'état initial à l'état final pour lequel le calcul est le plus simple, même si ce n'est pas le chemin réel suivi par le système.

a) Variation d'enthalpie due à un changement de température

On considère un échantillon de corps pur monophasé passant d'un état initial *i* caractérisé par les variables d'état (T_i, P_0) à l'état final *f* caractérisé par (T_f, P_0) .

Pour le calcul de la variation d'enthalpie, on peut considérer un chemin *isobare*. D'après la définition de la capacité thermique à pression constante C_P du système, la variation d'enthalpie est :

$$\Delta H = \int_{T_i}^{T_f} C_P(T) dT.$$

Dans le cadre des applications on suppose très souvent que la capacité thermique à pression constante est toujours indépendante de la température. La formule précédente se simplifie alors en :

$$\Delta H = C_P(T_f - T_i). \quad (23.17)$$

b) Variation d'enthalpie due à un changement d'état isotherme et isobare

On considère une transformation d'un échantillon de corps pur de masse totale *m* passant d'un état d'équilibre initial où il se trouve dans les phases *I* et *II* décrit par les variables d'état $(T_0, P_0, x_{II,i})$ à un état d'équilibre final du même type décrit par les variables $(T_0, P_0, x_{II,f})$.

Remarque

On a donc : $P_0 = P_{I-II}(T_0)$.

La variation d'enthalpie est :

$$\Delta H = m(h_I + x_{II,f}\Delta_{I-II}h) - m(h_I + x_{II,i}\Delta_{I-II}h),$$

soit :

$$\Delta H = m(x_{II,f} - x_{II,i})\Delta_{I-II}h. \quad (23.18)$$

En utilisant les enthalpies molaires et la quantité de matière on trouve de la même manière :

$$\Delta H = n(x_{II,f} - x_{II,i})\Delta_{I-II}H_m. \tag{23.19}$$

Quel est le bilan d'énergie de la transformation ? La transformation étant isobare, s'il n'y a pas d'autre travail que celui des forces de pression et s'il n'y a pas d'énergie cinétique, le transfert thermique reçu par le système est :

$$Q = \Delta H.$$

La transformation étant isobare, le travail des forces de pression reçu par le système est :

$$\begin{aligned} W &= -P\Delta V = -PV_f + PV_i \\ &= -Pm((1 - x_{II,f})v_I + x_{II,f}v_{II}) + Pm((1 - x_{II,i})v_I + x_{II,i}v_{II}) \\ &= -Pm(v_I + x_{II,f}(v_{II} - v_I)) + Pm(v_I + x_{II,i}(v_{II} - v_I)) = -Pm(x_{II,f} - x_{II,i})(v_{II} - v_I). \end{aligned}$$

Exemple

1. Une poche à glace contient une masse totale $m = 500$ g d'un mélange d'eau liquide et de glace avec une fraction $x_L = 0,200$ de liquide. Quel est la valeur maximale du transfert thermique qu'elle peut recevoir tout en restant à la température $T_0 = 273$ K, température d'équilibre eau-glace sous la pression ambiante $P_0 = 1,00 \cdot 10^5$ bar ?

Pour répondre il faut connaître l'enthalpie massique de fusion de la glace à la température T_0 . On trouve dans les tables : $\Delta_{\text{fus}}h = 335$ kJ·kg⁻¹.

La poche de glace reste à la température T_0 tant qu'il y a équilibre entre l'eau et la glace. Le transfert thermique est égal à la variation d'enthalpie puisque la transformation est isobare. Sa valeur maximale correspond à la fusion de toute la glace soit $x_{L,f} = 1$, donc :

$$Q_{\text{max}} = \Delta H_{\text{max}} = m(1 - x_L)\Delta_{\text{fus}}h = 0,500 \times (1 - 0,200) \times 335 \cdot 10^3 = 134 \cdot 10^3 \text{ J}.$$

2. On vaporise de manière isotherme et isobare, sous la pression $P_0 = 2,00 \cdot 10^5$ Pa, une masse $m = 10,0$ kg d'eau liquide. Quel transfert thermique doit-on apporter ? Quel est le travail des forces de pression ?

On trouve dans les tables la température d'ébullition de l'eau à la pression P_0 qui est $T_0 = 393$ K. À cette température :

- l'enthalpie massique de vaporisation de l'eau est $\Delta_{\text{vap}}h = 526$ kJ·kg⁻¹,
- le volume massique de la vapeur d'eau est $v_G = 885 \cdot 10^{-3}$ m³·kg⁻¹,
- le volume massique de l'eau liquide est $v_L = 1,06 \cdot 10^{-3}$ m³·kg⁻¹.

On prend pour système la masse d'eau et on considère la transformation isobare et isotherme menant de l'état initial $(T_0, P_0, x_{G,i} = 0)$ à l'état final $(T_0, P_0, x_{G,f} = 1)$.

La transformation est isobare, il n'y a pas de travail autre que celui des forces de pression et pas d'énergie cinétique, donc le transfert thermique reçu par le

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

système est :

$$Q = \Delta H = m\Delta_{\text{vap}}h = 10,0 \times 526.10^3 = 5,26.10^6 \text{ J.}$$

Le travail algébrique reçu par le système dans cette transformation isobare est :

$$\begin{aligned} W &= -P_0(V_f - V_i) = -P_0(mv_G - mv_L) \\ &= -2,0.10^5 \times 10 \times (885 - 1,06).10^{-3} = -1,8.10^6 \text{ J.} \end{aligned}$$

Ce travail est négatif : il est fourni par le système à l'extérieur. En fournissant du transfert thermique à ce système on obtient qu'il produise du travail mécanique : c'est le principe de base de la machine à vapeur.

c) Généralisation

On considère une transformation d'un échantillon de corps pur de masse totale m passant d'un état initial i diphasé (phases I et II) caractérisé par les variables d'état $(T_i, P_0, x_{II,i})$ à un état final f où il se trouve entièrement dans la phase II et caractérisé par les variables d'état (T_f, P_0) .

Pour calculer la variation d'enthalpie, on ne change donc pas le résultat en supposant que le système passe par un état intermédiaire dans lequel le corps pur est entièrement dans la phase II et caractérisé par les variables d'état (T_i, P_0) .

La première étape de la transformation est alors un changement d'état isotherme et isobare, x_{II} passant de $x_{II,i}$ à 1, pour lequel la variation d'enthalpie est donnée par la formule 23.18 du paragraphe précédent :

$$H_{\text{int}} - H_i = m(1 - x_{II,i})\Delta_{I-II}h.$$

La deuxième étape est une transformation isobare dans laquelle la température passe de T_i à T_f . Par définition de la capacité thermique à pression constante, la variation d'enthalpie au cours de cette étape est :

$$H_f - H_{\text{int}} = \int_{T_i}^{T_f} mc_{P,II}dT = mc_{P,II}(T_f - T_i),$$

si l'on suppose la capacité thermique massique $c_{P,II}$ de la phase II indépendante de la température. Finalement la variation d'enthalpie au cours de la transformation menant de l'état i à l'état f est :

$$\Delta H = H_f - H_i = (H_f - H_{\text{int}}) + (H_{\text{int}} - H_i) = m((1 - x_{II,i})\Delta_{I-II}h + c_{P,II}(T_f - T_i)). \quad (23.20)$$

3 Mesures de grandeurs thermodynamiques

Dans ce paragraphe on s'intéresse à des expériences permettant la mesure des grandeurs suivantes :

- capacité thermique à pression constante,
- enthalpie de changement d'état.

Ces expériences, réalisables en travaux pratiques, nécessitent un dispositif appelé calorimètre.

3.1 Le calorimètre

Un **calorimètre** (voir figure 23.3) est un récipient composé en général d'une paroi extérieure et d'une cuve intérieure, fermé par un couvercle percé de petites ouvertures permettant d'introduire un agitateur, un thermomètre, une résistance chauffante. La cuve intérieure étant séparée de la paroi extérieure par de l'air, le système est relativement bien isolé et on peut négliger sur la durée d'une expérience de travaux pratiques les échanges thermiques avec l'extérieur.

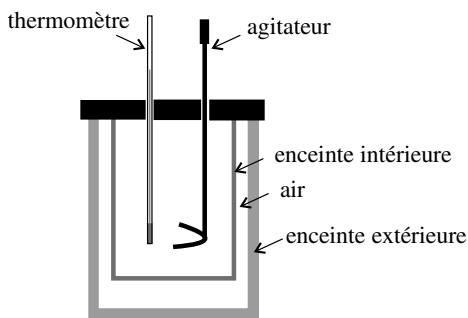


Figure 23.3 – Calorimètre.

Le système thermodynamique Σ étudié sera le système {calorimètre + son contenu}. Pour appliquer le premier principe on devra tenir compte de la capacité thermique du calorimètre parce qu'elle n'est pas négligeable en général devant la capacité thermique de ce qu'il contient. Par habitude, au lieu de donner la capacité thermique du calorimètre, on donne la masse d'eau qui aurait la même capacité thermique que l'on appelle **valeur en eau du calorimètre**. Ainsi, un calorimètre ayant une valeur en eau $\mu = 20 \text{ g}$ a une capacité thermique :

$$C = 4,18 \times 10^3 \times 20 \cdot 10^{-3} = 84 \text{ J} \cdot \text{K}^{-1}.$$

La valeur en eau du calorimètre tient compte de la capacité thermique de tous les instruments du calorimètre.

Quelle fonction d'état faut-il utiliser pour appliquer le premier principe au système Σ ? Les expériences dans un calorimètre se font à pression extérieure constante, le système étant en contact avec l'atmosphère par les petites ouvertures laissant passer le thermomètre et l'agitateur. Les transformations sont donc monobares. On utilisera donc l'enthalpie H , plutôt que l'énergie interne U , et le premier principe sous la forme (23.6).

Dans ce qui suit les températures seront notées θ et exprimées en degré Celsius.

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

3.2 Détermination d'une capacité thermique massique

On souhaite déterminer la capacité thermique massique c_{Fe} d'un échantillon fer. Le principe de l'expérience consiste à mesurer la température obtenue en mettant en contact thermique l'échantillon et une quantité connue d'eau, dont la capacité thermique à pression constante c_{eau} est connue. On réalise l'expérience dans le calorimètre pour avoir un système isolé thermiquement.

Expérience

On verse dans le calorimètre $m_{\text{eau}} = 400 \text{ g}$ d'eau très froide et on mesure la température qui se stabilise après quelques instants. On trouve $\theta_0 = 2,0^\circ\text{C}$. On introduit dans le calorimètre l'échantillon de fer, que l'on a préalablement pesé (sa masse est $m_{\text{Fe}} = 200 \text{ g}$) et qui est initialement à la température d'une étuve thermostatée, $\theta_1 = 85,0^\circ\text{C}$. On vérifie que l'échantillon est bien entièrement couvert d'eau. On attend que la température se stabilise et on mesure la température finale $\theta_F = 6,4^\circ\text{C}$.

On va appliquer le premier principe de la thermodynamique au système

$$\Sigma = \{ \text{calorimètre et instruments} + \text{eau} + \text{fer} \}$$

sur la transformation suivante :

État initial	$\text{Eau} + \text{calorimètre} : \theta_0$		$\text{Eau} + \text{calorimètre} : \theta_F$	
	$\text{Fer} : \theta_1$		$\text{Fer} : \theta_F$	

L'évolution est *monobare avec équilibre mécanique dans l'état initial et l'état final*, il n'y a pas d'énergie cinétique, ni de travail autre que celui de la pression. On a donc $\Delta H = Q$. Si le calorimètre est bien isolé, le système n'a pas d'échange thermique avec l'extérieur durant l'expérience donc :

$$\Delta H = 0.$$

L'enthalpie H est une fonction additive donc :

$$\Delta H = \Delta H_{\text{eau}} + \Delta H_{\text{calorimètre}} + \Delta H_{\text{Fe}}.$$

En supposant les différentes capacités thermiques à pression constantes indépendantes de la température, on a, d'après l'expression (23.17) page 832 :

$$\Delta H_{\text{eau}} = m_{\text{eau}} c_{\text{eau}} (\theta_F - \theta_0), \quad \Delta H_{\text{calorimètre}} = \mu c_{\text{eau}} (\theta_F - \theta_0), \quad \Delta H_{\text{Fe}} = m_{\text{Fe}} c_{\text{Fe}} (\theta_F - \theta_1),$$

en notant μ la valeur en eau du calorimètre. L'équation $\Delta H = 0$ conduit alors à :

$$\begin{aligned} c_{\text{Fe}} &= \frac{(m_{\text{eau}} + \mu) c_{\text{eau}} (\theta_F - \theta_0)}{m_{\text{Fe}} (\theta_1 - \theta_F)} \\ &= \frac{(400 + 20) \cdot 10^{-3} \times 4,18 \cdot 10^3 \times (6,4 - 2,0)}{200 \cdot 10^{-3} \times (85,0 - 6,4)} = 4,9 \cdot 10^2 \text{ J} \cdot \text{kg}^{-1}. \end{aligned}$$

La valeur que l'on trouve dans les tables est $c_{\text{Fe}} = 452 \text{ J} \cdot \text{kg}^{-1}$.

Remarque

- 1. Les pertes du calorimètre engendrent une erreur systématique sur le résultat. En effet, la température du calorimètre est en dessous de la température de la pièce donc les fuites thermiques sont vers l'intérieur du calorimètre et tendent à faire augmenter θ_F , ce qui fait grandir la valeur de c_{Fe} déterminée. Il est donc normal que l'on trouve une valeur trop élevée.
- 2. Pour améliorer la précision de la mesure, on a intérêt à ce que le changement de température de l'eau soit le plus grand possible.

3.3 Détermination d'une enthalpie de changement d'état

On souhaite déterminer l'enthalpie massique de fusion $\Delta_{fus,eau}h$ de l'eau. On dispose de glaçons sortant d'un congélateur à -18°C , d'eau, d'une balance, d'une bouilloire et d'un calorimètre.

L'idée est de mesurer la température finale d'un système dans lequel on a mis une quantité connue d'eau chaude (de température connue) et un glaçon (de masse et température connues). Il est souhaitable de choisir la quantité d'eau et sa température de telle manière, dans l'état final, que le système soit composé uniquement d'eau liquide, pour avoir une température qui s'homogénéise rapidement dans le système et dont la mesure détermine complètement l'état final du système.

Expérience

On verse dans le calorimètre $m_{eau} = 50\text{ g}$ d'eau chaude et on mesure la température qui se stabilise après quelques instants à la valeur $\theta_0 = 74^\circ\text{C}$.

On introduit ensuite dans le calorimètre un glaçon sortant du congélateur, à la température $\theta_1 = -18^\circ$, après en avoir déterminé la masse $m_{glace} = 19\text{ g}$. On attend que la température se stabilise et on mesure la température finale : $\theta_F = 38^\circ\text{C}$.

On va appliquer le premier principe de la thermodynamique au système

$$\Sigma = \{ \text{calorimètre et instruments} + \text{eau} + \text{glace} \}$$

sur la transformation suivante :

État initial		Eau liquide : m, θ_0		État final		Eau liquide : $m_{eau} + m_{glace}, \theta_F$
		Calorimètre : μ, θ_0				Calorimètre : θ_F
		Glace : m_{glace}, θ_1				

L'évolution est *monobare avec équilibre mécanique dans l'état initial et l'état final*, il n'y a pas d'énergie cinétique, ni de travail autre que celui de la pression. On a donc $\Delta H = Q$. Si le calorimètre est bien isolé, le système n'a pas d'échange thermique avec l'extérieur durant l'expérience donc :

$$\Delta H = 0.$$

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

L'enthalpie H est une fonction additive donc :

$$\Delta H = \Delta H_{\text{eau}} + \Delta H_{\text{calorimètre}} + \Delta H_{\text{glace}}.$$

En supposant les différentes capacités thermiques à pression constante indépendantes de la température, on a, d'après l'expression (23.17) :

$$\Delta H_{\text{eau}} = m_{\text{eau}} c_{\text{eau}} (\theta_F - \theta_0), \quad \Delta H_{\text{calorimètre}} = \mu c_{\text{eau}} (\theta_F - \theta_0),$$

en notant μ la valeur en eau du calorimètre.

Le calcul de ΔH_{glace} fait apparaître la capacité thermique de la glace : $c_S = 2,06 \text{ kJ} \cdot \text{kg}^{-1}$ et l'enthalpie massique de fusion de la glace $\Delta_{\text{fus}} h$ que l'on souhaite mesurer. On utilise la méthode donnée page 832. On imagine un chemin entre l'état initial (glace à θ_1) et l'état final (eau à θ_F) en trois étapes :

- transformation isobare dans laquelle la température passe de θ_1 à 0°C , l'enthalpie varie de : $m_{\text{glace}} c_S (0 - \theta_1)$,
- fusion isobare et isotherme de la glace à 0°C , l'enthalpie varie de : $m_{\text{glace}} \Delta_{\text{fus}} h$,
- transformation isobare dans laquelle l'eau liquide passe de 0°C à θ_F , l'enthalpie varie de : $m_{\text{glace}} c_{\text{eau}} (\theta_F - 0)$.

Ainsi :

$$\Delta H_{\text{glace}} = m_{\text{glace}} (c_{\text{eau}} \theta_F + \Delta_{\text{fus}} h - c_S \theta_1).$$

Finalement : $\Delta H = (m_{\text{eau}} + \mu) c_{\text{eau}} (\theta_F - \theta_0) + m_{\text{glace}} (c_{\text{eau}} \theta_F + \Delta_{\text{fus}} h - c_S \theta_1)$. Et de la relation $\Delta H = 0$ on déduit :

$$\Delta_{\text{fus}} h = \frac{(m_{\text{eau}} + \mu)}{m_{\text{glace}}} c_{\text{eau}} (\theta_0 - \theta_F) + c_S \theta_1 - c_{\text{eau}} \theta_F = 3,6 \cdot 10^5 \text{ J} \cdot \text{kg}^{-1}.$$

La valeur que l'on trouve dans les tables est $\Delta_{\text{fus}} h = 333 \text{ kJ} \cdot \text{kg}^{-1}$.

Remarque

Étant donné que la température finale θ_F est au-dessus de la température ambiante, les fuites thermiques font qu'elle est plus basse que ce que l'on calcule. L'expression de $\Delta_{\text{fus}} h$ ci-dessus montre qu'elle est alors surestimée. C'est une erreur systématique.

3.4 Mesure de la valeur en eau du calorimètre

Pour pouvoir réaliser les expériences précédentes il faut connaître la valeur en eau μ du calorimètre. Pour la déterminer on dispose d'une résistance électrique chauffante $R = 10 \Omega$, d'une alimentation continue 6 – 12 V et d'eau.

L'idée est de mesurer l'élévation de température du calorimètre lorsqu'on apporte une quantité d'énergie connue par la résistance.

Expérience

On remplit le calorimètre avec une masse $m_{\text{eau}} = 50 \text{ g}$ d'eau. On place la résistance dans le calorimètre sans la connecter. On laisse l'équilibre s'établir et on mesure la température : $\theta_0 = 21^\circ\text{C}$.

On branche la résistance sur l'alimentation réglée sur $U = 12 \text{ V}$ et pendant une durée $\tau = 2 \text{ min}$.

On attend que l'équilibre soit établi et on lit la température finale : $\theta_F = 27^\circ\text{C}$.

On va appliquer le premier principe de la thermodynamique au système

$$\Sigma = \{ \text{calorimètre et instruments} + \text{eau} + \text{résistance} \}$$

sur la transformation suivante :

$$\text{État initial} \mid \text{Eau + calorimètre : } \theta_0 \quad || \quad \text{État final} \mid \text{Eau + calorimètre : } \theta_F$$

Le système Σ subit une transformation *monobare avec équilibre mécanique dans l'état initial et l'état final*. On peut appliquer le premier principe sous la forme de l'équation (23.6) page 827.

Le calorimètre étant bien isolé, Σ ne reçoit aucun transfert thermique sur la durée de l'expérience : $Q = 0$. Il n'y a pas d'énergie cinétique. Par ailleurs Σ reçoit un travail autre que celui des forces de pression, le travail électrique fourni par le générateur : $W_{\text{autre}} = \mathcal{P}_{\text{Joule}} \tau = \frac{U^2}{R} \tau$. Le premier principe s'écrit donc :

$$\Delta H = \frac{U^2}{R} \tau.$$

La variation d'enthalpie se calcule comme au paragraphe 3.2 :

$$\Delta H_{\text{eau}} = m_{\text{eau}} c_{\text{eau}} (\theta_F - \theta_0), \quad \Delta H_{\text{calorimètre}} = \mu c_{\text{eau}} (\theta_F - \theta_0), \quad \Delta H_{\text{résistance}} \simeq 0,$$

en négligeant la capacité thermique de la résistance. Ainsi :

$$\Delta H = (m_{\text{eau}} + \mu) c_{\text{eau}} (\theta_F - \theta_0) = \frac{U^2}{R} \tau.$$

Il vient donc :

$$\mu = \frac{U^2 \tau}{R c_{\text{eau}} (\theta_F - \theta_0)} - m_{\text{eau}} = \frac{12^2 \times 2 \times 60}{10 \times 4,18.10^3 \times (27 - 21)} - 50.10^{-3} = 1,9.10^{-2} \text{ kg}.$$

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

SYNTHÈSE

SAVOIRS

- premier principe de la thermodynamique
- fonction d'état enthalpie
- capacité thermique à pression constante
- deuxième loi de Joule
- rapport des capacités thermiques d'un gaz parfait
- loi de Mayer
- enthalpie d'une phase condensée idéale
- ordre de grandeur de la capacité thermique massique de l'eau
- enthalpie de changement d'état

SAVOIR-FAIRE

- écrire le premier principe pour un système fermé
- distinguer le statut des termes d'échange (Q et W) de celui des termes de variation (ΔU et ΔH)
- calculer un transfert thermique à partir du premier principe
- exprimer le premier principe en terme d'enthalpie pour une transformation monobare avec équilibre mécanique dans l'état initial et l'état final
- exploiter l'additivité et l'extensivité de l'énergie interne
- exploiter l'additivité et l'extensivité de l'enthalpie
- faire un bilan d'énergie avec changement d'état
- mettre en œuvre une expérience de mesure d'une grandeur thermodynamique énergétique

MOTS-CLÉS

- | | | |
|---------------------|----------------------|-----------------------|
| • énergie | • enthalpie | • transfert thermique |
| • énergie cinétique | • capacité thermique | • calorimétrie |
| • énergie interne | • travail | |

S'ENTRAÎNER

23.1 Vrai ou faux ? (★)

Justifier les réponses et rectifier les affirmations fausses si c'est possible.

1. L'énergie interne d'un gaz ne dépend que de la température.
2. Au cours d'une transformation monobare, le transfert thermique est égal à la variation d'enthalpie.
3. La température d'un corps pur augmente quand on lui fournit un transfert thermique.
4. La température d'un système isolé reste constante.
5. Un système fournit adiabatiquement du travail.
 - a. Sa température ne varie pas.
 - b. Son énergie interne diminue.
6. Un gaz est contenu dans un cylindre fermé par un piston. On diminue son volume de moitié de manière soit adiabatique, soit isotherme. Le travail à fournir est :
 - a. plus grand si la compression est adiabatique que si elle est isotherme ;
 - b. égal à la variation d'énergie interne dans les deux cas.

23.2 Valeur en eau d'un calorimètre (★)

On mélange 95 g d'eau à 20°C et 71 g d'eau à 50°C dans un calorimètre.

1. Quelle est la température finale à l'équilibre, en négligeant l'influence du calorimètre ?
2. Expérimentalement on obtient 31,3°C. Expliquer.
3. En déduire la valeur en eau du calorimètre.

23.3 Détermination de la chaleur massique du cuivre (★)

Dans un calorimètre dont la valeur en eau est de 41 g, on verse 100 g d'eau. Une fois l'équilibre thermique atteint, on mesure une température de 20°C. On plonge alors une barre métallique dont la masse est 200 g et dont la température initiale est de 60°C. A l'équilibre, on mesure une température de 30°C. Déterminer la capacité thermique massique du métal.

On donne la capacité thermique massique de l'eau $c_e = 4,18 \text{ kJ.kg}^{-1}.\text{K}^{-1}$ et on suppose que les capacités thermiques massiques sont constantes dans le domaine de températures considérées.

23.4 Transformations polytropiques (★)

Une transformation polytropique est une transformation quasistatique vérifiant PV^k constante.

1. Calculer le travail des forces de pression pour un gaz parfait subissant une transformation polytropique entre (P_0, V_0, T_0) et (P_1, V_1, T_1) en fonction des pressions et volumes ainsi que de k .

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

2. On note $\gamma = \frac{C_p}{C_v}$ qui est une constante pour un gaz parfait. Trouver une expression du transfert thermique au cours de la transformation précédente de la forme $C(T_1 - T_0)$ où C est une constante.
3. Donner une interprétation physique de C .
4. Étudier en les interprétant physiquement les cas suivants : $k = \gamma$, $k = 0$, $k \rightarrow \infty$ et $k = 1$.

23.5 Compressions et détentes successives d'un gaz parfait (★★)

On étudie les compressions et détentes successives d'un gaz parfait ($\gamma = 1,4$). On suppose que les transformations sont mécaniquement réversibles.

1. Le gaz subit une compression isotherme (température $T_0 = 273$ K) de $P_0 = 1$ bar à $P_1 = 20$ bar, puis une détente adiabatique jusqu'à P_0 . Pour cette détente adiabatique on indique que la température et la pression sont liées par la loi : $T^\gamma P^{1-\gamma} = \text{constante}$. Calculer la température T_1 après les deux transformations.
2. On recommence l'opération. Calculer T_2 puis T_n après n séries de transformations. Effectuer l'application numérique pour $n = 5$.
3. Calculer la variation d'énergie interne au cours de la $n^{\text{ième}}$ double transformation en fonction de T_0 , P_0 et P_1 . Calculer le travail W et le transfert thermique Q reçus par le gaz. Effectuer les applications numériques pour une mole de gaz et pour $n = 5$.

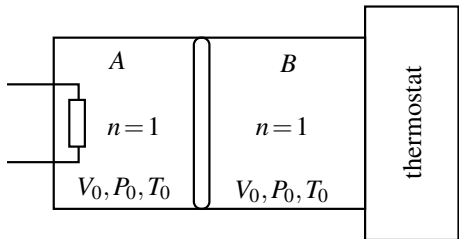
23.6 Canette autoréfrigérante (★)

A l'occasion de la Coupe du Monde de Football 2002, une canette autoréfrigérante a été mise au point. Elle comprend un réservoir en acier contenant le liquide réfrigérant. Lorsqu'on ouvre la canette, ce liquide est libéré, il se détend brusquement et se vaporise en traversant une spirale en aluminium qui serpente à travers la boisson à refroidir. Le volume de la boisson à refroidir est 33 mm^3 , et l'on considèrera pour simplifier qu'il s'agit d'eau de capacité thermique massique $c = 4,2 \text{ kJ.kg}^{-1}$. On considère que le corps réfrigérant est constitué d'une masse $m_r = 60 \text{ g}$ de N_2 dont l'enthalpie massique de vaporisation est $L_v = 200 \text{ kJ.kg}^{-1}$. Calculer la variation de température de la boisson.

APPROFONDIR

23.7 Chauffage d'une enceinte (★)

On étudie le système suivant :



On suppose que les enceintes contiennent des gaz parfaits et que l'enceinte A est parfaitement calorifugée. On note $\gamma = \frac{C_P}{C_V}$. On chauffe l'enceinte A jusqu'à la température T_1 par la résistance chauffante. Les transformations seront considérées comme quasistatiques.

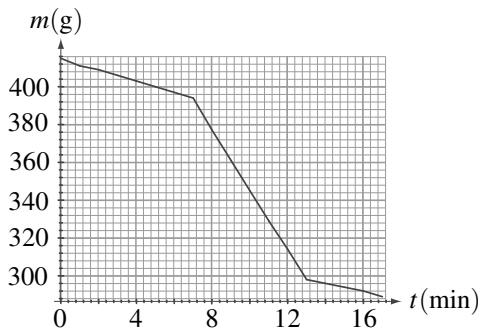
1. Déterminer les volumes finaux des deux enceintes ainsi que la pression finale.
2. Calculer la variation d'énergie interne de chacune des enceintes A et B ainsi que celle de l'ensemble $\{A + B\}$.
3. Quelle est la nature de la transformation de l'enceinte B ? En déduire le travail échangé entre les enceintes A et B et le transfert thermique Q_1 échangé entre B et le thermostat.
4. Déterminer le transfert thermique Q_2 fourni par la résistance.

23.8 Vaporisation de l'azote (★★)

On place de l'azote liquide dans un récipient posé sur une balance. On lui apporte du transfert thermique par une résistance parcourue par un courant de 2,95 A et soumise à une tension de 16,5 V.

On mesure alors l'évolution de la masse de l'azote au cours du temps et on obtient la courbe ci-contre.

1. Analyser cette évolution.
2. Montrer que la mesure de pentes sur la courbe permet de d'accéder à l'enthalpie de vaporisation de l'azote.
3. Calculer cette enthalpie de vaporisation.



23.9 Formation de la neige artificielle, d'après CCP PC (★★)

La neige artificielle est obtenue en pulvérisant de fines gouttes d'eau liquide à $T_l = 10^\circ\text{C}$ dans l'air ambiant à $T_a = -15^\circ\text{C}$.

1. Dans un premier temps la goutte d'eau supposée sphérique (rayon $R = 0,2\text{ mm}$) se refroidit en restant liquide. Elle reçoit de l'air extérieur un transfert thermique $h(T_a - T(t))$ par unité de temps et de surface, où $T(t)$ est la température de la goutte. On rappelle que la masse volumique de l'eau est $\rho = 10^3\text{ kg.m}^{-3}$ et sa capacité thermique massique $c = 4,18.10^3\text{ J.kg}^{-1}$.

- a. Établir l'équation différentielle vérifiée par $T(t)$.
- b. En déduire le temps t auquel $T(t)$ est égale à -5°C . Effectuer l'application numérique pour $h = 65\text{ W.m}^{-2}.\text{K}^{-1}$.

2. a. Lorsque la goutte atteint la température de -5°C , la surfusion cesse : la goutte est partiellement solidifiée et la température devient égale à 0°C . Calculer la fraction x de liquide restant à solidifier en supposant la transformation très rapide et adiabatique. On néglige aussi la variation de volume. L'enthalpie massique de fusion de la glace est $L_f = 335.10^3\text{ J.kg}^{-1}$.

- b. Au bout de combien de temps la goutte est-elle complètement solidifiée ?

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

CORRIGÉS

23.1 Vrai ou faux ?

1. Faux. Il faudrait que le gaz soit parfait.
2. Faux. Il faut en plus que l'état initial et l'état final soient des états d'équilibre avec l'extérieur.
3. Faux : si le corps pur est diphasé, l'apport de transfert thermique provoque un changement d'état et non une augmentation de la température. Autre exemple : au cours d'une détente isotherme d'un gaz parfait, le gaz fournit du travail, reçoit du transfert thermique, et sa température ne varie pas.
4. Faux. D'après le premier principe $\Delta U = 0$ mais U n'est pas fonction que de température en général, donc T n'est pas forcément constante. Pour un gaz parfait ou une phase condensée incompressible et indilatable en revanche, le résultat est vrai.
5. a. Faux. Adiabatique ne veut pas dire isotherme.
b. Vrai car $W < 0$ et $Q = 0$.
6. a. Vrai. Le travail à fournir est égal à l'aire sous la courbe représentative de la transformation dans le diagramme de Clapeyron (P, V). Or la pente de l'adiabatique étant supérieure à celle de l'isotherme (voir chapitre 24 page 857) à l'état initial, il faut fournir plus de travail pour comprimer le système de manière adiabatique.
b. Faux. $W = \Delta U$ uniquement si la transformation est adiabatique.

23.2 Valeur en eau d'un calorimètre

Le système considéré est constitué du calorimètre et des deux masses d'eau qu'on y met.

1. En négligeant la capacité thermique du calorimètre, le bilan énergétique se traduit par : $\Delta H_{\text{masse 1}} + \Delta H_{\text{masse 2}} = 0$ puisqu'il n'y a pas de travail ni d'échange avec l'extérieur. On obtient l'équation $m_1 c_0 (\theta_f - \theta_1) + m_2 c_0 (\theta_f - \theta_2) = 0$ soit $t_f = \frac{m_2 \theta_2 + m_1 \theta_1}{m_1 + m_2} = 32,8^\circ\text{C}$.
2. L'écart de température entre la valeur théorique et la valeur constatée prouve qu'on ne peut pas négliger la capacité thermique du calorimètre.
D'autre part, la température constatée étant inférieure à la température théorique, on peut en déduire qu'on commence par mettre l'eau la plus froide dans le calorimètre.
3. En notant m_0 la valeur en eau du récipient, le nouveau bilan énergétique s'écrit

$$(m_0 + m_1) c_0 (\theta'_f - \theta_1) + m_2 c_0 (\theta'_f - \theta_2) = 0 \quad \text{soit} \quad m_0 = m_2 \frac{\theta_2 - \theta'_f}{\theta'_f - \theta_1} - m_1 = 22,5 \text{ g.}$$

23.3 Détermination de la chaleur massique du cuivre

Le système considéré est constitué du calorimètre, de l'eau qu'il contient et de la barre métallique qu'on y introduit.

Comme il n'y a ni transfert thermique avec l'extérieur ni travail, le bilan énergétique se traduit par : $\Delta H = 0$ soit $(\mu + m_e) c_e (\theta_f - \theta_e) + m_m c_m (\theta_f - \theta_m) = 0$.

On en déduit : $c_m = c_e \frac{(\mu + m_e) (\theta_f - \theta_e)}{m_m (\theta_m - \theta_f)} = 982 \text{ J.K}^{-1}.\text{kg}^{-1}$.

23.4 Transformations polytropiques

1. Le calcul du travail pour un transformation polytropique mécaniquement réversible a été calculé au chapitre 23, page 807 : $W = \frac{1}{k-1} (P_1 V_1 - P_0 V_0) = \frac{nR}{k-1} (T_1 - T_0)$.

2. En appliquant le premier principe on en déduit :

$$Q = \Delta U + \Delta E_c - W = nC_V(T_1 - T_0) + 0 - \frac{nR}{k-1}(T_1 - T_0) = nR \left(\frac{1}{\gamma-1} - \frac{1}{k-1} \right) (T_1 - T_0),$$

$$\text{donc : } C = \frac{(k-\gamma)nR}{(\gamma-1)(k-1)}.$$

3. C est homogène à une capacité thermique.

4. Pour $k = \gamma$, $C = 0$, ce qui signifie qu'il n'y a pas d'échange thermique : la transformation est adiabatique.

Pour $k = 0$, la loi polytropique s'écrit $P = \text{constante}$ et $C = \frac{\gamma nR}{\gamma-1} = C_P$, ce qui est cohérent.

Pour k infini, la loi polytropique s'écrit $P^{\frac{1}{k}} V = V = \text{constante}$ et $C = \frac{nR}{\gamma-1} = C_V$, ce qui est bien cohérent.

Pour $k = 1$, la loi polytropique s'écrit $PV = \text{constante}$ soit $T = \text{constante}$, c'est une transformation isotherme. La formule donne C tendant vers l'infini ce qui signifie que le système peut emmagasiner tout transfert thermique sans variation de température.

23.5 Compressions et détentes successives d'un gaz parfait

Le système considéré est le gaz.

1. On décompose la première transformation en deux étapes :

- état initial : P_0, T_0 ;
- état intermédiaire : $P_1 = 20P_0, T_0$;
- état final : P_0, T_1 .

Le seconde étape est adiabatique et mécaniquement réversible donc :

$$P_1^{1-\gamma} T_0^\gamma = P_0^{1-\gamma} T_1^\gamma \Rightarrow T_1 = T_0 \left(\frac{P_1}{P_0} \right)^{\frac{1-\gamma}{\gamma}} = 115 \text{ K}.$$

$$2. \text{ On trouve de même : } T_2 = T_1 \left(\frac{P_1}{P_0} \right)^{\frac{1-\gamma}{\gamma}} = T_0 \left(\frac{P_1}{P_0} \right)^{2\frac{1-\gamma}{\gamma}}.$$

$$\text{Puis, par récurrence : } T_n = T_{n-1} \left(\frac{P_1}{P_0} \right)^{\frac{1-\gamma}{\gamma}} = T_0 \left(\frac{P_1}{P_0} \right)^{\frac{1-\gamma}{\gamma} n}.$$

L'application numérique donne, pour $n = 5$: $T_5 = 3,78 \text{ K}$. C'est une méthode qui peut sembler très efficace pour refroidir un gaz mais en fait on a supposé la transformation réversible et le gaz parfait. À très basse température, même si le système est encore gazeux, le modèle de gaz parfait est insuffisant.

3. On décompose la $n^{\text{ième}}$ transformation en deux étapes

- état initial : P_0, T_{n-1} ;
- état intermédiaire : $P_1 = 20P_0, T_{n-1}$;

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

- état final : P_1, T_n .

On note n_m le nombre de moles du gaz. La variation d'énergie interne ΔU_1 est nulle pour la transformation isotherme d'après la première loi de Joule. Pour la seconde transformation :

$$\Delta U_2 = \frac{n_m R}{\gamma - 1} (T_n - T_{n-1}) = \frac{n_m R}{\gamma - 1} T_0 \left(\frac{P_1}{P_0} \right)^{(n-1) \frac{1-\gamma}{\gamma}} \left(\left(\frac{P_1}{P_0} \right)^{\frac{1-\gamma}{\gamma}} - 1 \right)$$

Finalement, la variation d'énergie interne est : $\Delta U = \Delta U_2$ puisque $\Delta U_1 = 0$.

Pour la première transformation qui est isotherme et mécaniquement réversible, le travail est :

$$W_1 = n_m R T_n \ln \left(\frac{P_1}{P_0} \right) = n_m R T_0 \left(\frac{P_1}{P_0} \right)^{\frac{1-\gamma}{\gamma}} \ln \left(\frac{P_1}{P_0} \right)$$

et avec le premier principe :

$$Q_1 = \Delta U_1 - W_1 = -W_1$$

Pour la seconde transformation qui est adiabatique $Q_2 = 0$, donc $W_2 = \Delta U_2$. D'où finalement : $W = W_1 + W_2$ et $Q = Q_1$.

Application numérique : $\Delta U = -106 \text{ J}$, $W = -12,2 \text{ J}$ et $Q = -94,1 \text{ J}$

23.6 Canette autoréfrigérante

On considère comme système, le liquide à réfrigérer et le liquide réfrigérant. L'évolution est isobare (sous la pression atmosphérique) et rapide, donc on peut la supposer adiabatique. On exprime le premier principe à l'aide de la fonction enthalpie :

$$\Delta H_{\text{tot}} = Q \Rightarrow \Delta H_{\text{eau}} + \Delta H_{\text{refr}} = 0$$

On remplace les variations d'enthalpie par leurs expressions :

$$m_{\text{eau}} c (T_f - T_i) + m_{\text{refr}} L_v = 0 \Rightarrow T_f - T_i = -\frac{m_{\text{refr}} L_v}{m_{\text{eau}} c} = -8,7^\circ \text{C}$$

Dans la pratique, le refroidissement est de l'ordre de 10°C en 5 s. L'ordre de grandeur trouvé est donc correct.

23.7 Chauffage d'une enceinte

1. Dans l'état final, la température de l'enceinte A est T_1 puisqu'on a chauffé cette enceinte et la température de l'enceinte B toujours en contact avec le thermostat est T_0 . Par ailleurs, l'équilibre mécanique se traduit par l'égalité des pressions dans les deux enceintes soit avec l'équation des gaz parfaits $P_f = \frac{RT_1}{V_A} = \frac{RT_0}{V_B}$. On a conservation du volume c'est-à-dire que

les volumes V_A et V_B vérifient $V_A + V_B = 2V_0$. De l'égalité des pressions, on tire $V_B = V_A \frac{T_0}{T_1}$ donc en reportant dans la conservation du volume, on obtient $V_A = \frac{2V_0 T_1}{T_1 + T_0}$ puis $V_B = \frac{2V_0 T_0}{T_1 + T_0}$.

Finalement la pression finale vaut $P_f = \frac{R(T_0 + T_1)}{2V_0}$.

2. Comme les gaz sont parfaits, les variations d'énergie interne ne dépendent que des variations de température. Par conséquent, dans l'enceinte B, la variation d'énergie interne est nulle $\Delta U_B = 0$. Pour l'enceinte A, on a $\Delta U_A = C_V (T_1 - T_0)$.

De la relation de Mayer $C_P - C_V = nR$ et de la définition de γ rappelée dans l'énoncé, on obtient $C_V = \frac{nR}{\gamma-1}$ et finalement $\Delta U_A = \frac{R}{\gamma-1} (T_1 - T_0)$ (ici $n = 1$ mole). La variation d'énergie interne de l'ensemble s'obtient en appliquant l'extensivité de l'énergie interne soit :

$$\Delta U = \Delta U_A + \Delta U_B = \frac{R}{\gamma-1} (T_1 - T_0).$$

3. L'évolution de l'enceinte B est quasistatique et à température constante : elle est donc isotherme. On en déduit l'expression du travail reçu par l'enceinte B :

$$W = - \int_{V_0}^{V_B} P dV = \int_{V_0}^{V_B} RT_0 \frac{dV}{V} = -RT_0 \ln \frac{V_B}{V_0} = RT_0 \ln \frac{T_1 + T_0}{2T_0}.$$

On déduit le transfert thermique Q_1 reçu par l'enceinte B en provenance du thermostat (il n'y a pas d'échange avec l'enceinte A qui est calorifugée) en appliquant le premier principe :

$$\Delta U_B = 0 = Q_1 + W \quad \text{soit} \quad Q_1 = RT_0 \ln \frac{2T_0}{T_0 + T_1}.$$

4. On considère que le gaz de l'enceinte A et la résistance constitue le système calorifugé. L'application du premier principe au gaz de l'enceinte A donne $\Delta U_A = -W + Q_2$ où Q_2 est le transfert thermique fourni par la résistance au gaz de l'enceinte A soit :

$$Q_2 = \frac{R}{\gamma-1} (T_1 - T_0) + RT_0 \ln \frac{T_1 + T_0}{2T_0}.$$

23.8 Vaporisation de l'azote

1. Dans la première phase, on a évaporation de l'azote ; dans la seconde, on ajoute le chauffage, ce qui accélère le phénomène et dans la troisième phase, on arrête le chauffage et on retrouve la première phase.

2. On estime la puissance de fuite thermique dans la phase 1 ou 3 soit $P_f = \left(\frac{dm}{dt} \right)_0 L_V$. On effectue ensuite le bilan lorsque le chauffage est branché, ce qui permet d'obtenir $P = P_f + UI$. On en déduit $L_V = \frac{UI}{\left(\frac{dm}{dt} \right) - \left(\frac{dm}{dt} \right)_0}$.

3. Le calcul des pentes donne $\left(\frac{dm}{dt} \right) = -16 \text{ g.min}^{-1}$ et $\left(\frac{dm}{dt} \right)_0 = -3 \text{ g.min}^{-1}$. On déduit de la relation de la question précédente $L_V = 224 \text{ J.g}^{-1}$.

CHAPITRE 23 – PREMIER PRINCIPE. BILANS D'ÉNERGIE.

23.9 Formation de neige artificielle d'après CCP PC

1. a. On applique le premier principe à la goutte entre t et $t + dt$. Son volume ne varie pas donc elle ne reçoit que du transfert thermique. D'autre part, elle est liquide donc : $dU \simeq dH$, soit :

$$dU \simeq dH = \delta Q \Rightarrow \frac{4}{3}\pi R^3 \rho c dT = 4\pi R^2 h (T_a - T) dt \Rightarrow \frac{dT}{dt} + \frac{3h}{\rho R c} T = \frac{3h}{\rho R c} T_a$$

b. La solution de l'équation précédente est : $T(t) = T_a + (T_1 - T_a) \exp\left(-\frac{t}{\tau}\right)$ avec $\tau = \frac{\rho R c}{3h}$.

L'application numérique donne : $\tau = 4,3$ s et $t = \tau \ln \frac{T_1 - T_a}{T - T_a} = 3,9$ s.

2. a. On suppose d'après l'énoncé que l'eau liquide n'est pas complètement solidifiée, donc à l'état final la température du mélange est $T_f = 0^\circ\text{C}$. On imagine la suite de transformations suivante :

$$m \text{ liquide à } T_i = -5^\circ\text{C} \xrightarrow{(1)} m \text{ liquide à } T_f = 0^\circ\text{C}$$

puis :

$$m \text{ liquide à } T_f = 0^\circ\text{C} \xrightarrow{(2)} mx \text{ liquide} + m(1-x) \text{ glace à } T_f = 0^\circ\text{C}$$

On rappelle que la variation d'une fonction d'état est indépendante du chemin suivi. On applique le premier principe sur les deux transformations :

$$\Delta H = Q = 0 \Leftrightarrow mc(T_f - T_i) - m(1-x)L_f = 0 \Leftrightarrow x = 1 - \frac{c(T_f - T_i)}{L_f} = 0,94$$

On applique le premier principe à la goutte, sachant qu'elle reste à $T = 0^\circ\text{C}$ pendant la solidification. On appelle dm la masse d'eau solidifiée pendant dt et dH la variation d'enthalpie du système :

$$dH = 4\pi R^2 h (T_a - T) dt \text{ et } dH = -dm L_f \Rightarrow dm = \frac{4\pi R^2 h (T - T_a)}{L_f} dt$$

soit en intégrant :

$$m(t) = \frac{4\pi R^2 h (T - T_a)}{L_f} t + m_0$$

La masse de liquide qui reste à solidifier est : $m(t) - m_0 = \frac{4}{3}\pi R^3 \rho x$, soit après simplification :

$$t = \frac{x R L_f \rho}{3h(T - T_a)} = 21,5 \text{ s}$$

Deuxième principe. Bilans d'entropie.



Le deuxième principe de la thermodynamique fixe une condition pour le sens d'évolution d'un système et établit une différence entre les deux types d'échange d'énergie, travail et transfert thermique. Sa formulation fait apparaître une nouvelle fonction d'état du système, l'entropie, dont la valeur, mesurable à l'échelle macroscopique, traduit le caractère désordonné de la matière à l'échelle microscopique.

1 Le deuxième principe de la thermodynamique

1.1 Transformations irréversibles et transformations réversibles

a) Exemples

Détente de Joule - Gay Lussac Un récipient calorifugé rigide et indéformable comporte deux compartiments de volumes V_0 et V_1 . Dans l'état initial, l'un contient un gaz de variables d'état $(T_i, P_i, V_i = V_0)$ et l'autre est vide (voir figure 24.1). On ouvre le robinet permettant la communication entre les deux compartiments. Dans l'état d'équilibre final le gaz occupe les deux compartiments et a pour variables d'état $(T_f, P_f, V_f = V_0 + V_1)$.

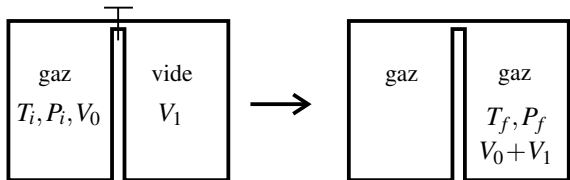


Figure 24.1 – Détente de Joule - Gay Lussac.

Cette transformation est un exemple de **transformation irréversible**. Il n'est pas possible de revenir en arrière. Pour ramener l'échantillon de gaz dans l'état initial il faut utiliser un matériel supplémentaire (compresseur).

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D’ENTROPIE.

Mise en contact avec un thermostat On considère un échantillon de métal initialement à la température 10°C que l’on met en contact avec un thermostat de température égale à 20°C. La figure 24.2(a) montre le résultat d’un calcul de simulation donnant la température θ du métal en fonction du temps (courbe gris foncé) : θ évolue de 10°C à 20°C, température imposée par le thermostat, en une durée de l’ordre de 5 minutes.

Pour ramener le métal à l’état initial on le met en contact avec un thermostat à 10°C. Le résultat de la simulation est la courbe gris clair de la figure 24.2 (a).

Les deux transformation sont-elles inverses l’une de l’autre ? La réponse est négative : par exemple la température met moins de temps pour diminuer de 20 à 19°C dans la deuxième expérience que pour augmenter 19 à 20°C dans la première. En fait, la première transformation est **irréversible** : elle ne peut être décrite « à l’envers ».

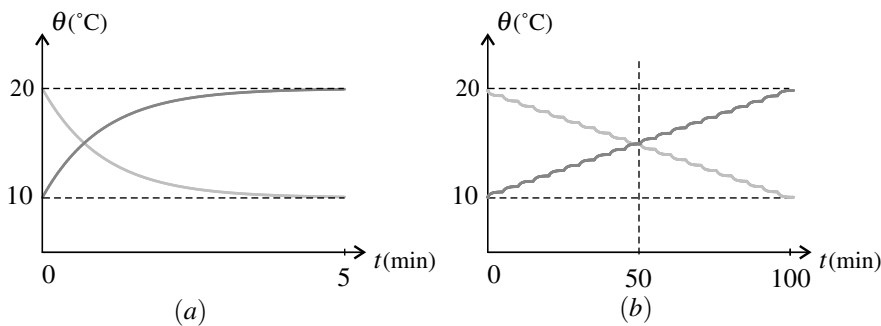


Figure 24.2 – Évolution de la température du métal entre 10°C et 20°C : (a) transformation en une étape ; (b) la transformation en 20 en étapes.

Peut-on faire évoluer le métal du même état initial vers le même état final par une transformation réversible ?

On voit sur la figure 24.2(b) la simulation d’une expérience en 20 étapes : le métal, initialement à 10°C, est mis successivement en contact, pendant 5 min, avec des thermostats de températures 10,5 - 11,0 - 11,5°C... qui augmentent à chaque étape de 0,5°C, jusqu’à 20°C. L’évolution de sa température est représentée en gris foncé. La courbe est constituée de 20 petites courbes analogues à la courbe en gris foncé de la figure 24.2(a) : la température du métal atteint la température de chacun des thermostats successifs en 5 minutes (cette durée est indépendante de la différence de température initiale entre le métal et le thermostat). Comme la différence de température entre deux thermostats successifs est petite, la différence entre la température du métal et la température du thermostat reste toujours petite.

On voit aussi sur la figure 24.2(b) l’évolution de la température du métal initialement à 20°C que l’on met en contact avec les mêmes thermostats dans l’ordre inverse (courbe en gris clair). Les courbes sont presque symétriques l’une de l’autre par rapport à la droite verticale $t = 50$ min. Les deux transformations sont quasiment identiques mais de sens inverse. Ces transformations sont presque réversibles. Cependant elles prennent un temps beaucoup plus long que les transformations avec un seul thermostat.

b) Causes d'irréversibilité

Comme on l'a vu au chapitre 22, un système initialement à l'équilibre dont on change les conditions extérieures évolue de son état d'équilibre initial i (dans lequel il ne peut plus rester) vers un état d'équilibre f . Une telle évolution est toujours irréversible.

L'irréversibilité n'est pas l'impossibilité de revenir de l'état f à l'état i , mais l'impossibilité de le faire sans changer fortement les conditions extérieures (par exemple pour la transformation de la figure 24.2(a), il faut changer de thermostat).

Tout déséquilibre d'un système thermodynamique provoque une évolution irréversible. Ainsi, une transformation est irréversible dans les circonstances suivantes :

1. Il n'y a pas équilibre mécanique : la pression n'est pas la même de part et d'autre d'une paroi mobile située à l'intérieur du système ou bien séparant le système de l'extérieur. C'est l'**irréversibilité mécanique**.
2. Il n'y a pas équilibre thermique : la température n'est pas homogène dans le système ou la température du système est différente de la température de l'extérieur avec lequel il échange du transfert thermique. C'est l'**irréversibilité thermique**.
3. Il n'y a pas équilibre de diffusion : la densité moléculaire est inhomogène ou un changement d'état a lieu à une température T sous une pression P telles que $P \neq P_{I-II}(T)$, pression d'équilibre entre les deux phases concernées.

Il existe aussi des phénomènes dont la présence rend les transformations irréversibles : l'effet Joule dans une résistance électrique, les frottements à l'intérieur d'un fluide (frottements visqueux), les frottements entre un fluide et un solide ou entre deux solides.

c) Transformations réversibles

Une transformation thermodynamique d'un système est **réversible** si :

- les contraintes extérieures varient continûment et suffisamment lentement pour que le système soit toujours à l'équilibre ;
- on peut inverser le sens de la transformation par un changement infinitésimal de ces contraintes.

La première condition définit une transformation **quasistatique**. Elle assure que les variables et fonctions d'état du système restent constamment définies tout au long de la transformation. Toute transformation réversible est quasistatique, mais la réciproque n'est pas vraie.

Une transformation réversible est nécessairement infiniment lente. C'est donc un modèle théorique de peu d'intérêt pratique *a priori*. Cependant, la théorie des machines thermiques (voir le chapitre suivant) montre qu'une machine fonctionnant avec des transformations réversibles a un rendement maximal. L'intérêt de ce modèle théorique est qu'il permet de savoir quelle performance maximale on peut atteindre. Pour se rapprocher de la performance maximale on cherche à réduire le plus possible les causes d'irréversibilité.

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

1.2 Le deuxième principe de la thermodynamique

Le deuxième principe de la thermodynamique introduit une distinction entre les transformations réversibles et les transformations irréversibles que le premier principe ne fait pas.

a) La fonction d'état entropie

Le deuxième principe de la thermodynamique suppose l'existence d'une nouvelle fonction d'état des systèmes thermodynamiques appelée **entropie**, fonction d'état **extensive** et **additive**, soit telle que :

$$S_{\Sigma_1 + \Sigma_2} = S_{\Sigma_1} + S_{\Sigma_2},$$

si on réunit deux systèmes Σ_1 et Σ_2 . Cette fonction est définie par le deuxième principe dont l'énoncé est donné ci-dessous et on peut lui donner une interprétation à l'échelle microscopique qui sera commentée à la fin de ce chapitre : il s'agit d'une mesure du désordre moléculaire.

L'entropie se mesure en $\text{J} \cdot \text{K}^{-1}$.

b) Énoncé

Le **deuxième principe** s'énonce ainsi :

Lorsqu'un système Σ subit une transformation d'un état initial i à un état final f , la variation $\Delta S = S_f - S_i$ de son entropie est :

$$\Delta S = S_{\text{éch}} + S_{\text{créée}}. \quad (24.1)$$

Dans cette formule $S_{\text{éch}}$ est un **terme d'échange** dont l'expression est :

$$S_{\text{éch}} = \frac{Q}{T_S},$$

où T_S est la température de la surface du système traversée par le transfert thermique.

Le **terme de création** $S_{\text{créée}}$ est tel que :

$$S_{\text{créée}} \geq 0,$$

avec :

- $S_{\text{créée}} > 0$, pour une transformation irréversible,
- $S_{\text{créée}} = 0$, pour une transformation réversible.

c) Cas d'un système isolé

Dans le cas d'un système isolé, système n'échangeant ni énergie ni matière avec l'extérieur, Q est nul donc le terme d'échange aussi : $S_{\text{éch}} = 0$. Par suite : $\Delta S = S_{\text{créée}} \geq 0$.

L'entropie S d'un système isolé ne peut que croître au cours d'une transformation.

Lors d'une évolution spontanée d'un système isolé, évolution nécessairement irréversible, son entropie croît. L'évolution s'arrête lorsque le système isolé a atteint un état dans lequel son entropie est maximale, qu'il ne peut plus quitter (car son entropie diminuerait).

L'état d'équilibre d'un système isolé est celui dans lequel son entropie est maximale.

d) Cas d'une transformation adiabatique

Dans le cas d'une transformation adiabatique, le transfert thermique reçu est nul donc l'entropie échangée aussi : $S_{\text{éch}} = 0$. Par suite, $\Delta S = S_{\text{créée}} \geq 0$ donc $\Delta S > 0$ si la transformation est irréversible et $\Delta S = 0$ si la transformation est réversible.

L'entropie d'un système augmente au cours d'une transformation adiabatique et irréversible.

L'entropie d'un système ne varie pas au cours d'une transformation adiabatique et réversible. Une transformation **adiabatique et réversible** est une transformation **isentropique**.

e) Remarque importante

L'équation (24.1) est relative à un système Σ et à une transformation de ce système. Chaque fois qu'on l'écrit il faut bien préciser le système choisi et la transformation considérée.

Le terme $\Delta S = S_f - S_i$ ne dépend que des états initial et final de la transformation et ne dépend pas du chemin suivi. C'est ce qui est contenu dans la notation abrégée « Δ ».

En revanche :

- le terme d'entropie échangée $S_{\text{éch}}$ dépend du chemin suivi entre l'état initial et l'état final comme le transfert thermique Q à partir duquel il se calcule ;
- par conséquent, le terme d'entropie créée $S_{\text{créée}}$ dépend aussi du chemin suivi.

Les trois termes de l'équation (24.1) n'ont pas le même statut, ce que la notation souligne par la présence ou l'absence du « Δ » devant les termes.



C'est une erreur grave de mettre un « Δ » devant $S_{\text{éch}}$ ou $S_{\text{créée}}$.

2 Entropie d'un échantillon de corps pur

On va donner les expressions de la fonction d'état entropie pour un échantillon de corps pur. Ces expressions seront systématiquement fournies au candidat si elles sont nécessaires lors des épreuves des concours. Toutefois les résultats encadrés en gris sont à connaître par cœur ou à savoir retrouver seul.

En utilisant l'extensivité et l'additivité de l'entropie, on peut calculer à partir de ces expressions l'entropie de systèmes plus compliqués.

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

2.1 Entropie d'un gaz parfait

a) Entropies molaire et massique

On considère un échantillon de gaz parfait de taille donnée par sa quantité de matière n ou sa masse m . Ce gaz parfait a un rapport des capacités thermiques molaires à pression constante et à volume constant γ **supposé indépendant de la température**. L'état du système est déterminé par les trois variables d'état (T, P, V) liées par l'équation d'état $PV = nRT$. L'entropie du système dans l'état (T, P, V) est, par extensivité :

$$S = n S_m(T, P) = m s(T, P),$$

où $S_m(T, P)$ est l'entropie molaire du gaz parfait et $s(T, P)$ l'entropie massique, à la température et sous la pression du système.

Les entropies molaire et massique sont des grandeurs intensives qui se mesurent respectivement en $\text{J} \cdot \text{K}^{-1} \cdot \text{mol}^{-1}$ et $\text{J} \cdot \text{K}^{-1} \cdot \text{kg}^{-1}$.

L'expression de l'entropie molaire en fonction des variables intensives (T, P) est :

$$S_m = \frac{R\gamma}{\gamma-1} \ln\left(\frac{T}{T_0}\right) - R \ln\left(\frac{P}{P_0}\right) + S_{m0}, \quad (24.2)$$

où T_0 et P_0 sont une température et une pression de référence et S_{m0} l'entropie molaire dans l'état (T_0, P_0) . De même :

$$s = \frac{R\gamma}{M(\gamma-1)} \ln\left(\frac{T}{T_0}\right) - \frac{R}{M} \ln\left(\frac{P}{P_0}\right) + s_0, \quad (24.3)$$

où M est la masse molaire du gaz parfait et s_0 l'entropie massique dans l'état (T_0, P_0) .

b) Entropie d'un échantillon de gaz parfait

En utilisant ce qui précède et l'équation d'état $PV = nRT$ des gaz parfaits, on peut exprimer l'entropie de l'échantillon de gaz :

- en fonction des variables d'état (T, P) :

$$S = \frac{nR\gamma}{\gamma-1} \ln\left(\frac{T}{T_0}\right) - nR \ln\left(\frac{P}{P_0}\right) + S_0, \quad (24.4)$$

en notant $S_0 = nS_{m0}$ l'entropie dans l'état (T_0, P_0) ;

- en fonction des variables d'état (T, V) :

$$S = \frac{nR}{\gamma-1} \ln\left(\frac{T}{T_0}\right) + nR \ln\left(\frac{V}{V_0}\right) + S_0, \quad (24.5)$$

où $V_0 = \frac{nRT_0}{P_0}$ et où S_0 est l'entropie dans l'état (T_0, V_0) ;

- en fonction des variables d'état (P, V) :

$$S = \frac{nR}{\gamma-1} \ln\left(\frac{P}{P_0}\right) + \frac{nR\gamma}{\gamma-1} \ln\left(\frac{V}{V_0}\right) + S_0, \quad (24.6)$$

où S_0 est l'entropie dans l'état (P_0, V_0) .

Sur toutes ces formules on voit que l'entropie augmente avec la température et le volume, ce qui se comprend intuitivement : si la température augmente la vitesse d'agitation thermique augmente donc le désordre moléculaire augmente. Si le volume accessible aux molécules augmente le désordre moléculaire augmente aussi.

En revanche, on ne peut rien dire *a priori* sur l'influence de la pression : d'après la formule (24.6) l'entropie augmente avec la pression quand on maintient le volume constant et d'après la formule (24.4) elle diminue avec la pression quand on maintient la température constante.

c) Loi de Laplace

La **loi de Laplace** est une relation vérifiée par des variables d'état d'un gaz parfait au cours d'une **transformation adiabatique et réversible**.

On a vu plus haut qu'une transformation adiabatique ($Q = 0$ donc $S_{\text{éch}} = 0$) et réversible ($S_{\text{créée}} = 0$) est une transformation **isentropique** : $\Delta S = 0$ soit $S_f = S_i$.

Si la transformation est réversible, le système est à chaque instant très proche d'un état d'équilibre, de sorte que ses variables d'état et ses fonctions d'état restent constamment définies. Ainsi, on peut dire que tout au long de la transformation adiabatique et réversible, l'entropie est définie et garde une valeur constante :

$$S = \text{constante} = S_i,$$

où S_i est sa valeur dans l'état initial.

Dans le cas d'un gaz parfait, on en déduit, en utilisant la relation (24.6) que :

$$\frac{nR}{\gamma - 1} \ln \left(\frac{P}{P_0} \right) + \frac{nR\gamma}{\gamma - 1} \ln \left(\frac{V}{V_0} \right) + S_0 = \frac{nR}{\gamma - 1} \ln \left(\frac{P_i}{P_0} \right) + \frac{nR\gamma}{\gamma - 1} \ln \left(\frac{V_i}{V_0} \right) + S_0,$$

soit, après multiplication des deux membres de l'égalité par $\frac{\gamma - 1}{nR}$ et simplification :

$$\ln \left(\frac{P}{P_0} \right) + \gamma \ln \left(\frac{V}{V_0} \right) = \ln \left(\frac{P_i}{P_0} \right) + \gamma \ln \left(\frac{V_i}{V_0} \right),$$

soit encore, en utilisant les formules $\alpha \ln x = \ln(x^\alpha)$ et $\ln x + \ln y = \ln(xy)$:

$$\ln \left(\frac{PV^\gamma}{P_0V_0^\gamma} \right) = \ln \left(\frac{P_iV_i^\gamma}{P_0V_0^\gamma} \right),$$

et finalement, en prenant l'exponentielle des deux membres (pour $x > 0$, $\exp(\ln x) = x$) et en multipliant les deux membres de l'égalité par $P_0V_0^\gamma$, on arrive à :

$PV^\gamma = P_iV_i^\gamma.$

 (24.7)

En partant de l'expression (24.4) de l'entropie et en appliquant la même méthode, on aboutit à la relation :

$T^\gamma P^{1-\gamma} = T_i^\gamma P_i^{1-\gamma}.$

 (24.8)

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

Enfin, en partant de l'expression (24.5) de l'entropie et en appliquant la même méthode on aboutit à la relation :

$$TV^{\gamma-1} = T_i V_i^{\gamma-1}. \quad (24.9)$$

Les relations (24.7), (24.8) et (24.9) sont les trois formes de la loi de Laplace, qu'il faut connaître. Cependant, un seul effort de mémoire est nécessaire car on retrouve facilement les deux autres à partir de l'une d'entre elles. Par exemple, si l'on sait que $PV^\gamma = P_i V_i^\gamma$, on trouve la relation entre T et P (par exemple) en utilisant l'équation d'état. Pour cela on remplace V par $\frac{nRT}{P}$ et V_i par $\frac{nRT_i}{P_i}$. Il vient :

$$P \left(\frac{nRT}{P} \right)^\gamma = P_i \left(\frac{nRT_i}{P_i} \right)^\gamma \quad \text{soit} \quad P^{1-\gamma} T^\gamma = P_i^{1-\gamma} T_i^\gamma,$$

après simplification du facteur $(nR)^\gamma$.

Si un gaz parfait, dont le rapport des capacités thermiques à pression et volume constants γ ne dépend pas de T , subit une transformation **adiabatique et réversible** ou bien **isentropique**, ses variables d'état (T_i, P_i, V_i) dans l'état initial, (T, P, V) dans un état quelconque au cours de la transformation et (T_f, P_f, V_f) dans l'état final vérifient :

$$\begin{aligned} P_i V_i^\gamma &= PV^\gamma = P_f V_f^\gamma, \\ T_i^\gamma P_i^{1-\gamma} &= T^\gamma P^{1-\gamma} = T_f^\gamma P_f^{1-\gamma}, \\ T_i V_i^{\gamma-1} &= TV^{\gamma-1} = T_f V_f^{\gamma-1}. \end{aligned}$$

La loi de Laplace, lorsqu'elle s'applique, peut être mise à profit pour trouver les variables d'état dans l'état final.

Elle permet aussi de tracer le chemin suivi par le système dans le diagramme de Clapeyron (P, V) pour une transformation adiabatique et réversible. La courbe correspondant est appelée **courbe adiabatique** (le mot courbe est souvent sous-entendu et on dit « une adiabatique »). Elle a une équation de la forme : $PV^\gamma = \text{constante}$.

Il faut savoir distinguer, dans le diagramme de Clapeyron, une courbe adiabatique d'une **courbe isotherme** (souvent appelée simplement « isotherme ») dont l'équation est, d'après la loi des gaz parfaits, de la forme : $PV = \text{constante}$. Les deux courbes sont des courbes décroissantes (P diminue si V augmente) mais lorsqu'elles se coupent en un point la **pente de la courbe adiabatique est plus forte en valeur absolue que la pente de la courbe isotherme**. En effet, si une adiabatique coupe une isotherme au point (P_0, V_0) (voir figure 24.3) :

- pour la courbe isotherme dont l'équation est $PV = P_0 V_0$ la pente au point (P_0, V_0) est :

$$\left(\frac{dP}{dV} \right)_{\text{isotherme}} = \frac{d}{dV} \left(\frac{P_0 V_0}{V} \right) = -\frac{P_0 V_0}{V_0^2} = -\frac{P_0}{V_0};$$

- pour la courbe isotherme dont l'équation est $PV^\gamma = P_0V_0^\gamma$ la pente au point (P_0, V_0) est :

$$\left(\frac{dP}{dV}\right)_{\text{adiabatique}} = \frac{d}{dV} \left(\frac{P_0V_0^\gamma}{V^\gamma} \right) = -\gamma \frac{P_0V_0^\gamma}{V_0^{\gamma+1}} = -\gamma \frac{P_0}{V_0}.$$

La pente de la courbe adiabatique est plus négative puisque $\gamma > 1$.

Cette propriété est illustrée sur la figure 24.3.

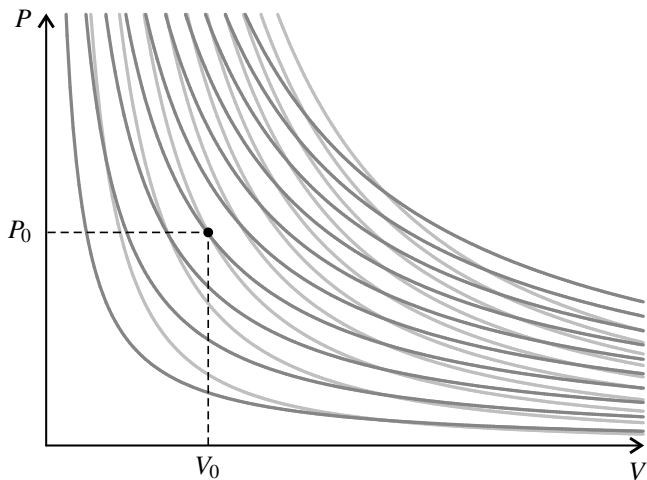


Figure 24.3 – Diagramme de Clapeyron pour un gaz parfait : les isothermes sont en gris foncé et les adiabatiques en gris clair.

2.2 Entropie d'une phase condensée indilatable et incompressible

Par extensivité, l'entropie d'un échantillon d'une phase condensée de quantité de matière n et de masse m est :

$$S = n S_m(T) = m s(T),$$

où $S_m(T)$ est l'entropie molaire et $s(T)$ l'entropie massique à la température T de l'échantillon.

L'entropie molaire S_m et l'entropie massique s d'une phase condensée et indilatable ne dépendent que de sa température (elles ne dépendent pas de la pression) et s'écrivent :

$$S_m(T) = C_m \ln \left(\frac{T}{T_0} \right) + S_{m,0}, \tag{24.10}$$

où C_m est la capacité thermique molaire et S_{m0} est l'entropie molaire à la température T_0 ;

$$s(T) = c \ln \left(\frac{T}{T_0} \right) + s_0, \tag{24.11}$$

où c est la capacité thermique massique et s_0 l'entropie molaire à la température T_0 .

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

2.3 Entropie d'un système diphasé

a) Expression de l'entropie

On considère un échantillon d'un corps pur sous deux phases *I* et *II*, de masse *m* et quantité de matière *n*. On appelle x_{II} la fraction molaire (et massique) de la phase *II* dans le système. Par additivité de l'entropie on a, comme au paragraphe 2.6 du chapitre 23, page 831 :

$$S = n(S_{m,I} + x_{II}(S_{m,II} - S_{m,I})) = m(s_I + x_{II}(s_{II} - s_I)). \quad (24.12)$$

b) Entropies de changement d'état

On appelle **entropie molaire de changement d'état** $\Delta_{I-II}S_m$ la variation d'entropie au cours de la transformation d'une mole de corps pur de l'état *I* à l'état *II* en un point du plan (*P*, *T*) où les phases *I* et *II* coexistent, soit :

$$\Delta_{I-II}S_m = S_{m,II} - S_{m,I}.$$

$\Delta_{I-II}S_m$ se mesure en $\text{J} \cdot \text{K}^{-1} \cdot \text{mol}^{-1}$.

On appelle **entropie massique de changement d'état** $\Delta_{I-II}s$ la variation d'entropie au cours de la transformation d'un kilogramme de corps pur de l'état *I* à l'état *II* en un point du plan (*P*, *T*) où les phases *I* et *II* coexistent, soit :

$$\Delta_{I-II}s = s_{II} - s_I.$$

$\Delta_{I-II}s$ se mesure en $\text{J} \cdot \text{K}^{-1} \cdot \text{kg}^{-1}$.

Les entropies molaire et massique de changement d'état ne dépendent que de la température *T* puisque la pression est imposée par la condition d'équilibre de diffusion, $P = P_{I-II}(T)$, nécessaire à la coexistence, à l'équilibre, des phases *I* et *II*.

L'entropie d'un échantillon de corps pur diphasé, comportant les phases *I* et *II*, peut se mettre sous les formes suivantes :

$$S = n(S_{m,I} + x_{II}\Delta_{I-II}S_m) = m(s_I + x_{II}\Delta_{I-II}s). \quad (24.13)$$

c) Variation d'entropie au cours d'un changement d'état isotherme et isobare

On considère une transformation d'un échantillon de corps pur de masse totale *m* passant d'un état d'équilibre initial où il se trouve dans les phases *I* et *II* et décrit par les variables d'état ($T_0, P_0, x_{II,i}$), à un état d'équilibre final du même type décrit par les variables ($T_0, P_0, x_{II,f}$).

Remarque

On a donc : $P_0 = P_{I-II}(T_0)$.

La variation d'entropie est :

$$\Delta S = n(x_{II,f} - x_{II,i})\Delta_{I-II}S_m = m(x_{II,f} - x_{II,i})\Delta_{I-II}s. \quad (24.14)$$

d) Lien entre l'enthalpie et l'entropie de changement d'état

Si on réalise la transformation en mettant en contact le système avec un milieu extérieur de même température T_0 et de même pression P_0 que lui, la transformation est *réversible* puisqu'il y a équilibre thermique, mécanique, et équilibre de diffusion (puisque $P_0 = P_{I-II}(T_0)$). Cette transformation est aussi *isobare* et *isotherme*.

D'après le deuxième principe, $S_{\text{créée}} = 0$, $\Delta S = S_{\text{éch}} + S_{\text{créée}} = S_{\text{éch}}$ et $S_{\text{éch}} = \frac{Q}{T_0}$ où Q est le transfert thermique reçu par le système. Ainsi :

$$\Delta S = \frac{Q}{T_0}.$$

D'après le paragraphe 2.3 du chapitre 23, la transformation étant isobare, le transfert thermique reçu par le système est :

$$Q = \Delta H = m(x_{II,f} - x_{II,i})\Delta_{I-II}h,$$

expression donnée par la formule (23.18), page 832. Il vient donc :

$$\Delta S = m(x_{II,f} - x_{II,i})\frac{\Delta_{I-II}h}{T_0}.$$

En comparant cette expression avec la formule (24.14) on trouve que : $\frac{\Delta_{I-II}h}{T_0} = \Delta_{I-II}s$.

L'enthalpie de changement d'état et l'entropie de changement d'état, massiques ou molaires, à la température T sont liées par la relation :

$$\Delta_{I-II}h(T) = T\Delta_{I-II}s(T) \quad \text{ou} \quad \Delta_{I-II}H_m(T) = T\Delta_{I-II}S_m(T). \quad (24.15)$$

e) Entropie et désordre moléculaire

Les enthalpies et entropies molaires de fusion d'un corps pur sont ainsi reliées par :

$$\Delta_{\text{fus}}S_m = \frac{1}{T_{\text{fus}}}\Delta_{\text{fus}}H_m.$$

Or l'expérience montre que $\Delta_{\text{fus}}H_m > 0$: il faut apporter de l'énergie (par exemple par transfert thermique) pour provoquer une fusion. Donc $\Delta_{\text{fus}}S_m > 0$: à la température de fusion, l'entropie molaire du liquide $S_{m,L}$ est supérieure à l'entropie molaire du solide $S_{m,S}$. Ceci correspond au fait que le liquide est moins ordonné que le solide (voir chapitre 21) et que l'entropie augmente avec le désordre moléculaire.

De la même manière :

$$\Delta_{\text{vap}}S_m = \frac{1}{T_{\text{vap}}}\Delta_{\text{vap}}H_m > 0 \quad \text{donc} \quad S_{m,G} > S_{m,L},$$

ce qui correspond au fait que le gaz est plus désordonné que le liquide.

Dans l'unité du système international, $\text{J}\cdot\text{K}^{-1}\cdot\text{mol}^{-1}$, l'entropie molaire d'un solide est de l'ordre de quelques unités, celle d'un liquide de l'ordre de quelques dizaines et celle d'un gaz voisine de 200.

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

3 Exemples de bilans d'entropie

3.1 Méthode générale

Faire le bilan d'entropie pour une transformation donnée d'un système Σ consiste à :

- exprimer ou calculer la variation d'entropie ΔS du système,
- exprimer ou calculer l'entropie échangée $S_{\text{éch}}$,
- en déduire l'entropie créée $S_{\text{créée}}$.

On peut ensuite conclure sur le caractère réversible ou irréversible de la transformation et analyser les causes de l'irréversibilité éventuelle.

Dans les paragraphes suivants on applique cette démarche à différents exemples.

3.2 Exemple 1 : détente de Joule - Gay Lussac

On revient sur la détente de Joule - Gay Lussac décrite au paragraphe 1.1.

On prend pour système le gaz que l'on suppose parfait, et de rapport des capacités thermiques γ indépendant de la température. On note n la quantité de gaz.

On fait l'hypothèse que la transformation est adiabatique (elle est rapide et le récipient est calorifugé). Le premier principe donne donc un renseignement sur l'état final. Il s'écrit : $\Delta U + \Delta E_c = W + Q$. On a $\Delta E_c = 0$ car le système est, à l'échelle macroscopique, au repos dans l'état initial et dans l'état final. De plus $W = 0$ parce que les parois sont indéformables, et $Q = 0$ par hypothèse. Ainsi :

$$\Delta U = 0.$$

Or, d'après la première loi de Joule, U ne dépend que de T donc si U ne varie pas, T ne varie pas non plus. Ainsi :

$$T_f = T_i.$$

Pour calculer ΔS on utilise l'expression (24.5) puisqu'on connaît T_i , V_i , T_f et V_f :

$$\begin{aligned}\Delta S &= \left(\frac{nR}{\gamma-1} \ln \left(\frac{T_f}{T_0} \right) + nR \ln \left(\frac{V_f}{V_0} \right) + S_0 \right) - \left(\frac{nR}{\gamma-1} \ln \left(\frac{T_i}{T_0} \right) + nR \ln \left(\frac{V_i}{V_0} \right) + S_0 \right) \\ &= \frac{nR}{\gamma-1} \ln \left(\frac{T_f}{T_i} \right) + nR \ln \left(\frac{V_f}{V_i} \right),\end{aligned}$$

en utilisant deux fois $\ln x - \ln y = \ln \left(\frac{x}{y} \right)$. Ainsi :

$$\Delta S = nR \ln \left(1 + \frac{V_1}{V_0} \right).$$

Comme la transformation est adiabatique, l'entropie échangée est nulle et :

$$S_{\text{créée}} = \Delta S = nR \ln \left(1 + \frac{V_1}{V_0} \right).$$

Elle est strictement positive. Comme dit plus haut, la transformation est irréversible. L'irréversibilité est due au déséquilibre mécanique.

Numériquement, pour $n = 1$ mol et $V_1 = V_0$, $\Delta S = S_{\text{créée}} = 5,76 \text{ J} \cdot \text{K}^{-1}$.

3.3 Exemple 2 : mise en contact avec un thermostat

Un échantillon de gaz de température initiale T_i est mis en contact avec un thermostat de température T_0 . Dans l'état final le gaz est à la température T_0 du thermostat. On envisage soit une évolution isochore, situation représentée sur la figure 22.1 page 796, soit une évolution isobare représentée sur la figure 22.2 page 797.

On suppose le gaz parfait et de rapport des capacités thermiques γ indépendant de la température.

On note n la quantité de matière de l'échantillon. L'état initial est caractérisé par les variables d'état (T_i, P_i, V_i) , l'état final par (T_f, P_f, V_f) avec $T_f = T_0$.

a) Cas isochore

On prend pour système le gaz et les parois du récipient qui le contient (voir figure 22.1 page 796). On supposera que ces parois ont une capacité thermique négligeable devant celle du gaz, ce qui permet d'assimiler les variations d'énergie interne et d'entropie du système aux variations de ces fonctions pour le gaz seulement.

On peut exprimer la variation d'entropie du gaz entre l'état initial et l'état final en utilisant la formule (24.5). Il vient comme plus haut :

$$\Delta S = \frac{nR}{\gamma - 1} \ln \left(\frac{T_f}{T_i} \right) + nR \ln \left(\frac{V_f}{V_i} \right).$$

Puisque $V_f = V_i$ et $T_f = T_0$ on a :

$$\Delta S = \frac{nR}{\gamma - 1} \ln \left(\frac{T_0}{T_i} \right) = C_V \ln \left(\frac{T_0}{T_i} \right).$$

Pour exprimer l'entropie échangée, on a besoin du transfert thermique que l'on obtient par application du premier principe avec $\Delta E_c = 0$: $Q = \Delta U - W$. On a : $\Delta U = C_V(T_f - T_i)$ puisque $C_V = \frac{nR}{\gamma - 1}$ est indépendante de la température comme γ et $W = 0$ car la transformation est isochore. Ainsi :

$$Q = \Delta U = C_V(T_f - T_i) = C_V(T_0 - T_i).$$

Par ailleurs la surface du système traversée par le transfert thermique est au contact du thermostat donc à la température T_0 (ceci n'aurait pas été exact si on n'avait pas mis les parois à l'intérieur du système). Ainsi :

$$S_{\text{éch}} = \frac{Q}{T_0} = C_V \left(1 - \frac{T_i}{T_0} \right).$$

On en déduit l'entropie créée :

$$S_{\text{créée}} = \Delta S - S_{\text{éch}} = C_V \left(\frac{T_i}{T_0} - 1 - \ln \left(\frac{T_i}{T_0} \right) \right),$$

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D’ENTROPIE.

en utilisant la formule $\ln\left(\frac{1}{x}\right) = -\ln x$. Numériquement : pour $n = 1$ mol et $\gamma = 1,4$ on calcule $C_V = 20,8 \text{ J}\cdot\text{K}^{-1}$, et pour $T_i = 293 \text{ K}$ et $T_0 = 373 \text{ K}$, on trouve $\Delta S = 5,0 \text{ J}\cdot\text{K}^{-1}$, $S_{\text{éch}} = 0,48 \text{ J}\cdot\text{K}^{-1}$ et $S_{\text{créée}} = 4,5 \text{ J}\cdot\text{K}^{-1}$.

L’entropie créée ne peut être que positive ou nulle. C’est bien le cas quelles que soient les températures T_0 et T_i (la capacité thermique C_V est positive) parce que, pour tout réel x , on a $x - 1 \geq \ln x$, inégalité qui exprime le fait que la courbe de la fonction logarithme est au dessous de sa tangente au point $(1, \ln 1 = 0)$ (voir figure ci-contre).

Pour T_i différent de T_0 , $S_{\text{créée}} > 0$ donc la transformation est irréversible.

L’irréversibilité est due au déséquilibre thermique entre le système et le thermostat.

Il est intéressant de regarder le comportement de $S_{\text{créée}}$ lorsque la température initiale du système est très proche de la température du thermostat. On pose alors : $T_i = T_0(1 + \varepsilon)$ avec $\varepsilon \ll 1$ et il vient :

$$S_{\text{créée}} = C_V(\varepsilon - \ln(1 + \varepsilon)) \simeq \frac{1}{2}C_V\varepsilon^2,$$

en utilisant le développement limité $\ln(1 + \varepsilon) \simeq \varepsilon - \frac{1}{2}\varepsilon^2$ (voir annexe mathématique). Cette formule montre bien que $S_{\text{créée}} > 0$ pour ε différent de 0. De plus, elle montre que le terme de création est du deuxième ordre en ε , alors que le transfert thermique, $Q = -C_V T_0 \varepsilon$, est du premier ordre.

On revient au cas où T_i n’est pas proche de T_0 . On peut porter le système à la température T_0 de manière quasiment réversible, en le mettant en contact successivement avec N thermostats de températures $T_1 = T_i + \frac{T_0 - T_i}{N}$, $T_2 = T_i + 2\frac{T_0 - T_i}{N}$, ..., $T_{N-1} = T_i + (N-1)\frac{T_0 - T_i}{N}$, $T_N = T_0$.

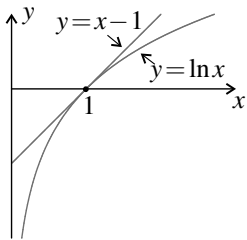
On divise ainsi par N l’écart de température entre le système et le thermostat avec lequel il est mis en contact et par N^2 la création d’entropie à chaque étape (pour N grand). Comme il y a N étapes, la création d’entropie globale tend vers 0 quand N tend vers ∞ . Pour N grand la transformation est quasiment réversible. C’est bien ce que l’on a observé sur la figure 24.2.

b) Cas isobare

Ce cas est tout à fait analogue au cas précédent. On prend pour système le gaz et les parois du récipient qui le contient (voir figure 22.2 page 797). On suppose que ces parois ont une capacité thermique négligeable devant celle du gaz, ce qui permet d’assimiler les variations d’enthalpie et d’entropie du système aux variations de ces fonctions pour le gaz seulement.

On peut exprimer la variation d’entropie du gaz entre l’état initial et l’état final en utilisant la formule (24.4) :

$$\begin{aligned} \Delta S &= \left(\frac{nR\gamma}{\gamma-1} \ln\left(\frac{T_f}{T_0}\right) - nR \ln\left(\frac{P_f}{P_0}\right) + S_0 \right) - \left(\frac{nR\gamma}{\gamma-1} \ln\left(\frac{T_i}{T_0}\right) + R \ln\left(\frac{P_i}{P_0}\right) + S_0 \right) \\ &= \frac{nR\gamma}{\gamma-1} \ln\left(\frac{T_f}{T_i}\right) - nR \ln\left(\frac{P_f}{P_i}\right). \end{aligned}$$



Puisque $P_f = P_i$ et $T_f = T_0$, on a :

$$\Delta S = \frac{nR\gamma}{\gamma-1} \ln \left(\frac{T_0}{T_i} \right) = C_P \ln \left(\frac{T_0}{T_i} \right).$$

Pour exprimer l'entropie échangée, on a besoin du transfert thermique que l'on obtient par le premier principe pour une transformation isobare : $Q = \Delta H + \Delta E_c - W_{\text{autre}}$. On a : $\Delta H = C_P(T_f - T_i)$ puisque $C_P = \frac{nR\gamma}{\gamma-1}$ est indépendante de la température comme γ , $\Delta E_c = 0$ et $W_{\text{autre}} = 0$ car seules les forces de pression travaillent. Ainsi :

$$Q = \Delta H = C_P(T_f - T_i) = C_P(T_0 - T_i).$$

Pour la même raison que dans le cas isochore :

$$S_{\text{éch}} = \frac{Q}{T_0} = C_P \left(1 - \frac{T_i}{T_0} \right).$$

On en déduit l'entropie créée :

$$S_{\text{créée}} = \Delta S - S_{\text{éch}} = C_P \left(\frac{T_i}{T_0} - 1 - \ln \left(\frac{T_i}{T_0} \right) \right).$$

La conclusion est la même : la transformation est irréversible si T_i est différent de T_0 à cause du déséquilibre thermique entre le thermostat et le système. On a une transformation réversible si on utilise un grand nombre de thermostats successifs dont les températures sont réparties entre T_i et T_0 .

3.4 Exemple 3 : compression d'un gaz parfait

a) Compression monotherme et irréversible

Un gaz, supposé parfait et de rapport des capacités thermiques γ indépendant de la température, est contenu dans un récipient maintenu à la température T_0 fermé par un piston adiabatique (c'est-à-dire qui ne laisse pas passer le transfert thermique), de surface S et de masse négligeable (voir figure 24.4).

Dans l'état initial il y a équilibre, les paramètres d'état du gaz sont (T_i, P_i, V_i) . On rompt cet équilibre en posant une masse m sur le piston. Le système évolue vers un nouvel état d'équilibre (P_f, V_f, T_f) .

On prend pour système le gaz et le piston. On appelle n la quantité de gaz. La transformation est monotherme car le gaz ne reçoit de transfert thermique que du thermostat de température T_0 . Elle n'est pas isotherme car elle est brutale et que la température du gaz n'est pas toujours définie.

On doit préciser les états initial et final du système. Dans l'état initial, comme dans l'état final, la condition d'équilibre thermique impose : $T_i = T_0$ et $T_f = T_0$.

La condition d'équilibre mécanique du piston s'écrit :

- dans l'état initial : $P_i = P_0$;
- dans l'état final : $P_f = P_0 + \frac{mg}{S}$.

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

Enfin l'équation d'état du gaz parfait donne : $V_i = \frac{nRT_0}{P_0}$ et $V_f = \frac{nRT_0}{P_0 + \frac{mg}{S}}$.

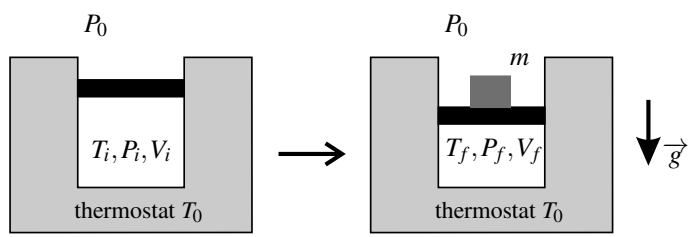


Figure 24.4 – Compression monotherme et irréversible d'un gaz.

La variation d'entropie du gaz peut être calculé à partir de l'expression (24.4). On trouve :

$$\Delta S = \frac{nR\gamma}{\gamma - 1} \ln \left(\frac{T_f}{T_i} \right) - nR \ln \left(\frac{P_f}{P_i} \right) = -nR \ln \left(1 + \frac{mg}{P_0 S} \right).$$

Le transfert thermique s'obtient par application du premier principe, avec ici $\Delta E_c = 0$: $Q = \Delta U - W$. $\Delta U = 0$ parce que $T_i = T_f$ et que l'énergie interne du gaz parfait ne dépend que de sa température (première loi de Joule), donc $Q = -W$.

Pour calculer le travail on peut remarquer que le poids de la surcharge correspond à un supplément de pression $\frac{mg}{S}$ sur le piston. La transformation est donc monobare avec $P_{\text{ext}} = P_0 + \frac{mg}{S}$ et le travail donné par la formule (22.4), page 804. Ainsi :

$$Q = -W = \left(P_0 + \frac{mg}{S} \right) (V_f - V_i) = \left(P_0 + \frac{mg}{S} \right) nRT_0 \left(\frac{1}{P_0 + \frac{mg}{S}} - \frac{1}{P_0} \right) = -nRT_0 \frac{mg}{P_0 S}.$$

On en déduit l'entropie échangée $S_{\text{éch}} = \frac{Q}{T_0}$, puisque la surface du système à travers laquelle le transfert thermique passe est à la température T_0 , soit :

$$S_{\text{éch}} = -nR \frac{mg}{P_0 S}.$$

Finalement, on obtient l'entropie créée en appliquant le deuxième principe :

$$S_{\text{créée}} = \Delta S - S_{\text{éch}} = nR \left(\frac{mg}{P_0 S} - \ln \left(1 + \frac{mg}{P_0 S} \right) \right).$$

$S_{\text{créée}} > 0$ quelle que soit la valeur non nulle de m , d'après l'inégalité $x - 1 > \ln x$ (avec $x = 1 + \frac{mg}{P_0 S}$). La transformation est donc irréversible. L'irréversibilité est due au déséquilibre mécanique du piston.

Avec $n = 1 \text{ mol}$, $S = 1,0 \cdot 10^{-2} \text{ m}^2$, $m = 10 \text{ kg}$, $g = 9,8 \text{ m} \cdot \text{s}^{-2}$ et $P_0 = 1,0 \cdot 10^5 \text{ Pa}$, on calcule : $\frac{mg}{P_0 S} = 9,8 \cdot 10^{-2}$, $\Delta S = -7,8 \cdot 10^{-1} \text{ J} \cdot \text{K}^{-1}$, $S_{\text{éch}} = -8,1 \cdot 10^{-1} \text{ J} \cdot \text{K}^{-1}$ et $S_{\text{créée}} = 3,7 \cdot 10^{-2} \text{ J} \cdot \text{K}^{-1}$.

Si la masse ajoutée est faible, plus précisément si $\frac{mg}{P_0 S} \ll 1$, on a comme plus haut :

$$S_{\text{créée}} \simeq \frac{1}{2} nR \left(\frac{mg}{P_0 S} \right)^2,$$

qui est du deuxième ordre. On doit donc pouvoir rendre la transformation réversible en ajoutant progressivement de très petites masses sur le piston.

b) Compression isotherme et réversible

Dans une autre expérience on ajoute la même masse m sur le piston mais en N étapes, en ajoutant la masse $\frac{m}{N}$ à chaque fois. On peut par exemple poser des petites billes une à une ou verser du sable très lentement (voir figure 24.5).

Comme on ajoute une très faible masse à chaque fois, le piston est toujours quasiment à l'équilibre, donc la transformation est mécaniquement réversible. Elle est aussi isotherme car le gaz, très peu perturbé à chaque étape, garde une température définie et toujours égale à T_0 .

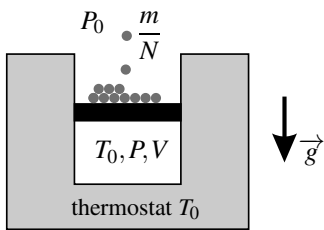


Figure 24.5 – Compression isotherme et réversible d'un gaz : un état intermédiaire.

On va faire le bilan d'entropie de cette transformation. L'état initial et l'état final sont les mêmes qu'au paragraphe précédent donc la variation d'entropie aussi.

Pour calculer le travail on peut appliquer la formule (22.7), page 807 valable pour une transformation isotherme d'un gaz parfait :

$$W_{\text{rev}} = nRT_0 \ln \left(\frac{P_f}{P_i} \right) = nRT_0 \ln \left(1 + \frac{mg}{P_0 S} \right).$$

On en déduit le transfert thermique : $Q = \Delta U + \Delta E_c - W = -W$, pour les mêmes raisons que plus haut, puis l'entropie échangée :

$$S_{\text{éch, rev}} = \frac{Q}{T_0} = nRT_0 \ln \left(1 + \frac{mg}{P_0 S} \right).$$

On constate que $S_{\text{éch, rev}} = \Delta S$ donc que l'entropie créée dans cette transformation est nulle. La transformation est réversible.

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D’ENTROPIE.

3.5 Exemple 4 : chauffage par effet Joule

On place une masse m d’eau liquide dans un calorimètre (dont on négligera la valeur en eau μ devant m) supposé parfaitement isolé. On plonge dans cette eau une résistance électrique R . Dans l’état initial l’eau est à la température T_i .

On fait passer un courant d’intensité I pendant une durée τ dans la résistance. Dans l’état final l’eau a une température T_f . Comment s’établit le bilan d’entropie ? La transformation est-elle réversible ?

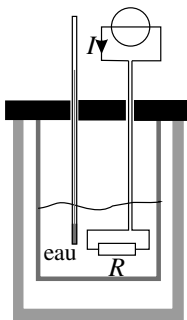


Figure 24.6 – Chauffage monobare par effet Joule.

On prend comme système le calorimètre et tout ce qu’il contient soit la masse m d’eau et la résistance. On suppose pour simplifier les écritures que les capacités thermiques de la résistance et du calorimètre sont négligeables devant celle de l’eau. Les variations d’enthalpie et d’entropie du système sont donc pratiquement égales à celles de l’eau seule.

Quel est l’état final de l’eau ? On va le déterminer en appliquant le premier principe sous la forme $\Delta H + \Delta E_c = W_{\text{autre}} + Q$, puisque la transformation est monobare, avec $\Delta E_c = 0$. $\Delta H = mc_{\text{eau}}(T_f - T_i)$ en notant c_{eau} la capacité thermique de l’eau. Le travail des forces autres que la pression est le travail électrique fourni par le générateur : $W_{\text{autre}} = UI\tau = RI^2\tau$. Enfin, le transfert thermique Q est nul si le calorimètre est très bien isolé thermiquement. Ainsi, le premier principe s’écrit :

$$mc_{\text{eau}}(T_f - T_i) = RI^2\tau \quad \text{d'où} \quad T_f = T_i + \frac{RI^2\tau}{mc_{\text{eau}}}.$$

La variation d’entropie de l’eau se calcule en utilisant l’expression (24.11) de l’entropie d’une phase condensée :

$$\Delta S = m \left(c_{\text{eau}} \ln \left(\frac{T_f}{T_0} \right) + s_0 \right) - m \left(c_{\text{eau}} \ln \left(\frac{T_i}{T_0} \right) + s_0 \right) = mc_{\text{eau}} \ln \left(\frac{T_f}{T_i} \right),$$

en utilisant $\ln x - \ln y = \ln \left(\frac{x}{y} \right)$. L’entropie échangée est nulle puisqu’il n’y a pas de transfert

thermique : $S_{\text{éch}} = 0$. Il vient finalement :

$$S_{\text{créée}} = \Delta S = mc_{\text{eau}} \ln \left(1 + \frac{RI^2 \tau}{mc_{\text{eau}} T_i} \right).$$

L'entropie créée est strictement positive, la transformation est donc irréversible.

L'irréversibilité est due à l'effet Joule.

Si $RI^2 \tau \ll mc_{\text{eau}} T_i$, $S_{\text{créée}} \simeq \frac{RI^2 \tau}{T_i} = \frac{W_{\text{autre}}}{T_i}$. Ainsi, pour une énergie dissipée par effet Joule de 300 J et $T_i = 293$ K, l'entropie créée est $S_{\text{créée}} \simeq 1 \text{ J} \cdot \text{K}^{-1}$.

3.6 Exemple 5 : solidification d'un liquide surfondu

Lorsqu'on refroidit progressivement un échantillon de corps pur liquide, dans un récipient en parfait état (il ne doit pas y avoir de rayures sur la paroi), le corps reste liquide même en dessous de la température de fusion T_{fus} du corps. Ce phénomène s'appelle la **surfusion** et on dit que le liquide est surfondu. Cet état d'équilibre est métastable : une très petite perturbation provoque la solidification d'une partie ou de la totalité du liquide surfondu.

Dans l'expérience considérée, un tube à essai contient une masse m d'eau surfondue (donc entièrement liquide) à une température T_i inférieure à $T_{\text{fus}} = 273$ K, température de fusion de l'eau sous $P_0 = 1$ bar. On fait cesser la surfusion en frappant légèrement sur le tube à essai avec un agitateur. Une partie de l'eau se solidifie presque instantanément. On se propose de faire le bilan d'entropie du système constitué par l'eau lors de cette transformation.

Dans l'état initial, l'eau est liquide à la température T_i , à la pression P_0 imposée par l'atmosphère. Dans l'état final, le système contient une masse $mx_{L,f}$ d'eau liquide et une masse $m(1 - x_{L,f})$ de glace, à la température T_{fus} , sous la pression P_0 . Il faut calculer $x_{L,f}$ pour connaître complètement l'état final.

La transformation est extrêmement rapide. On fait donc l'approximation qu'elle est adiabatique. Le système est en contact avec l'atmosphère donc elle est monobare.

Pour trouver x_L on peut appliquer le premier principe sous la forme $\Delta H + \Delta E_c = W_{\text{autre}} + Q$. Dans cette transformation $\Delta E_c = 0$, $W_{\text{autre}} = 0$ car il n'y a pas d'autre force que les forces de pression et $Q = 0$ par hypothèse. Il vient donc : $\Delta H = 0$.

Pour exprimer ΔH on utilise la méthode vue dans le paragraphe 2.7 du chapitre 23. On imagine un état intermédiaire dans lequel l'eau est totalement liquide à la température T_{fus} et on écrit :

$$\Delta H = H_f - H_i = (H_f - H_{\text{int}}) + (H_{\text{int}} - H_i) = -m(1 - x_{L,f})\Delta_{\text{fus}}h(T_{\text{fus}}) + mc_{\text{eau}}(T_{\text{fus}} - T_i),$$

en utilisant les formules (23.18) et (23.17), page 832. De la relation $\Delta H = 0$ on tire :

$$x_{L,f} = 1 - \frac{c_{\text{eau}}(T_{\text{fus}} - T_i)}{\Delta_{\text{fus}}h(T_{\text{fus}})}.$$

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

Remarque

On doit avoir $0 < x_{L,f} < 1$ ce qui suppose que : $T_{\text{fus}} - \frac{\Delta_{\text{fus}}h(T_{\text{fus}})}{c_{\text{eau}}} < T_i < T_{\text{fus}}$, soit, avec $\Delta_{\text{fus}}h(T_{\text{fus}}) = 334 \text{ kJ} \cdot \text{kg}^{-1}$ et $c_{\text{eau}} = 4,18 \text{ kJ} \cdot \text{K}^{-1} \cdot \text{kg}^{-1}$, $193 \text{ K} < T_i < 273 \text{ K}$. Cette condition est réaliste.

Pour calculer la variation d'entropie on utilise la même méthode :

$$\Delta S = S_f - S_i = (S_f - S_{\text{int}}) + (S_{\text{int}} - S_i) = -m(1 - x_{L,f})\Delta_{\text{fus}}s(T_{\text{fus}}) + mc_{\text{eau}} \ln \left(\frac{T_{\text{fus}}}{T_i} \right),$$

en utilisant les formules (24.14) et (24.11), page 858 et page 857.

La transformation étant supposée adiabatique, l'entropie échangée est nulle. Ainsi, l'entropie créée est $S_{\text{créée}} = \Delta S$. En utilisant la relation $T_{\text{fus}}\Delta_{\text{fus}}s(T_{\text{fus}}) = \Delta_{\text{fus}}h(T_{\text{fus}})$ et l'expression ci-dessus de $x_{L,f}$, on peut la mettre sous la forme :

$$S_{\text{créée}} = \Delta S = mc_{\text{eau}} \left(\frac{T_i}{T_{\text{fus}}} - 1 - \ln \left(\frac{T_i}{T_{\text{fus}}} \right) \right).$$

Il apparaît alors clairement que $S_{\text{créée}} > 0$ pour $T_i \neq T_{\text{fus}}$. La transformation est irréversible. L'irréversibilité est due au fait que la transformation de phase n'est pas faite dans les conditions de l'équilibre.

Pour $T_i = 265 \text{ K}$ et $m = 0,10 \text{ kg}$ on trouve $S_{\text{créée}} = 25 \cdot 10^{-3} \text{ J} \cdot \text{K}^{-1}$.

SYNTHÈSE

SAVOIRS

- causes d'irréversibilité
- conditions pour la réversibilité
- deuxième principe
- loi de Laplace
- définition de l'entropie de changement d'état
- relation entre l'entropie et l'enthalpie de changement d'état

SAVOIR-FAIRE

- analyser les causes d'irréversibilité
- faire un bilan d'entropie
- utiliser une expression fournie de l'entropie
- exploiter l'extensivité et l'additivité de l'entropie
- exprimer l'entropie d'un système diphasé en fonction des entropies molaires ou massiques de deux phases
- exprimer l'entropie d'un système diphasé en fonction de l'entropie molaire ou massique d'une des phases et de l'entropie de changement d'état

MOTS-CLÉS

- | | | |
|---------------------|--|---------------------|
| • réversible | • entropie échangée | • entropie massique |
| • irréversible | • entropie créée | • entropie molaire |
| • deuxième principe | • transformation adiabatique et réversible | |
| • entropie | | |

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

S'ENTRAÎNER

24.1 Bilan entropique (★)

Un morceau de fer de 2 kg, chauffé à blanc (à la température de 880 K) est jeté dans un lac à 5° C. Quelle est l'entropie créée ? On donne la capacité calorifique massique du fer : $c_{\text{fer}} = 4600 \text{ kJ.kg}^{-1}.\text{K}^{-1}$. Quelle est la cause de cette création d'entropie ?

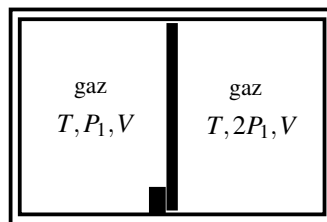
24.2 Égalisation des températures de deux systèmes (★)

Dans une enceinte thermiquement isolée, on met en contact thermique deux systèmes Σ_1 et Σ_2 constitués chacun d'un corps pur monophasé, de températures respectives T_1 et T_2 et de capacités thermiques à pression constante respectives C_1 et C_2 . La transformation est isobare.

1. Déterminer la température finale des deux systèmes.
2. Exprimer l'entropie créée dans la transformation.

24.3 Mise à l'équilibre (★)

Une enceinte indéformable aux parois calorifugées est séparée en deux compartiments par une cloison étanche de surface S , mobile et diathermane. Les deux compartiments contiennent chacun un gaz parfait. Dans l'état initial, le gaz du compartiment 1 est dans l'état ($T = 300 \text{ K}$, $P_1 = 1 \text{ bar}$, $V = 1 \text{ L}$), le gaz du compartiment 2 dans l'état ($T, 2P_1, V$), une cale bloque la cloison mobile. On enlève la cale et on laisse le système atteindre un état d'équilibre.



1. Déterminer l'état final.
2. Calculer l'entropie créée.

24.4 Mélange de deux gaz parfaits, d'après ENAC (★★)

Un récipient parfaitement isolé thermiquement est composé de deux compartiments séparés par une paroi isolante thermiquement. Initialement les compartiments contiennent des gaz parfaits différents de mêmes capacités thermiques molaires C_{Vm} et C_{Pm} . On note n_i , P_i et T_i le nombre de moles, la pression et la température du compartiment i . On enlève la paroi séparant les deux compartiments.

1. Calculer la température finale T_f du mélange qu'on considérera comme un gaz parfait.
2. Exprimer la pression finale P_f du mélange en fonction de P_1 , P_2 , T_1 , T_2 , n_1 et n_2 .
3. Exprimer la pression partielle finale $P_{f,i}$ de chaque gaz i (on rappelle que cette pression se calcule en faisant comme si le gaz i était seul) en fonction de P_f , n_1 et n_2 .

4. Montrer que la variation d'entropie du système constitué par les deux gaz, avec $n_1 = n_2 = n$, est :
$$\Delta S = nC_{Pm} \ln \frac{T_f^2}{T_1 T_2} + nR \ln \frac{P_1 P_2}{P_f^2} + 2nR \ln 2.$$

5. Commenter ce résultat dans le cas où les deux gaz sont identiques et dans le même état initial : $P_1 = P_2 = P_0$ et $T_1 = T_2 = T_0$. Que devrait valoir la variation d'entropie ?

24.5 Bilan d'entropie de l'effet Joule (★)

On considère une masse de 100 g d'eau dans laquelle plonge un conducteur de résistance $R = 20 \, \Omega$. Cette dernière est parcourue par un courant de 10 A pendant 1 s. On note Σ le système formé de l'eau et de la résistance. On donne :

- masse du conducteur : $m_c = 19 \, \text{g}$;
- capacité thermique massique du conducteur : $c_c = 0,42 \, \text{J.g}^{-1}.\text{K}^{-1}$;
- capacité thermique massique de l'eau : $c_e = 4,18 \, \text{J.g}^{-1}.\text{K}^{-1}$.

1. La température de l'ensemble est maintenue constante et égale à $20 \, ^\circ\text{C}$. Quelle est la variation d'entropie de Σ ? Quelle est l'entropie créée ? Quelle est la cause de la création d'entropie ?
2. Le même courant passe dans le conducteur pendant la même durée mais maintenant S) est isolé thermiquement. Calculer la variation d'entropie de Σ et l'entropie créée. Quelle est la cause de la création d'entropie ?

24.6 Fonte de glace dans l'eau (★)

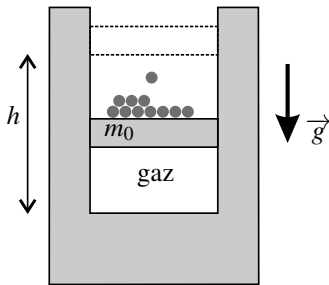
Dans un récipient parfaitement calorifugé, on met un morceau de glace à la température de 0°C dans un kilogramme d'eau initialement à la température de 20°C .

On donne la capacité thermique massique de l'eau $c = 4,2.10^3 \, \text{J.K}^{-1}.\text{kg}^{-1}$ et l'enthalpie massique de fusion de la glace $\Delta_{\text{fus}}h = 336.10^3 \, \text{J.kg}^{-1}$.

1. Déterminer la masse minimale de glace nécessaire pour que l'eau soit à la température de 0°C dans l'état final.
2. Calculer dans ce cas ΔS_e la variation d'entropie de l'eau initialement à l'état liquide.
3. Même question pour ΔS_{ge} pour l'eau initialement sous forme de glace.
4. En déduire le bilan d'entropie de l'évolution. Conclure.

24.7 Compression quasistatique ou non (★)

Un cylindre vertical à parois adiabatiques est fermé par un piston adiabatique, de section s et de masse m_0 , mobile sans frottement. Il contient un gaz parfait dont on supposera le rapport de capacités thermiques indépendant de la température et égal à $\gamma = 1,4$. Le cylindre est placé dans le vide, la pression du gaz étant équilibrée par le poids du piston. Initialement la température du gaz est $T_0 = 273 \, \text{K}$ et le piston se trouve à une hauteur $h = 10 \, \text{cm}$.



1. On ajoute progressivement de très petites masses sur le piston dont la somme est $m = 2m_0$. Exprimer, en fonction des données, la hauteur h_1 à laquelle est descendu le piston une fois l'équilibre final réalisé, ainsi que la température finale T_1 .
2. Partant du même état initial, on pose en une fois une masse m sur le piston. Exprimer, en fonction des données, la hauteur h_2 à laquelle est descendu le piston une fois l'équilibre final réalisé ainsi que la température finale T_2 .

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

APPROFONDIR

24.8 Vaporisation réversible ou irréversible (★★)

On vaporise une masse $m = 1$ g d'eau liquide des deux manières suivantes :

- la masse m est enfermée à 100°C sous la pression atmosphérique, dans un cylindre fermé par un piston. Par déplacement lent du piston, on augmente le volume à température constante et on s'arrête dès que toute l'eau est vaporisée. Le volume est alors égal à $V = 1,67$ L.
- On introduit rapidement la masse m d'eau liquide initialement à 100°C dans un récipient fermé de même température, de volume $1,67$ L, initialement vide. L'enthalpie massique de vaporisation de l'eau est $L_v = 2,25 \cdot 10^6 \text{ J} \cdot \text{kg}^{-1}$.

1. Pour chacun des processus précédents, calculer le transfert thermique fourni par le thermostat et les variations d'énergie interne, d'enthalpie et d'entropie de l'eau.
2. Calculer l'entropie créée lors du processus irréversible.

24.9 Sens d'un cycle monotherme (★)

Une mole de gaz parfait ($\gamma = 1,4$) subit la succession de transformations suivante :

- détente isotherme de $P_A = 2$ bar et $T_A = 300$ K jusqu'à $P_B = 1$ bar, en restant en contact avec un thermostat à $T_T = 300$ K ;
- évolution isobare jusqu'à $V_C = 20,5$ L toujours en restant en contact avec le thermostat à T_T ;
- compression adiabatique réversible jusqu'à l'état A.

1. Représenter ce cycle en diagramme (P, V). S'agit-il d'un cycle moteur ou récepteur ?
2. Déterminer l'entropie créée entre A et B.
3. Calculer la température en C, le travail W_{BC} et le transfert thermique Q_{BC} reçus par le gaz au cours de la transformation BC. En déduire l'entropie échangée avec le thermostat ainsi que l'entropie créée.
4. Calculer la valeur numérique de l'entropie créée au cours d'un cycle. Le cycle proposé est-il réalisable ? Le cycle inverse l'est-il ?

24.10 Entrée de matière dans un récipient (★★)

On considère un récipient vide cylindrique de volume V_1 , dont les parois sont calorifugées. On perce un trou de manière à ce que l'air ambiant (de pression P_0 et température, T_0) y pénètre de façon adiabatique (transformation rapide). On appelle V_0 le volume initialement occupé par l'air qui entre dans le récipient.

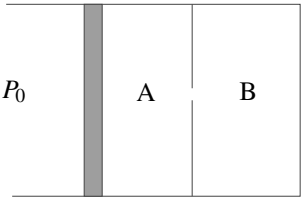
1. Calculer la température finale de l'air du récipient.
2. Déterminer l'entropie créée. Quelle est la cause de la création d'entropie ?

24.11 Détente irréversible d'un gaz parfait (★★)

Soit le dispositif de la figure ci-contre. Les parois et le piston sont adiabatiques. La paroi interne est fixe et diatherme (elle permet les échanges thermiques). Elle est percée d'un trou fermé par une fenêtre amovible.

La pression extérieure est $P_0 = 1$ bar. Initialement le volume A contient $n = 1$ mol d'un gaz parfait, dans les conditions de pression et température $P_0 = 1$ bar et $T_0 = 300$ K, et le volume B est vide.

Le rapport des capacités thermiques du gaz γ vaut 1,4.



1. On ouvre la fenêtre. Décrire qualitativement ce qui se passe suivant le volume de l'enceinte B. En déduire l'existence d'un volume critique de B : V_C que l'on ne demande pas de calculer.
2. On suppose $V_B < V_C$.
 - a. On appelle V_1 le volume final occupé par le gaz. Déterminer le travail reçu par le gaz.
 - b. Déterminer l'état final du gaz (P_1, V_1, T_1) en fonction de P_0, V_A et V_B .
 - c. Calculer l'entropie créée. Conclure. Quelle est la cause de la création d'entropie ?
 - d. Déterminer V_C . Effectuer l'application numérique.
3. Reprendre les questions 2 dans le cas $V_B > V_C$
4. On suppose maintenant que seul le piston est adiabatique, et que le dispositif est maintenu à T_0 par un thermostat. On appelle V'_C le nouveau volume critique.
 - a. Déterminez le nouvel état final pour $V_B < V'_C$.
 - b. Déterminez le nouveau volume critique V'_C .
 - c. Calculer l'entropie créée quand $V_B < V'_C$.

24.12 Étude de quelques transformations du gaz parfait, d'après CCP TSI (★)

On étudie différentes transformations de n moles d'un gaz parfait. On note P sa pression, V son volume, T sa température, C_{Vm} sa capacité calorifique molaire à volume constant, c_p sa capacité calorifique molaire à pression constante et $\gamma = \frac{c_p}{C_{Vm}}$.

1. On enferme le gaz dans une enceinte diathermane (permettant les transferts thermiques) dont une paroi de masse négligeable est mobile verticalement sans frottement. Le milieu extérieur se comporte comme un thermostat de température T_1 . Initialement le gaz est caractérisé par une pression P_1 , un volume V_1 et une température T_1 . On suppose que la paroi est dans un premier temps bloquée. On la débloque et on la déplace de manière quasistatique jusqu'à ce que le volume V'_1 offert au gaz soit égal à $2V_1$ et on la rebloque. Déterminer la pression P'_1 du gaz à l'état final.
2. Déterminer le travail W_1 mis en jeu au cours de cette transformation en fonction de n, R et T_1 .
3. Calculer la variation d'énergie interne au cours de la transformation et en déduire le transfert thermique reçu par le gaz.
4. Donner l'expression de l'entropie échangée avec le milieu extérieur.

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

5. Établir la variation d'entropie au cours de la transformation.
6. En déduire l'expression de l'entropie créée. Que peut-on en conclure ?
7. On suppose maintenant que les parois sont athermanes (ne permettant pas les échanges thermiques). Le récipient est divisé en deux compartiments de volumes identiques, le gaz se trouvant dans l'un d'eux et le deuxième étant vide. On enlève la séparation sans fournir de travail.
8. Établir l'expression de la variation d'énergie interne au cours de la transformation.
9. Déterminer les valeurs de la température T'_1 et de la pression P'_1 à l'équilibre.
10. Déterminer la variation d'entropie au cours de l'évolution.
11. Quelle est l'entropie échangée ?
12. En déduire l'entropie créée. Que peut-on en conclure ?
13. Comparer ces deux transformations.

CORRIGÉS

24.1 Bilan entropique

On peut considérer le lac comme un thermostat. Dans l'état final, le fer est à la même température que le lac.

La variation d'entropie du fer est : $\Delta S_{\text{fer}} = mc_{\text{fer}} \ln \frac{T_{\text{lac}}}{T_0} = -10,6 \cdot 10^6 \text{ J.K}^{-1}$.

L'entropie échangée est $S_{\text{éch}} = \frac{Q}{T_{\text{lac}}} = \frac{mc_{\text{fer}}(T_{\text{lac}} - T_0)}{T_{\text{lac}}} = -19,9 \cdot 10^6 \text{ J.K}^{-1}$.

L'entropie créée est donc : $S_{\text{créée}} = \Delta S_{\text{fer}} - S_{\text{éch}} = 9 \cdot 10^6 \text{ J.K}^{-1} > 0$: la transformation est irréversible.

Cette création d'entropie est due la mise en contact de deux corps à températures différentes. La température de l'ensemble est inhomogène, il y a donc un transfert thermique du morceau de fer vers le lac, ce phénomène irréversible donc crée de l'entropie.

24.2 Égalisation des températures de deux systèmes

1. Le système $\{\Sigma_1 + \Sigma_2\}$ est thermiquement isolé et a une évolution isobare. Le premier principe s'écrit donc, en notant T_f la température finale commune aux deux solides et en utilisant l'additivité de la fonction d'état enthalpie :

$$\Delta H_{\Sigma_1} + \Delta H_{\Sigma_2} = 0 \quad \text{soit} \quad C_1(T_f - T_1) + C_2(T_f - T_2) = 0.$$

On en tire : $T_f = \frac{C_1 T_1 + C_2 T_2}{C_1 + C_2}.$

2. Le système $\{\Sigma_1 + \Sigma_2\}$ est thermiquement isolé donc le deuxième principe s'écrit, en utilisant l'additivité de la fonction d'état entropie :

$$\Delta S_{\Sigma_1} + \Delta S_{\Sigma_2} = S_{\text{créée}}.$$

Quelle que soit la nature des systèmes (gaz parfait ou bien phase condensée incompressible et indilatable) leurs variations d'entropie dans cette transformation isobare s'écrivent :

$$\Delta S_1 = C_1 \ln \left(\frac{T_f}{T_1} \right) \quad \text{et} \quad \Delta S_2 = C_2 \ln \left(\frac{T_f}{T_2} \right).$$

Il vient donc finalement : $S_{\text{créée}} = C_1 \ln \left(\frac{T_f}{T_1} \right) + C_2 \ln \left(\frac{T_f}{T_2} \right).$

24.3 Mise à l'équilibre

1. Le système considéré contient : les gaz à l'intérieur des deux compartiments, l'enceinte et le piston.

- Gaz (1) : état initial (P_1, V, T) , état final (P', V'_1, T') .
- Gaz (2) : état initial (P_2, V, T) , état final (P', V'_2, T') .

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

Le nombre de moles dans le compartiment de gauche est $n_2 = 2n_1$, puisque les volumes et températures sont égaux et que $P_2 = 2P_1$.

La paroi étant diathermane, la température finale est la même dans les deux compartiments. La cloison étant mobile, la pression finale est la même pour les deux gaz. On applique le premier principe pour déterminer l'état final :

$$\Delta U_1 + \Delta U_2 + \Delta U_{\text{enceinte}} + \Delta U_{\text{cloison}} + \Delta E_{\text{ccloison}} \simeq \Delta U_1 + \Delta U_2 = W + Q$$

Le système est isolé thermiquement et l'enceinte est indéformable donc :

$$\Delta U_1 + \Delta U_2 = 0 \Leftrightarrow \frac{n_1 R}{\gamma - 1} (T' - T) + \frac{2n_1 R}{\gamma - 1} (T' - T) = 0 \Leftrightarrow T' = T$$

Pour déterminer les pressions et les volumes, on utilise l'équation d'état du gaz parfait et l'invariance du volume total, égal à $2V$. A l'état final, $P'V'_1 = n_1 RT$ et $P'V'_2 = n_2 RT$. En

calculant le rapport des deux expressions : $\frac{V'_2}{V'_1} = \frac{n_2}{n_1} = 2$. Sachant que $V'_1 + V'_2 = V$, on en

déduit $V'_1 = \frac{V}{3}$ et $V'_2 = \frac{2V}{3}$. Pour la pression, les équations d'état donnent : $P' = \frac{P_1 + P_2}{2} = 1,5 \text{ bar}$.

2. Puisque l'évolution est adiabatique, l'entropie échangée est nulle et l'entropie créée est égale à la variation d'entropie totale des gaz que l'on peut calculer en utilisant l'expression de l'entropie molaire d'un gaz parfait en fonction des variables (T, V) :

$$S_{\text{créée}} = \Delta S_T = n_1 R \ln \frac{V'_1}{V} + n_2 R \ln \frac{V'_2}{V} = 2,83 \cdot 10^{-2} \text{ J.K}^{-1}$$

L'entropie créée est strictement positive ce qui prouve que la transformation est irréversible.

24.4 Mélange de deux gaz parfaits, d'après ENAC

1. Du fait de l'isolation thermique, le transfert thermique reçu par le système $\Sigma_1 + \Sigma_2$ formé par les deux gaz est nul : $Q = 0$. D'autre part, $W = 0$ car les parois sont immobiles. Par conséquent, le premier principe s'écrit $\Delta U_{\Sigma_1 + \Sigma_2} = 0$. En utilisant l'additivité de l'énergie interne cela s'écrit :

$$n_1 C_{Vm} (T_f - T_1) + n_2 C_{Vm} (T_f - T_2) = 0 \quad \text{d'où} \quad T_f = \frac{n_1 T_1 + n_2 T_2}{n_1 + n_2}.$$

2. D'après l'équation d'état $P_1 V_1 = n_1 R T_1$, $P_2 V_2 = n_2 R T_2$ et $P_f (V_1 + V_2) = (n_1 + n_2) R T_f$. On peut donc écrire $P_f = \frac{(n_1 + n_2) R T_f}{V_1 + V_2}$ soit en explicitant l'expression de T_f :

$$P_f = \frac{(n_1 + n_2) R (n_1 T_1 + n_2 T_2)}{(n_1 + n_2) \left(\frac{n_1 R T_1}{P_1} + \frac{n_2 R T_2}{P_2} \right)} = \frac{P_1 P_2 (n_1 T_1 + n_2 T_2)}{n_1 T_1 P_2 + n_2 T_2 P_1}.$$

3. En utilisant les définitions de l'énoncé, on obtient $P_{f,1} (V_1 + V_2) = n_1 R T_f$ soit :

$$P_{f,1} = \frac{n_1}{n_1 + n_2} P_f \quad \text{et de même} \quad P_{f,2} = \frac{n_2}{n_1 + n_2} P_f.$$

4. On utilise l'expression de l'entropie molaire en fonction de la température et de la pression. Pour les deux gaz parfaits : $S_m = C_{Pm} \ln \frac{T}{T_0} - R \ln \frac{P}{P_0} + S_0$. Ainsi, en utilisant l'extensivité et l'additivité de la fonction d'état entropie on a :

$$\Delta S_{\Sigma_1 + \Sigma_2} = n_1 \left(C_{Pm} \ln \frac{T_f}{T_1} - R \ln \frac{P_{f1}}{P_1} \right) + n_2 \left(C_{Pm} \ln \frac{T_f}{T_2} - R \ln \frac{P_{f2}}{P_2} \right).$$

Dans le cas où $n_1 = n_2 = n$, $P_{f1} = P_{f2} = \frac{1}{2} P_f$, et :

$$\Delta S_{\Sigma_1 + \Sigma_2} = n C_{Pm} \ln \frac{T_f^2}{T_1 T_2} - n R \ln \frac{P_f}{2 P_1} - n R \ln \frac{P_f}{2 P_2} = C_{Pm} \ln \frac{T_f^2}{T_1 T_2} - n R \ln \frac{P_f^2}{P_1 P_2} + 2 n R \ln 2.$$

5. Avec les hypothèses proposées, on a $P_f = P_0$, $T_f = T_0$ et $\Delta S = 2R \ln 2$. La variation d'entropie n'est pas nulle alors qu'il n'y a pas de transformation : il s'agit du paradoxe de Gibbs qui se résout en corrigeant l'expression de l'entropie pour tenir compte du fait que les molécules de même formule chimique sont indiscernables.

24.5 Bilan d'entropie de l'effet Joule

1. Puisque l'eau et la résistance formant le système Σ sont complètement décrits par leur température, les fonctions d'état U , H et S ne dépendent que de T , donc ces fonctions restent constantes pendant la transformation qui est isotherme : $\Delta U = 0$, $\Delta H = 0$ et $\Delta S = 0$.

Puisque l'ensemble est maintenu à la température T , c'est que Σ est en contact avec un thermostat à la même température, donc l'entropie échangée est :

$$S_{\text{éch}} = \frac{Q}{T}$$

On applique le premier principe pour déterminer Q (il s'agit d'un liquide et d'un solide, donc $\Delta U \simeq \Delta H$, on utilise indifféremment l'une ou l'autre de ces deux fonctions) :

$$\Delta H = W_{\text{élec}} + Q = 0 \Rightarrow Q = -W_{\text{élec}} = -R i^2 t = -2 \text{ kJ}$$

d'où $S_{\text{éch}} = \frac{Q}{T} = -6,82 \text{ J.K}^{-1}$. On en déduit l'entropie créée :

$$S_{\text{créée}} = \Delta S - S_{\text{éch}} = S_{\text{éch}} = 6,82 \text{ J.K}^{-1} > 0$$

L'origine de l'irréversibilité de la transformation est l'effet Joule. Il traduit la puissance cédée par le champ électromagnétique à la matière (voir cours de seconde année). Cette puissance est toujours positive dans le cas d'un conducteur, elle correspond à un phénomène irréversible. Au niveau microscopique, on peut le modéliser par un frottement exercé sur les électrons lors de leur déplacement.

2. L'évolution est maintenant adiabatique donc : $\Delta H = W_{\text{élec}}$, soit :

$$(m_e c_e + m_c c_c)(T_f - T_i) = R I^2 t$$

La température finale est $T_f = 297,7 \text{ K}$.

Puisque l'évolution est adiabatique, l'entropie créée est égale à la variation d'entropie :

$$S_{\text{créée}} = \Delta S_T = (m_e c_e + m_c c_c) \ln \frac{T_f}{T_i} = 6,77 \text{ J.K}^{-1} > 0$$

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

24.6 Fonte de glace dans l'eau

1. Soit m la masse de glace qui aura fondu et M_e la masse initiale d'eau. Dans l'état final pour lequel on a juste la quantité de glace nécessaire pour avoir de l'eau à $0,0^\circ\text{C}$ sans qu'il reste de glace, on a une masse d'eau $m + M_e$ et plus de glace. Le bilan énergétique de l'ensemble s'écrit : $\Delta H = 0$ soit $M_e c (\theta_f - \theta_i) + mL = 0$ donc la masse de glace nécessaire est

$$m = -\frac{M_e c (\theta_f - \theta_i)}{L} = 250 \text{ g.}$$

Il s'agit de la masse minimale puisqu'au-delà de cette masse critique permettant d'avoir une température finale de $0,0^\circ\text{C}$, on a un mélange eau - glace à $0,0^\circ\text{C}$ qui est en équilibre et il n'y a donc plus d'évolution.

2. Pour l'eau, la variation d'entropie est :

$$\Delta S_e = M_e c \ln \frac{T_f}{T_i} = -297 \text{ J.K}^{-1}.$$

3. Pour la glace se transformant en eau, on a $\Delta S_{eg} = m \Delta_{\text{fus}} s = \frac{m \Delta_{\text{fus}} h}{T_0} = 308 \text{ J.K}^{-1}$.

4. Finalement la variation d'entropie pour l'ensemble est

$$\Delta S = \Delta S_e + \Delta S_{ge} = 11 \text{ J.K}^{-1}.$$

Le système étant calorifugé, il n'y a pas d'entropie échangée et l'entropie créée est égale à la variation d'entropie soit $S_{\text{cr}} = \Delta S > 0$ et la transformation est irréversible.

24.7 Compression quasistatique ou non

1. Le système Σ considéré est constitué du gaz, du piston, du cylindre et de la masse posée sur le piston.

A l'état initial, le gaz est à la pression P_0 , température T_0 , volume $V_0 = hs$. A l'état final, le gaz est à la pression P_1 , température T_1 , volume $V_1 = h_1 s$. Le piston ayant une masse m_0 et n'étant pas surmonté par un gaz, $P_0 = \frac{m_0 g}{s}$ et $P_1 = \frac{3m_0 g}{s}$ (on rappelle que pour déterminer les pressions, il suffit d'écrire l'équilibre mécanique du piston).

La transformation est adiabatique réversible, donc on peut appliquer la loi de Laplace au gaz parfait contenu dans le cylindre :

$$P_0 V_0^\gamma = P_1 V_1^\gamma \Leftrightarrow h_1 = h \left(\frac{1}{3} \right)^{\frac{1}{\gamma}} = 4,6 \text{ cm.}$$

Pour obtenir la température, on applique aussi une des lois de Laplace :

$$P_0^{1-\gamma} T_0^\gamma = P_1^{1-\gamma} T_1^\gamma \Leftrightarrow T_1 = T_0 3^{\frac{\gamma-1}{\gamma}} = 374 \text{ K.}$$

2. On étudie le même système. L'état initial est le même que précédemment et dans l'état final, le gaz est à la pression P_2 , son volume est V_2 et sa température T_2 . Puisque la masse totale déposée est la même, la pression finale P_2 est égale à P_1 . D'autre part, la masse est posée dès le début de l'évolution, donc la pression extérieure exercée sur le gaz est constante et égale à P_1 .

On applique le premier principe au système Σ :

$$\Delta U_{\text{gaz}} + \Delta U_{\text{enceinte}} + \Delta U_{\text{piston}} + \Delta E_{\text{cpiston}} = W + Q.$$

Le piston est immobile au départ et à l'arrivée et on suppose que les capacités thermiques du piston et de l'enceinte sont négligeables. La transformation est adiabatique donc $Q = 0$. La pression extérieure est constante donc $W = P_1(V_0 - V_2)$. On obtient alors :

$$\Delta U_{\text{gaz}} = \frac{nR}{\gamma - 1}(T_2 - T_0) = P_1(V_0 - V_2).$$

Or $nRT_0 = P_0V_0$ et $nRT_2 = P_1V_2$, d'où : $T_2 = \frac{3(\gamma - 1) + 1}{\gamma}T_0 = 429 \text{ K}$ et $h_2 = 5,2 \text{ cm}$.

24.8 Vaporisation réversible ou irréversible

1. Le système considéré est l'eau contenue dans le cylindre.

La première transformation est réversible. Puisque le déplacement du piston est lent, on peut supposer que la pression à l'intérieur du cylindre est maintenue en permanence égale à la pression atmosphérique. Dans ce cas la transformation a lieu à température et pression constantes. On peut alors écrire :

$$\Delta H = Q_1 = mL_v = 2,25 \text{ kJ} \quad \text{et} \quad \Delta S = \frac{Q}{T} = \frac{\Delta H}{T} = 6 \text{ J.K}^{-1}$$

On calcule maintenant la variation d'énergie interne :

$$\Delta U = \Delta H - \Delta(PV) = \Delta H - P_s(V_{\text{final}} - V_{\text{initial}}) \simeq \Delta H - P_s V_{\text{final}} = 2,1 \text{ kJ}$$

La deuxième transformation n'est pas réversible, mais les états initial et final de l'eau sont les mêmes que pour la première transformation. Comme les variations des fonctions d'état ne dépendent que des états extrêmes, les variations de U , H et S sont les mêmes pour les deux transformations. Dans cette transformation, par contre la pression n'est pas constante mais, puisque l'enceinte est supposée de volume constant, le travail des forces de pression est nul. En appliquant le premier principe, on trouve :

$$\Delta U = Q_2 \Rightarrow Q_2 = \Delta H - P_s V_{\text{final}} = mL_v - P_s V_{\text{final}}$$

On remarque sur ces deux exemples que le transfert thermique dépend du chemin suivi.

- 2. • Premier cas : $S_{\text{éch}} = \frac{Q_1}{T} = \frac{mL_v}{T}$ et $S_{\text{créée}} = \Delta S - S_{\text{éch}} = 0$, et on retrouve que la transformation est réversible.
- Deuxième cas : $S_{\text{éch}} = \frac{Q_2}{T} = \frac{mL_v - P_s V_{\text{final}}}{T}$ et $S_{\text{créée}} = \Delta S - S_{\text{éch}} = \frac{P_s V_{\text{final}}}{T}$, et on retrouve que la transformation est irréversible puisque l'entropie créée est positive.

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

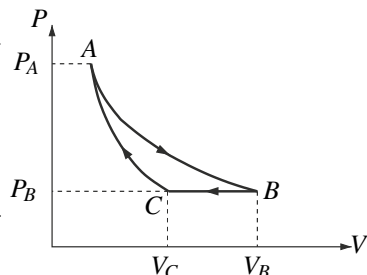
24.9 Sens d'un cycle monotherme

Le système étudié est le gaz qui décrit le cycle.

1. Pour calculer V_B , on applique l'équation du gaz parfait en B sachant que $T_B = T_A$, soit :

$$V_B = \frac{nRT_A}{P_B} = 24,9 \text{ L}$$

On a donc $V_B > V_C$. Le cycle tourne dans le sens horaire : c'est un cycle moteur.



2. En utilisant l'expression de l'entropie en fonction des variables (T, V) ou en fonction des variables (T, P) on trouve, pour la transformation isotherme AB : $\Delta S_{AB} = nR \ln \frac{V_B}{V_A} = nR \ln \frac{P_A}{P_B}$.

La transformation étant isotherme et le gaz parfait, $\Delta U_{AB} = 0$ donc $Q = -W$. L'évolution du gaz est de plus mécaniquement réversible et isotherme, dans ce cas le calcul fait au chapitre 23 donne : $W = -nRT_A \ln \frac{V_B}{V_A} = nRT_A \ln \frac{P_B}{P_A}$. Finalement : $Q = -nRT_A \ln \frac{P_B}{P_A}$.

On en déduit l'entropie échangée avec le thermostat à $T_T = T_A$: $S_{\text{éch}} = -nR \ln \frac{P_B}{P_A}$.

On remarque donc que $S_{\text{éch}} = \Delta S$: l'entropie créée est nulle ce qui prouve que le processus est réversible.

3. Pour calculer T_C , on applique l'équation du gaz parfait : $T_C = \frac{P_B V_C}{nR} = 246,6 \text{ K}$.

Il s'agit d'une transformation isobare donc : $W_{BC} = P_B(V_B - V_C) = 440 \text{ J}$; et avec le premier principe : $Q_{BC} = \Delta H_{BC} = \frac{nR\gamma}{\gamma-1}(T_C - T_B) = -1,55 \text{ kJ}$.

On peut aussi utiliser $Q_{BC} = \Delta U_{BC} - W_{BC}$, on obtient bien évidemment le même résultat.

On déduit du calcul précédent : $S_{\text{éch}} = \frac{Q}{T_T} = -5,17 \text{ J.K}^{-1}$.

Pour calculer l'entropie créée, il faut d'abord calculer la variation d'entropie du gaz entre B et C . En utilisant l'expression de l'entropie du gaz parfait en variables (T, P) (puisque la transformation est isobare) on a : $\Delta S_{BC} = \frac{nR\gamma}{\gamma-1} \ln \frac{T_C}{T_B} = -5,7 \text{ J.K}^{-1}$.

On peut alors calculer l'entropie créée :

$$S_{\text{créée}} = \Delta S - S_{\text{éch}} = -0,54 \text{ J.K}^{-1}.$$

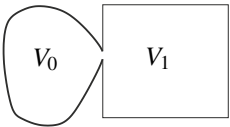
Ce résultat est en contradiction avec le second principe, puisque l'entropie créée ne peut être que positive : cette transformation est donc irréalisable.

4. Ce cycle est donc irréalisable pratiquement puisque l'entropie créée serait négative dans une des transformations. En revanche le cycle récepteur correspondant ($ACBA$) est possible. On verra effectivement dans le chapitre sur les machines thermiques qu'un cycle moteur monotherme (c'est-à-dire en contact avec un seul thermostat) ne peut exister.

24.10 Entrée de matière dans un récipient

1. On considère comme système l'air qui va entrer dans le récipient et le récipient lui-même dont on suppose la capacité thermique négligeable.

Le gaz entre dans le récipient jusqu'à ce que la pression soit égale à la pression extérieure P_0 en étant poussé par le gaz extérieur. Le gaz qui entre dans le récipient occupe, quand il est à l'extérieur, le volume V_0 . Il passe de la température T_0 (à l'extérieur) à la température T_1 (à l'intérieur). Le travail des forces de pression est le travail exercé par le gaz extérieur, et le volume balayé par le pourtour du gaz est V_0 . On en déduit que le travail reçu par le gaz est : $W = P_0 V_{\text{balayé}} = P_0 V_0$.



On applique le premier principe au système :

$$\Delta U = W + Q = W \Rightarrow \frac{nR}{\gamma - 1}(T_1 - T_0) = P_0 V_0.$$

Or $P_0 V_0 = nRT_0$, d'où $(T_1 - T_0) = (\gamma - 1)T_0$ et $T_1 = \gamma T_0$.

2. Comme l'évolution est adiabatique, l'entropie échangée est nulle et l'entropie créée est égale à la variation d'entropie du système. Pour calculer celle-ci, on retrouve rapidement l'expression de l'entropie d'un gaz parfait en variables (T, P) , en utilisant l'identité thermodynamique sous la forme $dH = TdS + VdP$ avec $dH = nC_{pm}dT = \frac{n\gamma R}{\gamma - 1}dT$ et $dP = 0$ puisque

la pression initiale est égale à la pression finale. On obtient : $dS = \frac{n\gamma R}{\gamma - 1} \frac{dT}{T}$ soit :

$$S_{\text{créée}} = \Delta S = \frac{n\gamma R}{\gamma - 1} \ln \frac{T_1}{T_0} = \frac{nR\gamma}{\gamma - 1} \ln \gamma > 0$$

donc l'évolution est irréversible comme on pouvait s'y attendre.

L'origine de l'irréversibilité est la différence de densité entre l'intérieur et l'extérieur du récipient.

24.11 Détentes irréversibles d'un gaz

1. Dès que la fenêtre est ouverte le gaz de l'enceinte A diffuse vers l'enceinte B . Si l'enceinte B est suffisamment grande (V_B supérieure à un volume critique appelé V_C), tout le gaz de A se retrouve dans B et dans l'état final, le piston est contre la paroi. Si $V_B < V_C$, il reste du gaz dans A et le piston s'immobilise avant la cloison. Dans ce cas, comme le piston est à l'équilibre dans l'état final, la pression dans le gaz est P_0 .

2. a. Le travail reçu par le gaz est dû au déplacement du piston, il est égal à $P_0 V_{\text{balayé}}$. Le volume balayé est : $V_{\text{balayé}} = V_A + V_B - V_1$, donc le travail est : $W = P_0(V_A + V_B - V_1)$.

b. Puisque $V_B < V_C$, la pression finale du gaz est P_0 . Le système Σ considéré est constitué du gaz, du cylindre et du piston. L'état initial et l'état final du gaz sont :

Etat initial	$\left \begin{array}{c} P_0 \\ V_A \\ T_0 \end{array} \right.$	Etat final	$\left \begin{array}{c} P_1 = P_0 \\ V_1 \\ T_1 \end{array} \right.$
--------------	---	------------	---

L'application du premier principe à Σ donne :

$$\Delta U_{\text{gaz}} + \Delta U_{\text{enceinte}} + \Delta U_{\text{piston}} + \Delta E_{\text{cpiston}} = W + Q$$

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D'ENTROPIE.

On suppose que les capacités thermiques du piston et de l'enceinte sont négligeables, donc $\Delta U_{\text{piston}} \simeq 0$ et $\Delta U_{\text{enceinte}} \simeq 0$. Le piston est immobile dans les états extrêmes donc $\Delta E_{\text{cpiston}} = 0$. Enfin la transformation est adiabatique donc $Q = 0$. Finalement :

$$\Delta U = W \quad \text{soit} \quad \frac{nR}{\gamma-1}(T_1 - T_0) = P_0(V_A + V_B - V_1)$$

L'équation du gaz parfait donne $nRT_1 = P_0V_1$. En combinant les deux expressions précédentes, on obtient pour T_1 : $T_1 = \frac{P_0}{nR} \left(\frac{\gamma-1}{\gamma} V_B + V_A \right)$.

Et pour V_1 : $V_1 = \frac{\gamma-1}{\gamma} V_B + V_A$.

c. On trouve en utilisant l'expression de l'entropie en variables (T, P) (puisque la pression est la même dans l'état final et l'état initial) : $\Delta S = \frac{nR\gamma}{\gamma-1} \ln \frac{T_1}{T_0}$. Et, avec l'expression de T_1

trouvée plus haut : $\Delta S = \frac{nR\gamma}{\gamma-1} \ln \left(\frac{P_0}{nRT_0} \left(\frac{\gamma-1}{\gamma} V_B + V_A \right) \right) = \frac{nR\gamma}{\gamma-1} \ln \left(1 + \frac{\gamma-1}{\gamma} \frac{V_B}{V_A} \right)$.

Comme l'évolution est adiabatique, l'entropie échangée est nulle et donc l'entropie créée est égale à la variation d'entropie : $S_{\text{créée}} = \Delta S = \frac{nR\gamma}{\gamma-1} \ln \left(1 + \frac{\gamma-1}{\gamma} \frac{V_B}{V_A} \right) > 0$.

La transformation est bien irréversible. L'origine de l'irréversibilité est la différence de pression entre l'intérieur et l'extérieur du récipient A (irréversibilité mécanique).

d. A la limite, si $V_B = V_C$, le volume final V_1 est égal à V_B , soit : $\frac{\gamma-1}{\gamma} V_C + V_A = V_C$, ce qui donne : $V_C = \gamma V_A = \gamma \frac{nRT_0}{P_0} = 0,035 \text{ m}^3 = 35 \text{ L}$.

3. Dans ce cas, on ne connaît pas la pression dans l'état final puisque le piston est retenu par la cloison, mais on connaît le volume final. L'état initial et l'état final du gaz sont :

$$\begin{array}{cc} \text{Etat initial} & \left| \begin{array}{c} P_0 \\ V_A \\ T_0 \end{array} \right. & \text{Etat final} & \left| \begin{array}{c} P_1 \\ V_B \\ T_1 \end{array} \right. \end{array}$$

Le volume balayé est V_A donc le travail est $W = P_0V_A$. Le premier principe donne :

$$\Delta U_{\text{gaz}} = W \Rightarrow \frac{nR}{\gamma-1}(T_1 - T_0) = P_0V_A$$

Or, $P_0V_A = nRT_0$, donc : $T_1 = \gamma T_0$ et $P_1 = \frac{nRT_1}{V_B} = \gamma P_0 \frac{V_A}{V_B}$.

On remarque que $P_1 < P_0$, ce qui confirme la condition de validité de l'hypothèse $V_B > V_C$.

Pour calculer la variation d'entropie, on utilise l'expression de l'entropie d'un gaz parfait en variables (T, V) (par exemple) :

$$\Delta S = \frac{nR}{\gamma-1} \ln \frac{T_1}{T_0} + nR \ln \frac{V_B}{V_A} = \frac{nR}{\gamma-1} \ln \frac{T_1 V_B^{\gamma-1}}{T_0 V_A^{\gamma-1}} = \frac{nR}{\gamma-1} \ln \left(\gamma \left(\frac{V_B}{V_A} \right)^{\gamma-1} \right)$$

L'évolution est adiabatique, donc l'entropie créée est égale à ΔS . On sait que $V_B > V_C$ donc $V_B > \gamma V_A$, et : $\frac{V_B}{V_A} > \gamma \Rightarrow \gamma \left(\frac{V_B}{V_A} \right)^{\gamma-1} > \gamma^{\gamma-1} > 1$. L'argument du logarithme est plus grand que 1, $S_{\text{créée}} > 0$ et la transformation est irréversible, pour la même raison que précédemment.

4. a. Le système est le même que précédemment et puisque $V_B < V'_C$, le piston s'arrête avant la cloison donc la pression finale est P_0 . L'état initial et l'état final du gaz sont :

$$\text{Etat initial} \left| \begin{array}{c} P_0 \\ V_A \\ T_0 \end{array} \right. \qquad \text{Etat final} \left| \begin{array}{c} P_0 \\ V_1 \\ T_0 \end{array} \right.$$

On en déduit le volume final : $V_1 = \frac{nRT_0}{P_0} = V_0$. Le piston n'a fait que transvaser le gaz sans changer son état.

b. Le nouveau volume critique vérifie $V_1 = V'_C$, soit $V'_C = V_A$.

c. Puisque les variables d'état du gaz sont inchangées, la variation d'entropie du gaz est nulle : $\Delta S = 0 \Rightarrow S_{\text{créée}} = -S_{\text{éch}}$.

On calcule l'entropie échangée avec le thermostat à la température T_0 : $S_{\text{éch}} = \frac{Q}{T_0}$.

On utilise le premier principe pour calculer Q . Sachant que $\Delta U = 0$ puisque T est constante : $Q = -W = -P_0 V_A = -nRT_0$, donc : $S_{\text{éch}} = -nR$ et $S_{\text{créée}} = nR > 0$. La transformation est toujours irréversible.

24.12 Étude de quelques transformations d'un gaz parfait d'après CCP TSI

1. L'état final est défini par $T'_1 = T_1$ et $V'_1 = 2V_1$. De l'équation d'état, on en déduit $P'_1 = \frac{P_1}{2}$.
2. La transformation est quasistatique donc on peut identifier la pression à la pression extérieure, ce qui permet d'écrire $\delta W = -PdV = -nRT_1 \frac{dV}{V}$ en utilisant le fait que la transformation est isotherme. On a donc $W_1 = -nRT_1 \ln 2$.
3. La variation d'énergie interne est nulle puisque $\Delta U_1 = nC_{Vm}\Delta T = 0$ du fait du caractère isotherme de la transformation. On en déduit le transfert thermique $Q_1 = -W_1 = nRT_1 \ln 2$.
4. L'extérieur étant un thermostat, l'entropie échangée vaut $S_{\text{éch}} = \frac{Q_1}{T_1} = nR \ln 2$.
5. D'après la relation de la question 4, la variation d'entropie vaut $\Delta S_1 = nR \ln 2$.
6. Des deux questions précédentes, on déduit l'expression de l'entropie créée $S_{\text{cr}} = \Delta S - S_{\text{éch}} = 0$. La transformation est donc réversible.
7. Les parois sont fixes donc il n'y a pas de travail $W = 0$. De plus, l'enceinte étant calorifugée, il n'y a pas de transfert thermique $Q = 0$. On en déduit que la variation d'énergie interne est nulle $\Delta U_2 = 0$.
8. Comme les gaz sont parfaits, la nullité de la variation d'énergie interne entraîne l'absence de variation de température $\Delta T = 0$ soit $T_1 = T'_1$. L'état final vérifie d'autre part $V'_1 = 2V_1$ donc par la relation d'état, on en déduit $P'_1 = \frac{P_1}{2}$.
9. Pour un gaz non parfait, l'énergie interne U ne dépend pas uniquement de la température donc la variation de température ne serait pas nulle $\Delta T \neq 0$.
10. On a la même chose qu'au cas précédent puisque les états initial et final sont les mêmes.
11. L'entropie échangée est nulle $S'_{\text{éch}} = 0$ car le système est calorifugé.
12. On en déduit l'entropie créée $S'_{\text{cr}} = \Delta S_2 > 0$ donc la transformation est irréversible.
13. Les transformations correspondent aux mêmes états initial et final mais la transformation est réversible dans un cas et pas dans l'autre.

CHAPITRE 24 – DEUXIÈME PRINCIPE. BILANS D’ENTROPIE.

Machines thermiques

25

Les machines thermiques sont des systèmes thermodynamiques avec lesquels on modélise de nombreux appareils et installations réels : moteurs à essence et Diesel, réfrigérateurs, pompes à chaleur, centrales électriques thermiques, usines d’incinération...

1 Machine monotherme

Une machine thermique est un système thermodynamique (M) échangeant du travail avec un système mécanique SM (ou électrique) et du transfert thermique avec un ou plusieurs thermostats au cours de transformations successives formant un cycle : quand il a subi toutes les transformations, le système est revenu dans son état initial.

La machine thermique la plus simple échange du transfert thermique avec un unique thermostat TH de température T_0 . On note W le travail algébrique qu’elle reçoit de la part du système mécanique (ou électrique) SM et Q le transfert thermique algébrique qu’elle reçoit de la part du thermostat au cours du cycle (voir figure 25.1) .

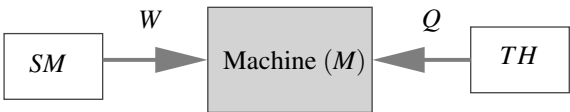


Figure 25.1 – Schéma synoptique d’une machine monotherme. W et Q sont positifs s’ils sont effectivement reçus par la machine (M) et négatifs s’ils sont cédés par (M).

On applique les deux principes de la thermodynamique à la machine (M) sur le cycle. Pour cette transformation, l’état final est identique à l’état initial i , donc les variations des fonctions d’état de (M) sont nulles : $\Delta U = U_i - U_i = 0$ et $\Delta S = S_i - S_i = 0$. Il vient donc :

$$\begin{cases} \Delta U = W + Q = 0 \\ \Delta S = \frac{Q}{T_0} + S_{\text{créée}} = 0, \end{cases}$$

car la température de la surface du système en contact avec le thermostat est à la température

CHAPITRE 25 – MACHINES THERMIQUES

du thermostat T_0 . On en tire :

$$Q = -T_0 S_{\text{créée}} \leq 0 \quad \text{et} \quad W = T_0 S_{\text{créée}} \geq 0.$$

Ainsi, la machine ne peut que recevoir du travail et donner du transfert thermique. Il s’agit par exemple d’un radiateur électrique qui reçoit du travail de l’installation électrique et fournit du transfert thermique à la pièce.

Si l’on veut une machine pouvant fournir du travail il faut nécessairement au moins deux sources. La suite du chapitre sera consacrée aux machines dithermes.

2 Machines thermiques dithermes

2.1 Généralités sur les machines dithermes

a) Notations

Une machine ditherme (M) échange du transfert thermique avec deux thermostats :

- un thermostat TH_{ch} de température T_{ch} appelé **source chaude** ;
- un thermostat TH_{fr} de température T_{fr} appelé **source froide**.

Comme le vocabulaire employé l’indique on suppose que : $T_{ch} > T_{fr}$.

On note W le travail algébrique reçu par la machine de la part du système SM , Q_{ch} et Q_{fr} les transferts thermiques reçus par la machine de la part des thermostats TH_{ch} et TH_{fr} respectivement. Les conventions de signe pour ces échanges énergétiques sont schématisées sur la figure 25.2.

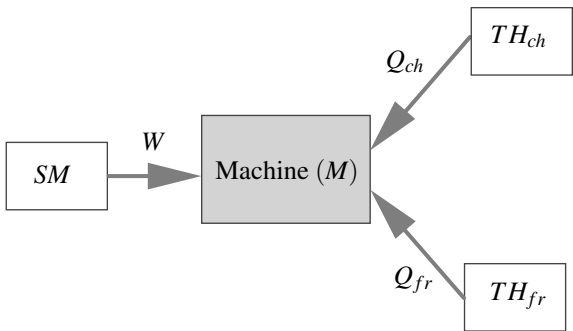


Figure 25.2 – Schéma synoptique d’une machine ditherme. W , Q_{ch} et Q_{fr} sont positifs s’ils sont effectivement reçus par la machine (M) et négatifs s’ils sont cédés par (M).

b) Inégalité de Clausius

On applique les deux principes de la thermodynamique à la machine (M). La transformation étant un cycle, les variations des fonctions d’état de (M) sont nulles. Il vient donc :

$$\begin{cases} \Delta U = W + Q_{ch} + Q_{fr} = 0 \\ \Delta S = \frac{Q_{ch}}{T_{ch}} + \frac{Q_{fr}}{T_{fr}} + S_{\text{créée}} = 0, \end{cases}$$

car la température de la surface du système à travers laquelle il reçoit le transfert thermique d'un thermostat est à la température de ce thermostat.

Le deuxième principe précise que $S_{cr   } \geq 0$, on a donc l'**in  galit   de Clausius** :

$$\frac{Q_{ch}}{T_{ch}} + \frac{Q_{fr}}{T_{fr}} \leq 0. \quad (25.1)$$

Cette in  galit   est une   galit   si et seulement si le cycle est une suite de transformations r  versibles (on dit plus rapidement que le cycle est r  versible).

c) Les deux types de machines dithermes

La premi  re possibilit   est que la machine re  oive de l'  nergie de la source chaude, en donne    la source froide. Ce qu'elle re  oit en plus par rapport    ce qu'elle c  de est transform   en travail : la machine est dans ce cas un moteur.

Un **moteur thermique** fournit du travail ($W < 0$) et l'  change thermique a lieu dans le sens naturel : la source chaude donne du transfert thermique au moteur ($Q_{ch} > 0$) tandis que la source froide re  oit du transfert thermique du moteur ($Q_{fr} < 0$).

La deuxi  me possibilit   est que la machine re  oive du travail. Il est alors possible qu'elle donne du transfert thermique    la source chaude et re  oive du transfert thermique de la source froide. La machine sert dans ce cas    chauffer la source chaude ou bien    refroidir la source froide.

Une machine thermique destin  e    refroidir (**machine frigorifique, climatiseur**) ou bien    chauffer (**pompe    chaleur**) re  oit du travail ($W > 0$), c  de du transfert thermique    la source chaude ($Q_{ch} < 0$) et prend du transfert thermique    la source froide ($Q_{fr} > 0$). Ce **transfert thermique de sens contraire au sens naturel** n  cessite l'apport de travail    la machine.

2.2 Moteur thermique

a) Rendement du moteur

La d  finition g  n  rale d'un **rendement** est :

$$\rho = \left| \frac{\text{  nergie utile}}{\text{  nergie co  teuse}} \right|.$$

Pour un moteur, c'est le travail qui est utile et c'est la source chaude qui est on  reuse. On d  finit donc le rendement du moteur de la mani  re suivante :

$$\rho_{\text{moteur}} = \left| \frac{W}{Q_{ch}} \right| = -\frac{W}{Q_{ch}}, \quad (25.2)$$

puisque $W < 0$ et $Q_{ch} > 0$. D'apr  s le premier principe, $W = -(Q_{ch} + Q_{fr})$, l'expression pr  c  dente devient donc :

$$\rho_{\text{moteur}} = 1 + \frac{Q_{fr}}{Q_{ch}}.$$

CHAPITRE 25 – MACHINES THERMIQUES

En multipliant l'équation du second principe par $\frac{T_{fr}}{Q_{ch}}$ on trouve : $\frac{T_{fr}}{T_{ch}} + \frac{Q_{fr}}{Q_{ch}} + \frac{T_{fr}S_{cr    }}{Q_{ch}} = 0$.
On en d  duit :

$$\rho_{\text{moteur}} = 1 - \frac{T_{fr}}{T_{ch}} - \frac{T_{fr}S_{cr    }}{Q_{ch}}. \tag{25.3}$$

b) Th  or  me de Carnot

D'apr  s le deuxi  me principe $S_{cr    } \geq 0$ et pour un moteur $Q_{ch} > 0$ donc $\frac{T_{fr}S_{cr    }}{Q_{ch}} \geq 0$ et d'apr  s (25.3) :

$$\rho_{\text{moteur}} \leq 1 - \frac{T_{fr}}{T_{ch}}.$$

L'  galit   est r  alis  e si et seulement si $S_{cr    } = 0$, c'est-  -dire si le cycle est r  versible. Ce r  sultat est le **th  or  me de Carnot** :

Le rendement d'un moteur ditherme r  versible est :

$$\rho_{\text{moteur,rev}} = 1 - \frac{T_{fr}}{T_{ch}}, \tag{25.4}$$

o   T_{ch} est la temp  rature de la source chaude et T_{fr} la temp  rature de la source froide. Ce rendement est appel   **rendement de Carnot**. C'est la valeur maximale du rendement d'un moteur thermique fonctionnant avec ces sources.

Un moteur ditherme fonctionnant sur un cycle comportant au moins une transformation irr  versible a un rendement plus faible que le rendement de Carnot.

Pour avoir le rendement de Carnot le plus   lev   possible il faut avoir deux sources de temp  ratures aussi   loign  es que possible.

c) Exemples

Une centrale   lectrique nucl  aire peut   tre mod  lis  e par une machine thermique fournissant du travail   lectrique et travaillant avec, comme source chaude, le r  acteur et, comme source froide, l'eau d'une rivi  re. Pour les valeurs typiques $T_{ch} = 600$ K et $T_{fr} = 300$ K le rendement de Carnot est   gal    0,5. En pratique le rendement est compris entre 30 et 40%. Il est plus faible que le rendement de Carnot en raison des irr  versibilit  s et de diverses pertes.

Dans le cas d'un moteur de voiture, la source froide est l'air atmosph  rique, de temp  rature typique $T_{fr} = 300$ K. Le transfert thermique est apport   au moteur par les gaz en combustion dont la temp  rature peut valoir $T_{ch} = 3000$ K. Pour ces valeurs le rendement de Carnot vaut 0,9. En pratique une valeur typique de rendement est 35% pour un moteur    essence et 45% pour un moteur Diesel.

2.3 Machine frigorifique

a) Efficacité de la machine frigorifique

Le but d'une machine frigorifique est de produire du froid. Il s'agit donc de prendre du transfert thermique à la source froide et la grandeur intéressante est Q_{fr} . La grandeur coûteuse est le travail W fourni à la machine. On définit l'**efficacité** de la machine frigorifique par :

$$e_{\text{frigo}} = \left| \frac{Q_{fr}}{W} \right| = \frac{Q_{fr}}{W}, \quad (25.5)$$

puisque $Q_{fr} > 0$ et $W > 0$.

D'après le premier principe, $W = -(Q_{fr} + Q_{ch})$ donc :

$$e_{\text{frigo}} = -\frac{Q_{fr}}{Q_{fr} + Q_{ch}} = -\frac{1}{1 + \frac{Q_{ch}}{Q_{fr}}}.$$

En multipliant l'équation du second principe par $\frac{T_{ch}}{Q_{fr}}$ on trouve : $\frac{Q_{ch}}{Q_{fr}} + \frac{T_{ch}}{T_{fr}} + \frac{T_{ch}S_{\text{créée}}}{Q_{fr}} = 0$.

On en déduit :

$$e_{\text{frigo}} = \frac{1}{\frac{T_{ch}}{T_{fr}} - 1 + \frac{T_{ch}S_{\text{créée}}}{Q_{fr}}}. \quad (25.6)$$

D'après le deuxième principe $S_{\text{créée}} \geq 0$ et pour une machine frigorifique $Q_{fr} > 0$ (la source froide donne du transfert thermique), donc $\frac{T_{ch}S_{\text{créée}}}{Q_{fr}} \geq 0$. Par suite :

$$e_{\text{frigo}} \leq \frac{1}{\frac{T_{ch}}{T_{fr}} - 1} = \frac{T_{fr}}{T_{ch} - T_{fr}}.$$

L'égalité est réalisée si et seulement si $S_{\text{créée}} = 0$, c'est-à-dire si le cycle est réversible. Ce résultat est le **théorème de Carnot** :

L'efficacité d'une machine frigorifique réversible est :

$$e_{\text{frigo,rev}} = \frac{T_{fr}}{T_{ch} - T_{fr}}, \quad (25.7)$$

où T_{ch} est la température de la source chaude et T_{fr} la température de la source froide. C'est la valeur maximale de l'efficacité d'une machine frigorifique fonctionnant avec ces sources.

Une machine frigorifique fonctionnant sur un cycle non réversible a une efficacité plus faible que $e_{\text{frigo,rev}}$.

D'après la formule (25.7), l'efficacité est plus grande quand les températures des deux sources sont plus proches.

CHAPITRE 25 – MACHINES THERMIQUES

b) Exemples

Un congélateur domestique est modélisable par une machine thermique ditherme avec pour source froide l'intérieur du congélateur, à la température $T_{fr} = -18^\circ\text{C} = 255\text{ K}$, et pour source chaude l'air de la pièce de température $T_{ch} = 300\text{ K}$. Pour ces températures, $e_{\text{frigo,rev}} = 5,7$. Dans la pratique le coefficient d'efficacité est au mieux voisin de 2.

L'efficacité d'un réfrigérateur domestique peut valoir jusqu'à 8. Cette valeur meilleure, s'explique par le fait que la température de la source froide (intérieur du réfrigérateur) est plus proche de la température de la source chaude (air de la pièce) que dans le cas du congélateur.

2.4 Pompe à chaleur

a) Efficacité d'une pompe à chaleur

Le but d'une machine pompe à chaleur est de chauffer la source chaude. La grandeur intéressante est donc Q_{ch} . La grandeur coûteuse est le travail W fourni à la machine. On définit donc l'**efficacité** de la pompe à chaleur :

$$e_{\text{pac}} = \left| \frac{Q_{ch}}{W} \right| = -\frac{Q_{ch}}{W}, \quad (25.8)$$

puisque $Q_{ch} < 0$ et $W > 0$.

D'après le premier principe, $W = -(Q_{fr} + Q_{ch})$ donc :

$$e_{\text{pac}} = \frac{Q_{ch}}{Q_{fr} + Q_{ch}} = \frac{1}{1 + \frac{Q_{fr}}{Q_{ch}}}.$$

En multipliant l'équation du second principe par $\frac{T_{fr}}{Q_{ch}}$ on trouve : $\frac{T_{fr}}{T_{ch}} + \frac{Q_{fr}}{Q_{ch}} + \frac{T_{fr}S_{\text{créée}}}{Q_{ch}} = 0$.

On en déduit :

$$e_{\text{pac}} = \frac{1}{1 - \frac{T_{fr}}{T_{ch}} - \frac{T_{fr}S_{\text{créée}}}{Q_{ch}}}. \quad (25.9)$$

D'après le deuxième principe $S_{\text{créée}} \geq 0$ et pour une pompe à chaleur $Q_{ch} < 0$ (la source chaude reçoit du transfert thermique), donc $\frac{T_{fr}S_{\text{créée}}}{Q_{ch}} \leq 0$. Par suite :

$$e_{\text{pac}} \leq \frac{1}{1 - \frac{T_{fr}}{T_{ch}}} = \frac{T_{ch}}{T_{ch} - T_{fr}}.$$

L'égalité est réalisée si et seulement si $S_{\text{créée}} = 0$, c'est-à-dire si le cycle est réversible. Ce résultat est le **théorème de Carnot** :

L'efficacité d'une machine pompe à chaleur réversible est :

$$e_{\text{pac,rev}} = \frac{T_{ch}}{T_{ch} - T_{fr}}, \quad (25.10)$$

où T_{ch} est la température de la source chaude et T_{fr} la température de la source froide. C'est la valeur maximale de l'efficacité d'une pompe à chaleur fonctionnant avec ces sources.

Une pompe à chaleur fonctionnant sur un cycle irréversible a une efficacité plus faible que $e_{\text{pac,rev}}$.

D'après la formule (25.10), l'efficacité est plus grande quand les températures des deux sources sont plus proches. L'efficacité est aussi supérieure à 1 (alors que le rendement d'un moteur est inférieur à 1). C'est pour cela que la pompe à chaleur est intéressante : en utilisant le travail W on peut fournir à la source chaude un transfert thermique égal à $e_{\text{pac}}W$, plus grand que W . La différence est l'énergie fournie par la source froide.

b) Exemple

Une pompe à chaleur, utilisée pour chauffer une maison en hiver, travaille avec l'eau du circuit de chauffage pour source chaude et l'air à l'extérieur de la maison pour source froide. Ainsi, on chauffe la maison en refroidissant le jardin. L'efficacité e_{pac} diminue avec l'écart des températures des deux sources. C'est pourquoi il est préférable d'avoir un chauffage par le sol (eau à $T_{ch} = 35^\circ\text{C}$) plutôt que par radiateur (eau à $T_{ch} = 60^\circ\text{C}$). On trouve dans la notice d'une pompe à chaleur le coefficient d'efficacité, appelé COP dans ce contexte, correspondant à $T_{ch} = 35^\circ\text{C}$ et $T_{fr} = 7^\circ\text{C}$. Le COP varie entre 3 et 5. La pompe à chaleur est de classe A selon les normes européennes si son COP est supérieur à 3,65. La valeur maximale théorique correspondant aux températures précédentes est : $e_{\text{pac,rev}} = \frac{273 + 35}{35 - 7} = 11$. La différence entre $e_{\text{pac,rev}}$ et le COP réel est due au fait que la machine réelle n'est pas réversible, mais aussi à la consommation d'énergie pour des tâches annexes.

3 Étude de cycles théoriques réversibles

3.1 Cycle de Carnot pour un gaz parfait

On appelle **cycle de Carnot** un cycle ditherme réversible. Lorsque le système échange du transfert thermique avec la source chaude (respectivement froide), la condition de réversibilité thermique impose que la température du système soit égale à T_{ch} (respectivement T_{fr}). Ces échanges thermiques ne peuvent donc avoir lieu que lors de transformations isothermes. En dehors de ces transformations, le système ne doit pas échanger de transfert thermique donc il ne peut avoir que des transformations adiabatiques et réversibles. Le cycle de Carnot le plus simple est le suivant :

- une transformation isotherme à T_{ch} ,
- une transformation adiabatique et réversible (donc isentropique) dans laquelle la température passe de T_{ch} à T_{fr} ,

CHAPITRE 25 – MACHINES THERMIQUES

- une transformation isotherme à T_{fr} ,
- une transformation adiabatique et réversible (donc isentropique) dans laquelle la température passe de T_{fr} à T_{ch} , ramenant le système à l'état initial.

Dans ce paragraphe le système Σ sera un échantillon de gaz parfait de quantité de matière n . La figure 25.3 montre, dans un diagramme de Clapeyron, un cycle de Carnot $ABCD$ pour ce système :

- AB : détente isotherme à T_{ch} ,
- BC : détente adiabatique et réversible,
- CD : compression isotherme à T_{fr} ,
- DA : compression adiabatique et réversible.

Le cycle est décrit dans le sens horaire et, comme il a été dit au chapitre 22, le gaz fournit du travail dans ce cas. Il s'agit donc du cycle d'un moteur ditherme.

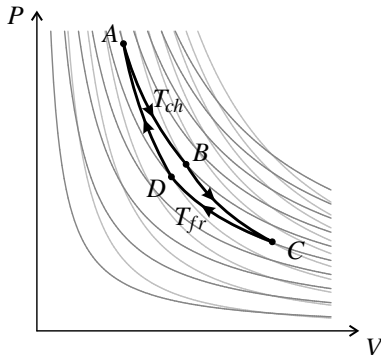


Figure 25.3 – Cycle de Carnot pour un échantillon de gaz parfait.

On va exprimer les transferts d'énergie Q_{ch} , Q_{fr} et W de Σ sur le cycle en fonction de la quantité de gaz n , et des caractéristiques du cycle :

les températures T_{ch}, T_{fr} , et le rapport des volumes $\alpha = \frac{V_B}{V_A}$.

Le transfert thermique Q_{ch} est échangé au cours de la transformation isotherme AB . D'après le premier principe $Q_{ch} = \Delta U_{AB} + \Delta E_{c,AB} - W_{AB}$. $\Delta U_{AB} = 0$ car la transformation est isotherme et l'énergie interne du gaz parfait ne dépend que de sa température (première loi de Joule); $\Delta E_c = 0$; le travail reçu par un gaz parfait dans une transformation isotherme est donné par la formule 22.7 page 807 : $W_{AB} = -nRT_{ch} \ln \left(\frac{V_B}{V_A} \right) = -nRT_{ch} \ln \alpha$. Ainsi :

$$Q_{ch} = nRT_{ch} \ln \alpha.$$

Le transfert thermique Q_{fr} est échangé au cours de la transformation isotherme CD . Le même raisonnement conduit à : $Q_{fr} = nRT_{fr} \ln \left(\frac{V_D}{V_C} \right)$. Or : $\frac{V_D}{V_C} = \frac{V_D}{V_A} \frac{V_A}{V_B} \frac{V_B}{V_C}$. Pour trouver $\frac{V_D}{V_A}$ et $\frac{V_B}{V_C}$ on peut appliquer la loi de Laplace aux transformations DA et BC qui sont adiabatiques et réversibles :

$$T_D V_D^{\gamma-1} = T_A V_A^{\gamma-1} \quad \text{d'où} \quad \frac{V_D}{V_A} = \left(\frac{T_A}{T_D} \right)^{\frac{1}{\gamma-1}} = \left(\frac{T_{ch}}{T_{fr}} \right)^{\frac{1}{\gamma-1}},$$

et de même :

$$\frac{V_B}{V_C} = \left(\frac{T_C}{T_B} \right)^{\frac{1}{\gamma-1}} = \left(\frac{T_{fr}}{T_{ch}} \right)^{\frac{1}{\gamma-1}}.$$

Ainsi : $\frac{V_D}{V_C} = \frac{1}{\alpha}$ et :

$$Q_{fr} = -nRT_{fr} \ln \alpha.$$

Au lieu de calculer directement le travail cédé par le système au cours du cycle : $W = W_{AB} + W_{BC} + W_{CD} + W_{DA}$, on peut le trouver en appliquant le premier principe au système sur le cycle :

$$W = -Q_{ch} - Q_{fr} = -nR(T_{ch} - T_{fr}) \ln \alpha.$$

Pour conclure on peut calculer le rendement de ce système en tant que moteur ditherme :

$$\rho_{\text{moteur}} = -\frac{W}{Q_{ch}} = \frac{T_{ch} - T_{fr}}{T_{ch}} = 1 - \frac{T_{fr}}{T_{ch}} = \rho_{\text{moteur,rev}}.$$

On retrouve, bien sûr, le rendement de Carnot.

Remarque

Le cycle étant réversible, il peut être aussi parcouru par le système Σ dans le sens $ADCB$. Tous les échanges énergétiques sont alors opposés à ceux que l'on vient de calculer et Σ est une machine frigorifique ou une pompe à chaleur.

3.2 Cycle de Carnot pour un système diphasé

Dans ce paragraphe le système Σ sera un échantillon de corps pur sous forme liquide et gaz. La figure 25.4 montre, dans un diagramme de Clapeyron, un cycle de Carnot $ABCD$ pour ce système :

- AB : transformation adiabatique et réversible,
- BC : vaporisation partielle isotherme réversible à T_{fr} sous la pression $P_{\text{sat}}(T_{fr})$,
- CD : transformation adiabatique et réversible,
- DA : liquéfaction totale isotherme réversible à T_{ch} sous la pression $P_{\text{sat}}(T_{ch})$.

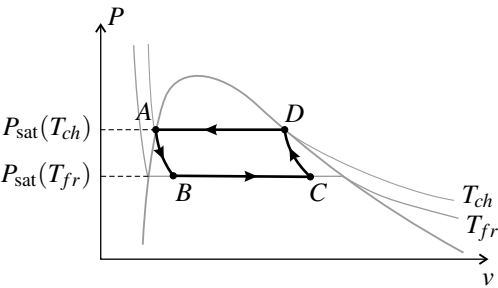


Figure 25.4 – Un cycle de Carnot pour un corps pur diphasé.

En A on a du liquide saturant seul et en D de la vapeur saturante sèche.

Le cycle est décrit dans le sens trigonométrique et, comme il a été dit au chapitre 22, dans ce cas le fluide reçoit du travail. Il s'agit donc du cycle d'une machine frigorifique ou d'une pompe à chaleur.

On va exprimer les transferts d'énergie Q_{ch} , Q_{fr} et W de Σ sur le cycle en fonction de la masse de fluide de gaz m , des températures T_{ch} , T_{fr} et de grandeurs caractéristiques du fluide (capacités thermiques massiques, enthalpie ou entropie de vaporisation).

On va commencer par déterminer les fractions massiques en gaz $x_{G,B}$ et $x_{G,C}$ aux points B et C. La transformation AB est adiabatique et réversible, donc isentropique. Ainsi :

$$S_A = S_B \quad \text{soit} \quad ms_L(T_{ch}) = m(s_L(T_{fr}) + x_{G,B}\Delta_{\text{vap}}s(T_{fr})),$$

CHAPITRE 25 – MACHINES THERMIQUES

en utilisant la formule (24.12) du chapitre 24, page 858. On en tire :

$$x_{G,B} = \frac{s_L(T_{ch}) - s_L(T_{fr})}{\Delta_{\text{vap}}s(T_{fr})} = \frac{c_L}{\Delta_{\text{vap}}s(T_{fr})} \ln \left(\frac{T_{ch}}{T_{fr}} \right),$$

en utilisant la formule (24.11), page 857, et en notant c_L la capacité thermique du liquide. De la même manière on obtient $x_{G,C}$ par l'équation :

$$S_C = S_D \quad \text{soit} \quad m(s_L(T_{fr}) + x_{G,C}\Delta_{\text{vap}}s(T_{fr})) = m(s_L(T_{ch}) + \Delta_{\text{vap}}s(T_{ch})),$$

$$\text{d'où : } x_{G,C} = \frac{c_L \ln \left(\frac{T_{ch}}{T_{fr}} \right) + \Delta_{\text{vap}}s(T_{ch})}{\Delta_{\text{vap}}s(T_{fr})}. \text{ Ainsi : } x_{G,C} - x_{G,B} = \frac{\Delta_{\text{vap}}s(T_{ch})}{\Delta_{\text{vap}}s(T_{fr})}.$$

Le fluide échange du transfert thermique avec la source froide au cours de la transformation isobare BC donc $Q_{fr} = Q_{BC} = \Delta H_{BC} = m(x_{G,C} - x_{G,B})\Delta_{\text{vap}}h(T_{fr})$, d'après la formule (23.18), page 832. Ainsi :

$$Q_{fr} = m\Delta_{\text{vap}}h(T_{fr}) \frac{\Delta_{\text{vap}}s(T_{ch})}{\Delta_{\text{vap}}s(T_{fr})} = m\Delta_{\text{vap}}h(T_{ch}) \frac{T_{fr}}{T_{ch}},$$

d'après la relation (24.15), page 859.

Le fluide échange du transfert thermique avec la source chaude au cours de la transformation isobare DA et de même $Q_{ch} = Q_{DA} = \Delta H_{DA}$ soit :

$$Q_{ch} = -m\Delta_{\text{vap}}h(T_{ch}).$$

On obtient le travail échangé par le fluide au cours du cycle en appliquant le premier principe :

$$W = -Q_{fr} - Q_{ch} = m\Delta_{\text{vap}}h(T_{ch}) \left(1 - \frac{T_{fr}}{T_{ch}} \right).$$

Pour conclure, on peut calculer le rendement de ce système en tant que machine frigorifique :

$$e_{\text{frigo}} = \frac{Q_{fr}}{W} = \frac{T_{fr}}{T_{ch} - T_{fr}}.$$

On retrouve bien le rendement de la machine réversible.

4 Étude de machines thermiques réelles

4.1 Moteur à explosion

Les moteurs à essence fonctionnent suivant un cycle théorique proposé par le physicien français Beau de Rochas en 1862. Le moteur fut réalisé par l'allemand Otto une quinzaine d'années plus tard. Ce moteur est appelé *moteur à explosion* car il est nécessaire de produire une étincelle à l'aide d'une bougie pour provoquer l'inflammation du mélange air-carburant.

On a représenté sur la figure 25.5 le cycle théorique et sur la figure 25.6 le cycle réel qui, comme on le voit, se rapproche du cycle théorique.

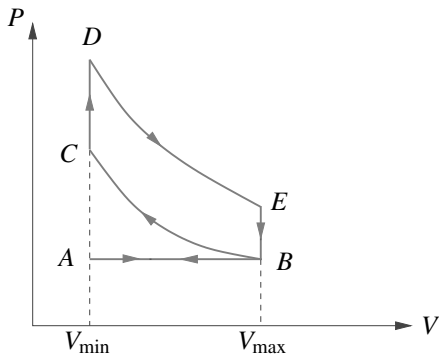


Figure 25.5 – Cycle théorique du moteur à quatre temps.

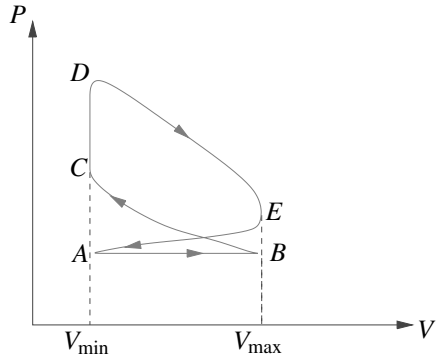


Figure 25.6 – Cycle réel du moteur à quatre temps.

En A le piston est en bout de course et le cylindre offre le volume minimal $V_{\min}$. L'évolution est la suivante :

- 1^{er} temps (AB) : admission, la soupape d'admission est ouverte, le piston descend en aspirant le mélange air-carburant jusqu'au volume $V_{\max}$.
- 2^e temps (BCD) : le piston remonte et comprime le gaz adiabatiquement jusqu'en C puis l'étincelle est produite par la bougie provoquant la combustion. La pression augmente très rapidement mais le piston n'a pas le temps de bouger (évolution isochore CD).
- 3^e temps (DE) : les gaz brûlés sous forte pression repoussent le piston. C'est une détente adiabatique avec production de travail.
- 4^e temps (EBA) : la soupape d'échappement s'ouvre, provoquant une rapide baisse de pression isochore (EB), puis le piston remonte pour refouler les gaz brûlés (BA).

Les mouvements du piston sont représentés sur la figure 25.7.

On s'aperçoit donc que, dans un moteur à quatre temps, le piston fait deux allers et retours pour décrire un cycle. Dans un moteur de voiture, les différents cylindres fonctionnent avec un décalage de manière à ce que le piston remonte dans certains cylindres quand il descend dans d'autres. Cela est dû au fait que les bielles des différents cylindres ne sont pas fixées toutes du même côté du vilebrequin.

Pour calculer le rendement du moteur on suppose le gaz parfait, de rapport des capacités thermiques γ indépendant de la température. On suppose de plus les transformations BC et DE adiabatiques et réversibles pour pouvoir appliquer la loi de Laplace. On note n la quantité de gaz contenue dans le système.

Le transfert thermique est échangé avec la source chaude lors de la transformation isochore CD donc :

$$Q_{ch} = Q_{CD} = \Delta U_{CD} = \frac{nR}{\gamma - 1} (T_D - T_C).$$

Le transfert thermique est échangé avec la source froide lors de la transformation isochore EB donc :

$$Q_{fr} = Q_{EB} = \Delta U_{EB} = \frac{nR}{\gamma - 1} (T_B - T_E).$$

CHAPITRE 25 – MACHINES THERMIQUES

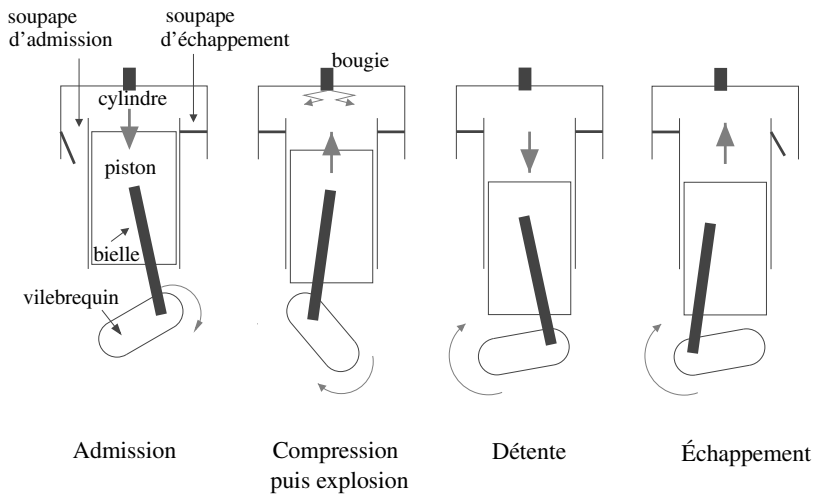


Figure 25.7 – Fonctionnement d'un moteur à 4 temps.

D'autre part, d'après la loi de Laplace :

$$T_C = T_B \left(\frac{V_{\max}}{V_{\min}} \right)^{\gamma-1} \quad \text{et} \quad T_D = T_E \left(\frac{V_{\max}}{V_{\min}} \right)^{\gamma-1}.$$

On en déduit que : $Q_{ch} = -Q_{fr} \left(\frac{V_{\max}}{V_{\min}} \right)^{\gamma-1}$, puis en appliquant le premier principe au système sur le cycle :

$$W = -Q_{ch} - Q_{fr} = - \left(1 - \left(\frac{V_{\max}}{V_{\min}} \right)^{1-\gamma} \right) Q_{ch}.$$

Ainsi le rendement du moteur est :

$$\rho_{\text{moteur}} = -\frac{W}{Q_{ch}} = 1 - \left(\frac{V_{\max}}{V_{\min}} \right)^{1-\gamma}. \quad (25.11)$$

Il dépend du rapport $\frac{V_{\max}}{V_{\min}}$ qui est appelé *taux de compression*. Comme $1 - \gamma < 0$, l'expression précédente montre que ρ_{moteur} est d'autant plus grand que le taux de compression est important. Une valeur type du taux de compression est 10 et le rendement donné par la formule (25.11) (avec $\gamma = 1,4$) est 0,60. Les carburants sont conçus de manière à supporter un fort taux de compression sans exploser avant l'étincelle de la bougie.

4.2 Machine frigorifique (MPSI)

a) Principe de fonctionnement

Les systèmes frigorifiques et les pompes à chaleur sont en général des systèmes à condensation dont le principe est représenté sur la figure 25.8. Un fluide, dit *frigorigène* ou *caloporteur* suivant l'utilisation, suit un circuit comportant :

- un compresseur C dans lequel il reçoit du travail et n'a pas d'échange thermique (dans le compresseur la température du fluide augmente),
- un condenseur dans lequel il est en contact avec la source chaude à laquelle il cède du transfert thermique,
- un détendeur D dans lequel il ne reçoit ni travail, ni transfert thermique (dans le détendeur la température du fluide diminue),
- un évaporateur dans lequel il est en contact avec la source froide de laquelle il reçoit du transfert thermique.

Les transformations du fluide dans le condenseur et l'évaporateur sont isothermes, l'énergie perdue ou gagnée par le fluide correspondant à un changement d'état.

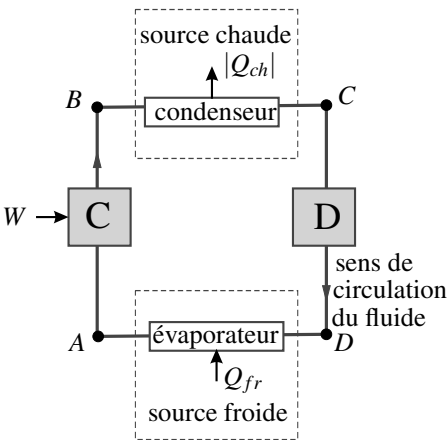


Figure 25.8 – Schéma d'un système à condensation.

b) Premier principe pour un fluide en écoulement

Le système thermodynamique fermé auquel on peut appliquer les équations du paragraphe 2 est le système constitué par la totalité du fluide contenu dans le circuit.

Pour obtenir une relation faisant apparaître les échanges énergétiques dans l'un des éléments du circuit (compresseur, condenseur, détendeur ou évaporateur), il faut appliquer le premier principe pour un fluide en écoulement dont la démonstration est donnée ci-après.

On considère, de manière générale, un fluide en écoulement lent, passant dans un élément actif à l'intérieur duquel il peut échanger du travail et/ou du transfert thermique. Entre l'entrée et la sortie de cet élément, les grandeurs thermodynamiques massiques du fluide, enthalpie massique h , énergie interne massique u , volume massique v changent. On note ces grandeurs

CHAPITRE 25 – MACHINES THERMIQUES

à l'entrée avec un e en indice, et à la sortie avec un s en indice. On note aussi P_e et P_s les pressions à l'entrée et à la sortie.

Soit w et q le travail et le transfert thermique reçus par l'unité de masse de fluide qui traverse l'élément actif. Le travail w est échangé par le fluide avec des pièces mobiles, à l'intérieur de l'élément actif.

On considère un système Σ **fermé** représenté sur la figure 25.9. Dans l'état initial, Σ contient une masse m de fluide située devant l'entrée de l'élément actif ainsi que le fluide qui remplit l'élément actif. Dans l'état final, Σ contient la même masse m de fluide à la sortie de l'élément actif et le fluide qui remplit l'élément actif.

On suppose l'**écoulement permanent** : l'état du fluide en un point donné de la canalisation est le même à chaque instant (même si, à deux instant différents, ce n'est pas le même fluide puisqu'il s'écoule). Ainsi, dans Σ à l'état final, le fluide qui est à l'intérieur de l'élément actif a exactement les mêmes propriétés que celui qui se trouve au même endroit, dans Σ à l'état initial.

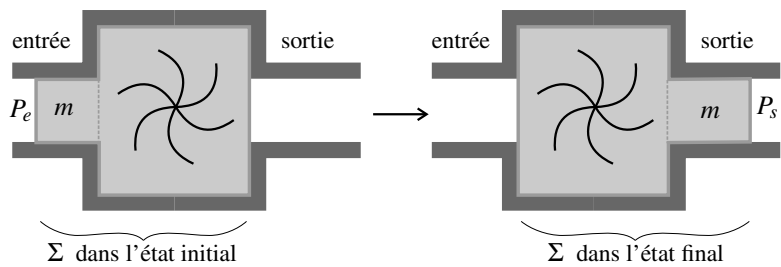


Figure 25.9 – Démontstration du premier principe pour un fluide en écoulement stationnaire.

Quelle est la variation d'énergie interne de Σ entre l'état initial et l'état final ? La différence provient de la masse m de fluide qui, dans l'état initial, a une énergie interne massique u_e et, dans l'état final, a une énergie interne massique u_s , donc :

$$\Delta U = mu_s - mu_e.$$

Pour simplifier, on fait l'hypothèse que le fluide s'écoule lentement et que la variation d'énergie cinétique est négligeable devant la variation d'énergie interne précédente. On fait donc l'approximation :

$$\Delta E_c \simeq 0.$$

Au cours de sa transformation le système Σ reçoit un travail de la part des forces de pression, qui le poussent à l'entrée et le repoussent à la sortie. Ce travail a été calculé au chapitre 22. On adapte la formule (22.3), page 803 : le volume balayé à l'entrée est mv_e , volume occupé par la masse m à l'entrée, et le volume balayé à la sortie est mv_s (voir figure 25.9). Ainsi :

$$W_{\text{pression}} = P_e(mv_e) - P_s(mv_s).$$

Ce travail n'est pas le seul reçu par Σ et ce n'est pas non plus le plus intéressant, car il s'agit d'un travail de forces internes au fluide. Σ reçoit dans l'élément actif un travail appelé **travail utile** donné par :

$$W_u = mw_u.$$

Il reçoit aussi un transfert thermique :

$$Q = mq.$$

Ainsi, le premier principe pour Σ , entre l'état initial et l'état final considérés, s'écrit :

$$\Delta U + \Delta E_c = W_{\text{pression}} + W_u + Q \quad \text{soit} \quad mu_s - mu_e = P_e(mu_e) - P_s(mu_s) + mw_u + mq,$$

soit encore, en simplifiant par m et en regroupant les termes :

$$(u_s + P_s v_s) - (u_e + P_e v_e) = w_u + q.$$

Il apparaît dans cette formule la variation d'enthalpie massique du fluide entre l'entrée et la sortie :

$$\Delta h = h_s - h_e = (u_s + P v_s) - (u_e + P v_e).$$

Pour un fluide en **écoulement stationnaire**, traversant un élément actif à l'intérieur duquel il reçoit, de parties mobiles, un travail massique w_u , et dans lequel il reçoit le transfert thermique massique q , le **premier principe** s'écrit, en négligeant la variation d'énergie cinétique :

$$\Delta h = w_u + q, \tag{25.12}$$

où Δh est la variation d'enthalpie massique entre l'entrée et la sortie de l'élément actif.



L'intérêt de cette formulation du premier principe est qu'elle ne fait pas intervenir le travail des forces de pression, travail interne au fluide, mais uniquement le travail utile, travail échangé par le fluide avec les parties mobiles de l'élément actif.

Ce résultat général s'applique à chacun des quatre éléments de la machine :

- pour le compresseur, dans lequel le fluide reçoit des pièces mobiles le travail massique w_{comp} et ne reçoit aucun transfert thermique :

$$\Delta h_{AB} = h_B - h_A = w_{\text{comp}},$$

- pour le condenseur, dans lequel il n'y a pas de pièces mobiles et dans lequel le fluide reçoit le transfert thermique massique $q_{ch} < 0$ de la source chaude (donc lui cède le transfert thermique $-q_{ch} > 0$) :

$$\Delta h_{BC} = h_C - h_B = q_{ch},$$

- pour le détenteur dans lequel il n'y a pas de pièces mobiles et dans lequel le fluide ne reçoit aucun transfert thermique :

$$\Delta h_{CD} = h_D - h_C = 0 \quad \text{soit} \quad h_C = h_D,$$

CHAPITRE 25 – MACHINES THERMIQUES

- pour l'évaporateur dans lequel il n'y a pas de pièces mobiles et dans lequel le fluide reçoit le transfert thermique massique q_{fr} de la source froide ;

$$\Delta h_{DA} = h_A - h_D = q_{fr}.$$

L'efficacité de la machine est, suivant sa fonction :

$$e_{frigo} = \frac{q_{fr}}{w_{comp}} = \frac{h_A - h_D}{h_B - h_A} \quad \text{ou} \quad e_{pac} = -\frac{q_{ch}}{w_{comp}} = \frac{h_B - h_C}{h_B - h_A}. \quad (25.13)$$

Remarque

Pour le système Σ comprenant la totalité du fluide contenu dans le circuit, de masse m_Σ , les échanges énergétiques sur un cycle (c'est-à-dire le circuit complet) sont :

$$W = m_\Sigma w_{comp}, \quad Q_{ch} = m_\Sigma q_{ch} \quad \text{et} \quad Q_{fr} = m_\Sigma q_{fr}.$$

Les expressions des efficacités ci-dessus sont bien compatibles avec (25.6) et (25.8).

c) Diagramme des frigoristes

Pour étudier ces machines on utilise habituellement un diagramme appelé **diagramme des frigoristes**. Dans ce diagramme, on porte la pression P en ordonnée et l'enthalpie massique h en abscisse pour le fluide utilisé. C'est, dans le principe, un diagramme (P, h) . Cependant, pour couvrir une plus large gamme de pressions, l'échelle des pressions est logarithmique : c'est, dans la pratique, un diagramme $(\log P, h)$.

Ce diagramme, comme le diagramme Clapeyron, comporte une zone d'équilibre liquide-vapeur qui est délimitée par la courbe de saturation, à droite il y a la zone du gaz et à gauche la zone du liquide (voir figure 25.10). Le sommet de la courbe de saturation est le point critique.

Sur le diagramme $(\log P, h)$ on trace des réseaux de courbes sur lesquelles les différentes grandeurs thermodynamiques intensives du fluide sont constantes :

- isothermes,
- isobares,
- isochores (volume massique v constant),
- isenthalpes (enthalpie massique h constante),
- isentropes (entropie massique s constante),
- isotitres (titre en vapeur x_V constant).

Les **isobares** sont des droites horizontales, les **isenthalpes** sont des droites verticales (voir figure 25.10 par exemple). Noter que les isobares ne sont pas équidistantes à cause de l'échelle logarithmique.

Les **isothermes** ont une forme plus compliquée (voir figure 25.11) :

- Dans la zone d'équilibre liquide-vapeur, si la température T est constante la pression l'est aussi, puisque $P = P_{L-G}(T)$, donc chaque isotherme a un palier horizontal. La largeur de ce palier est égale à l'enthalpie massique de vaporisation $\Delta_{vap}h(T)$ du fluide. Pour ne pas surcharger le diagramme, il est d'usage de ne dessiner que les extrémités des paliers sur la courbe de saturation.

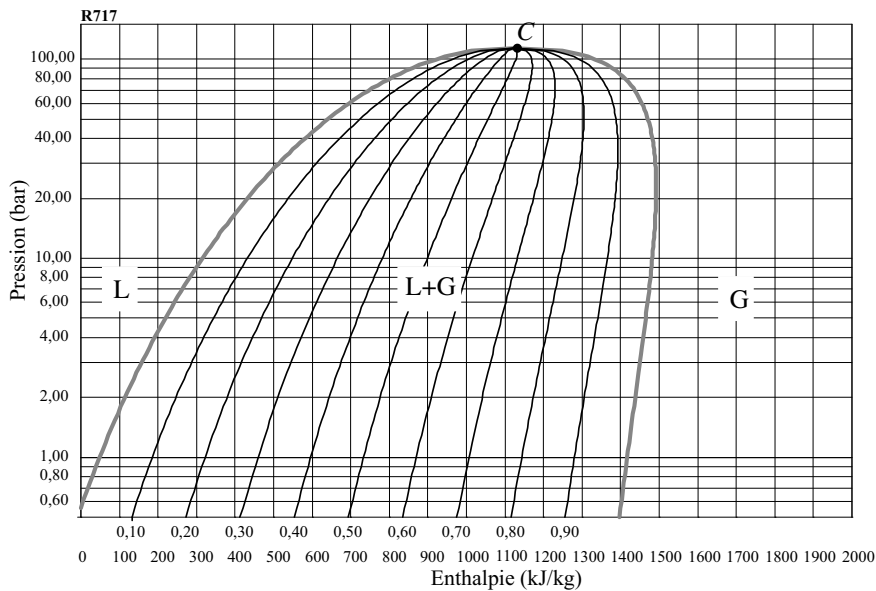


Figure 25.10 – Diagramme des frigorisites pour l'ammoniac (fluide frigorigère R717) : isobares (horizontales), isenthalpes (verticales), courbe de saturation (en gris) et réseau d'isotitres (en noir).

- Dans la zone du liquide les isothermes sont quasiment des droites verticales : pour une phase condensée l'enthalpie ne dépend pratiquement que de T , donc si T est constante, h est constante aussi. Pour la même raison de clarté on ne représente que le départ de cette droite verticale sur la courbe d'ébullition.
- Dans la zone de la vapeur, les isothermes sont courbées. Toutefois, aux basses pressions, elles ressemblent à des droites verticales car, pour les basses pressions, la vapeur est assimilable à un gaz parfait dont l'enthalpie ne dépend que de la température.

Les courbes **isotitres** n'existent que dans la zone d'équilibre liquide-vapeur. Elles partent du point critique, au sommet de la courbe de saturation, et vont jusqu'à l'axe des abscisses (voir figure 25.10).

Les **isentropes** n'ont pas de rupture de pente à la frontière du domaine d'équilibre liquide-vapeur. Dans la zone du liquide, ce sont pratiquement des droites verticales (voir figure 25.12) parce que, dans le modèle du liquide incompressible et indilatable, l'entropie ne dépend que de la température, donc si s est constante, T l'est aussi et donc h l'est aussi. Cette portion verticale n'est pas toujours représentée.

Les **isochores** sont très peu utilisées en pratique.

Sur chaque courbe on peut lire la valeur de la grandeur thermodynamique qui lui est associée. Si l'on connaît deux des grandeurs thermodynamiques intensives ci-dessus on place facilement le point représentant l'état du fluide sur le diagramme (il faut le plus souvent interpoler entre deux courbes) et on peut lire les valeurs des autres grandeurs. Le diagramme est donc une véritable table de données thermodynamiques concernant le fluide.

CHAPITRE 25 – MACHINES THERMIQUES

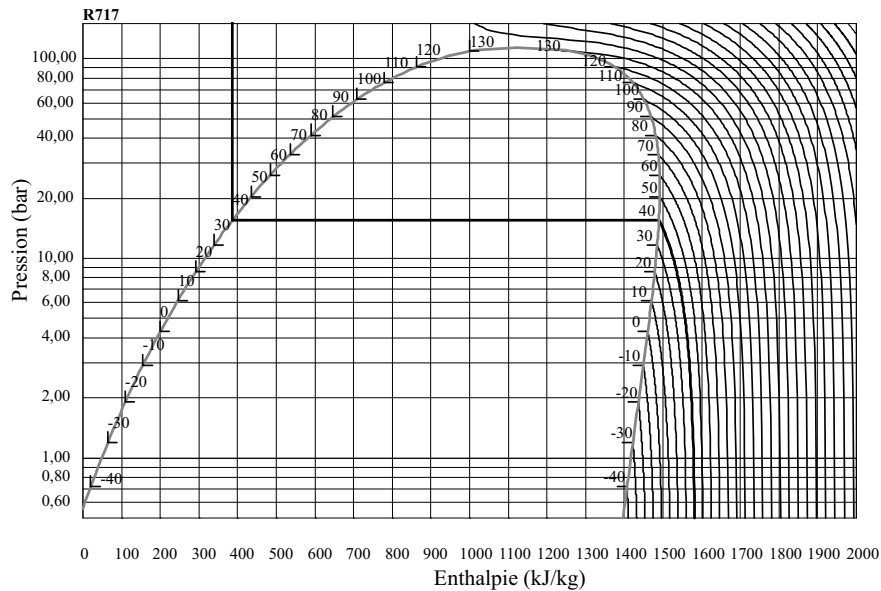


Figure 25.11 – Diagramme des frigorisites pour l'ammoniac : réseau d'isothermes. Les parties horizontales ou verticales des isothermes ne sont pas représentées par le logiciel. On a complété l'isotherme $T = 40^{\circ}\text{C}$.

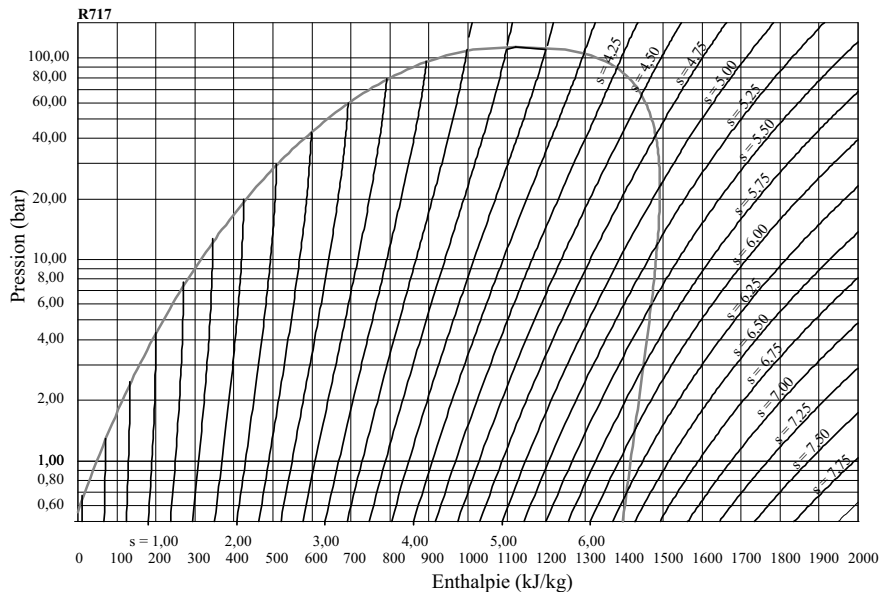


Figure 25.12 – Diagramme des frigorisites pour l'ammoniac : réseau d'isentropes. La valeur indiquée sur chaque isentrope est en $\text{kJ}\cdot\text{K}^{-1}\cdot\text{kg}^{-1}$. Les parties horizontales ou verticales des isentropes ne sont pas représentées.

d) Étude du cycle dans le diagramme (P, h)

Le diagramme est un outil puissant pour calculer les performance du cycle. On place sur le diagramme les points A, B, C et D représentant les états successifs du fluide, puis on lit leur abscisses h_A, h_B, h_C et h_D pour calculer le rendement.

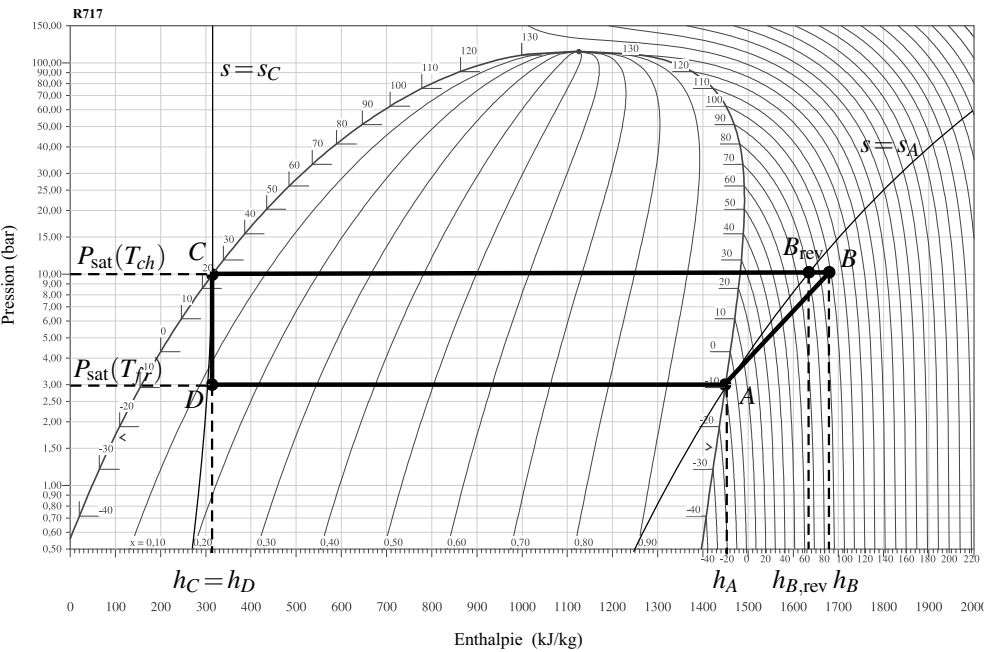


Figure 25.13 – Cycle d'une machine frigorifique : $T_{ch} = 298\text{ K}$, $T_{fr} = 263\text{ K}$.

Le cycle représenté sur la figure 25.13 est le cycle d'une machine réelle destinée à produire du froid. En A on a de la vapeur sèche et en C du liquide juste saturant. On lit sur le diagramme, $h_A = 1452 \pm 2\text{ J}\cdot\text{kg}^{-1}$, $h_B = 1680 \pm 2\text{ J}\cdot\text{kg}^{-1}$, $h_C = h_D = 313 \pm 2\text{ J}\cdot\text{kg}^{-1}$. On en déduit l'efficacité :

$$e_{\text{frigo}} = \frac{1452 - 313}{1680 - 1452} = 5,00 \pm 0,09.$$

L'efficacité d'une machine frigorifique réversible travaillant entre les mêmes températures serait :

$$e_{\text{frigo,rev}} = \frac{T_{fr}}{T_{ch} - T_{fr}} = \frac{263}{298 - 263} = 7,51.$$

L'efficacité réelle est plus petite, signe que le cycle n'est pas réversible. Ceci se voit d'ailleurs sur le diagramme : la compression AB est adiabatique, si elle était réversible B serait sur l'isentrope passant par A , ce qui n'est pas le cas. Sur la figure 25.13, B est à droite de l'isentrope $s = s_A$ ce qui montre que $s_B > s_A$. Si cette transformation était réversible elle aboutirait

CHAPITRE 25 – MACHINES THERMIQUES

au point B_{rev} avec $h_{B,\text{rev}} = 1635 \pm 2 \text{ J}\cdot\text{kg}^{-1}$ et l’efficacité de la machine serait :

$$e'_{\text{frigo}} = \frac{1452 - 313}{1635 - 1452} = 6,55 \pm 0,15,$$

valeur supérieure à l’efficacité réelle parce les irréversibilités font baisser l’efficacité. Cependant e'_{frigo} est encore inférieure à $e_{\text{frigo,rev}}$, ce qui provient du fait que la détente isenthalpique CD est, elle aussi, irréversible. En effet c’est une détente adiabatique et $s_D > s_C$, puisque D est à droite de l’isentrope $s = s_C$.

L’irréversibilité de la compression AB est une irréversibilité mécanique. L’irréversibilité de la détente CD est due aux frottements internes au fluide.

SYNTHÈSE

SAVOIRS

- sens des échanges d’énergie d’une machine thermique ditherme suivant son mode de fonctionnement
- définition du rendement ou de l’efficacité d’une machine thermique ditherme selon son rôle
- le rendement est maximum pour une machine réversible
- expression du rendement ou de l’efficacité d’une machine ditherme réversible dans les trois cas
- ordre de grandeur du rendement ou de l’efficacité d’une machine réelle
- premier principe dans un écoulement stationnaire

SAVOIR-FAIRE

- écrire les deux principes pour une machine ditherme sur un cycle
- établir l’expression du rendement dans le cas réversible
- analyser un dispositif concret et le modéliser par une machine cyclique ditherme.
- utiliser le premier principe dans un écoulement stationnaire
- utiliser un diagramme ($\log P, h$)

MOTS-CLÉS

- | | | |
|------------------------|-------------------------|-----------------------------|
| • moteur thermique | • cycle ditherme | • diagramme des frigoristes |
| • machine frigorifique | • rendement, efficacité | |
| • pompe à chaleur | • rendement de Carnot | |

S'ENTRAÎNER

25.1 Moteur réel (★)

Un moteur réel fonctionnant entre deux sources de chaleur, l'une à $T_{fr} = 400 \text{ K}$, l'autre à $T_{ch} = 650 \text{ K}$, produit 500 J par cycle, pour 1500 J de transfert thermique fourni.

1. Comparer son rendement à celui d'une machine de Carnot fonctionnant entre les deux mêmes sources.
2. Calculer l'entropie créée par cycle, notée $S_{crée}$.
3. Montrer que la différence entre le travail fourni par la machine de Carnot et la machine réelle est égale à $T_{fr}S_{crée}$, pour une dépense identique.

25.2 Perte de performance d'un congélateur (★)

Un congélateur neuf a un coefficient d'efficacité $e = 2,0$. Un appareil dans lequel on a laissé s'accumuler une couche de glace a une efficacité réduite. On suppose que l'effet de la couche de glace est de multiplier par 2 l'entropie créée pour un même transfert thermique pris à la source froide. L'intérieur du congélateur est à -20°C et la pièce dans laquelle il se trouve à 19°C .

1. Calculer numériquement α , rapport entre l'efficacité du congélateur neuf et l'efficacité d'une machine réversible fonctionnant avec les mêmes sources.
2. Montrer que ce rapport devient, pour le réfrigérateur usagé : $\alpha' = \frac{\alpha}{2 - \alpha}$. Calculer α' et l'efficacité réduite e' .

25.3 Pompe à chaleur avec source de température variable (★)

Une pompe à chaleur fonctionne avec pour source froide l'eau d'un lac à température constante $T_{fr} = 273,15 \text{ K} + \theta_{fr}$ avec $\theta_{fr} = 5^\circ\text{C}$. La pompe est utilisée pour faire passer une maison de capacité thermique C de la température $T_1 > T_{fr}$ à la température $T_2 > T_1$. On suppose la maison parfaitement isolée thermiquement et on note T sa température.

1. Le temps caractéristique de variation de T est très grand devant la durée d'un cycle de la pompe à chaleur. Dans ce cas, les principes de la thermodynamique, pour la machine sur un cycle, s'écrivent comme si T était constante. On appelle W_{cycle} le travail consommé sur un cycle et ΔT la petite variation de T entre le début et la fin du cycle. Montrer que : $W_{\text{cycle}} \geq C \Delta T f(T, T_{fr})$, où $f(T, T_{fr})$ est une fonction que l'on précisera.
2. On admet que le travail consommé par la pompe à chaleur pour faire passer T de T_1 à $T_2 > T_1$, vérifie $W \geq W_{\min} = C \int_{T_1}^{T_2} f(T, T_{fr}) dT$. Exprimer $W_{\min}$. Cette valeur minimale peut-elle être atteinte ?
3. Définir l'efficacité de la pompe à chaleur sur l'opération de chauffage de T_1 à T_2 . Exprimer sa valeur maximale et la calculer numériquement pour $\theta_1 = 10^\circ\text{C}$ et $\theta_2 = 20^\circ\text{C}$.

APPROFONDIR

25.4 Pompe à chaleur avec un gaz parfait (★)

Une pompe à chaleur effectue le cycle de Joule inversé suivant. L'air pris dans l'état A de température T_0 et de pression P_0 est comprimé suivant une adiabatique quasistatique jusqu'au point B où il atteint la pression P_1 . L'air est ensuite refroidi à pression constante et atteint la température finale de la source chaude T_1 correspondant à l'état C . L'air est encore refroidi dans une turbine suivant une détente adiabatique quasistatique pour atteindre l'état D de pression P_0 . Il se réchauffe enfin à pression constante au contact de la source froide et retrouve son état initial.

L'air est considéré comme un gaz parfait de rapport des capacités thermiques $\gamma = 1,4$ indépendant de la température. On pose $\beta = 1 - \frac{1}{\gamma}$ et $a = \frac{P_1}{P_0}$.

On prendra $T_0 = 283 \text{ K}$, $T_1 = 298 \text{ K}$, $a = 5$ et $R = 8,31 \text{ J.K}^{-1}.\text{mol}^{-1}$.

1. Représenter le cycle parcouru par le gaz dans un diagramme (P, V) .
2. Rappeler les conditions nécessaires pour assurer la validité des formules de Laplace. Donner la formule de Laplace relative à la pression et à la température.
3. En déduire l'expression des températures T_B et T_D des états B et D en fonction de T_0 , T_1 , a et β . Préciser leurs valeurs numériques.
4. Exprimer l'efficacité e de la pompe à chaleur en fonction des transferts thermiques.
5. En déduire l'expression de e en fonction de a et β . Donner sa valeur numérique.
6. Quelles doivent être les transformations du gaz si on fait fonctionner la pompe à chaleur suivant un cycle de Carnot réversible entre les températures T_0 et T_1 ?
7. Établir l'expression de son efficacité e_r . Donner sa valeur numérique.
8. Comparer e et e_r . Proposer une explication à ce résultat.
9. Déterminer l'expression de l'entropie créée S_c pour une mole d'air au cours du cycle de Joule en fonction de R , β et $x = a^\beta \frac{T_0}{T_1}$.
10. Étudier le signe de S_c en fonction de x . Était-ce prévisible ?
11. Calculer sa valeur ici.
12. Sachant qu'en régime permanent, les fuites thermiques s'élèvent à $Q_f = 20 \text{ kW}$, calculer la puissance du couple compresseur - turbine qui permet de maintenir la température de la maison constante.

25.5 Optimisation du chauffage d'un local (★)

On souhaite maintenir la température d'un local à $\theta_1 = 20^\circ\text{C}$ alors que la température extérieure est $\theta_2 = -2^\circ\text{C}$. L'énergie thermique nécessaire est de 32 MJ par heure. Par un système de chauffage central, on brûle a litres de fuel par jour (a est de l'ordre de 25 litres, sa valeur exacte n'intervient pas dans l'exercice).

1. On utilise une pompe à chaleur fonctionnant réversiblement entre le local et l'extérieur, le fuel servant à faire fonctionner le moteur de cette pompe, comme nous le décrit la question 2.

Calculer l'efficacité de la pompe à chaleur. En déduire la puissance consommée par le moteur de cette pompe.

2. Deux conseillers proposent chacun un dispositif qu'ils déclarent thermodynamiquement plus avantageux que le chauffage central :

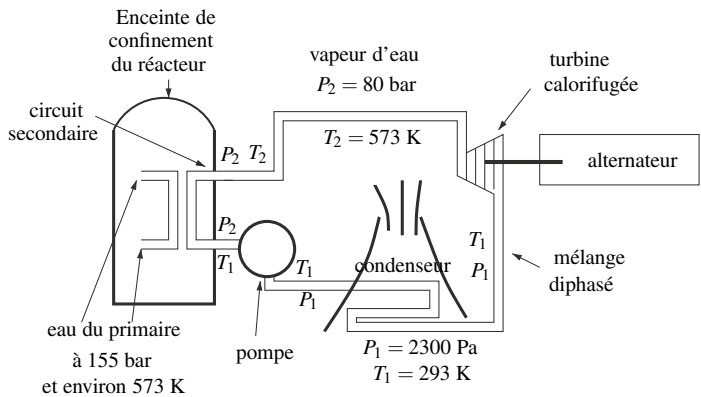
- *dispositif du conseiller n° 1* : les a litres de fuel sont brûlés, l'énergie thermique Q récupérée permet d'assurer la vaporisation de l'eau d'une chaudière auxiliaire à la température $\theta_3 = 210^\circ\text{C}$ qui sert de source chaude à un moteur ditherme réversible dont la source froide est le local, le travail fourni servant à faire fonctionner la pompe à chaleur étudiée à la question 1 ;
- *dispositif du conseiller n° 2* : le principe est le même mais la chaudière auxiliaire est à la température $\theta_4 = 260^\circ\text{C}$ et le moteur fonctionne entre cette chaudière et l'air extérieur.

Dans les deux cas, on suppose que toute l'énergie thermique Q obtenue par combustion du fuel est fournie par la chaudière auxiliaire au fluide du moteur.

- a. Dans les deux cas, calculer la durée Δt (en jours) pendant lequel le chauffage sera assuré avec les a litres de fuel.
- b. Lequel de ces deux systèmes est le plus économique ? Le système le plus économique tire-t-il son avantage de la température à laquelle fonctionne la chaudière auxiliaire ou cet avantage se maintient-il pour $\theta_3 = \theta_4$? Expliquer.
- c. La durée de chauffage, pour une quantité de fuel donnée, augmentant avec la température de la chaudière auxiliaire, un troisième conseiller prétend qu'il pourra, en utilisant le même dispositif, augmenter la durée de chauffage. A-t-il tort ou raison ? Pourquoi ?

25.6 Étude de la turbine d'un réacteur à eau pressurisée (REP) (★)

Le parc de production nucléaire français est composé de centrales de la filière REP :



On étudie l'eau circulant en circuit fermé dans le circuit secondaire. On propose de modéliser son évolution au prix de quelques approximations par le cycle suivant :

- État A : l'eau qui sort du condenseur est liquide sous la pression P_1 et à la température T_1 .
- Évolution AB : elle subit dans la pompe une compression durant laquelle sa température ne varie pratiquement pas. On négligera les échanges thermiques lors de cette compression qui l'amène dans l'état B sous la pression P_2 et à la température T_1 .

CHAPITRE 25 – MACHINES THERMIQUES

- Évolution BD : elle passe ensuite dans un échangeur qui permet les transferts thermiques entre le circuit primaire et le circuit secondaire. On peut décomposer en deux transformations ce qui se passe alors :
 - l'eau liquide s'échauffe de manière isobare (sous la pression P_2), jusqu'à l'état C (pression P_2 , température T_2) ;
 - l'eau liquide se vaporise entièrement, jusqu'à l'état D sous la pression P_2 et à la température T_2 .
- Évolution DE : la vapeur d'eau se détend de manière réversible dans une turbine calorifugée jusqu'à la pression P_1 et la température T_1 (état E). Durant cette détente, une fraction $(1 - x)$ de l'eau redevient liquide, et x reste gazeuse : x est donc le titre en vapeur dans l'état E .
- Évolution EA : la vapeur restant se condense à la température T_1 .

L'écoulement est stationnaire.

Dans le tableau suivant, on donne pour l'eau à 293 K et à 573 K : la pression de vapeur saturante P_{sat} en bar, les volumes massiques v_L du liquide et v_G de la vapeur en $\text{m}^3 \cdot \text{kg}^{-1}$, les enthalpies massiques h_L du liquide et h_G de la vapeur en $\text{kJ} \cdot \text{kg}^{-1}$ et enfin les entropies massiques s_L du liquide et s_G de la vapeur en $\text{kJ} \cdot \text{K}^{-1} \cdot \text{kg}^{-1}$.

$T(\text{K})$	P_{sat}	v_L	v_G	h_L	h_G	s_L	s_G
293	0,023 (P_1)			85	2540	0,3	8,7
573	80 (P_2)	$1,31 \cdot 10^{-3}$	0,026	1290	2890		6,0

La masse molaire de l'eau est $M = 18 \text{ g} \cdot \text{mol}^{-1}$. On rappelle la valeur de la constante des gaz parfaits : $R = 8,314 \text{ J} \cdot \text{K}^{-1} \cdot \text{mol}^{-1}$.

1. a. Quelle est la variation d'entropie d'un corps pur et sa variation d'enthalpie lors d'une vaporisation totale ?
b. Calculer l'entropie massique de l'eau liquide, $s_L(573)$, à la température de 573 K et sous une pression de 80 bar.
c. Sachant que la vapeur d'eau sous une pression de 0,023 bar et à la température de 293 K peut être considérée comme un gaz parfait, calculer son volume massique $v_G(293)$.
2. Tracer le cycle de l'eau sur un diagramme de Clapeyron (P, v) en plaçant les points correspondant aux états A, B, C, D et E .
3. a. Démontrer que la transformation DE est isentropique.
b. Calculer le titre en vapeur dans l'état E .
4. On notera w_a le travail reçu par l'alternateur par unité de masse du fluide passant dans la turbine (on négligera tout frottement). Calculer la valeur numérique de w_a .
5. a. Pour l'eau liquide, la capacité thermique massique est $c = 4,2 \text{ kJ} \cdot \text{K}^{-1} \cdot \text{kg}^{-1}$. Calculer le transfert thermique (par unité de masse de fluide écoulé) du système avec l'échangeur du circuit primaire q_{BD} .
b. On définit l'efficacité $e = \frac{w_a}{q_{BD}}$. La calculer. Que néglige-t-on dans cette définition ?
c. Calculer l'efficacité maximale qu'on aurait pu avoir avec les mêmes sources. Quelle conclusion peut-on en tirer ?

25.7 Moteur Diesel (★★)

Dans l'étude théorique du moteur Diesel, on suppose que le fonctionnement est décrit par le cycle de Diesel dans lequel la combustion est supposée s'effectuer à pression constante. Dans les moteurs actuels et notamment dans ceux dits rapides, on a cherché à réaliser une combustion procédant sous deux régimes : la combustion commence à volume constant et se termine à pression constante. Le cycle consiste alors en une compression adiabatique AB , un échauffement isochore BC , un échauffement isobare CD , une détente adiabatique DE et une compression isochore EA .

L'objet de ce problème est d'étudier un tel moteur. On considère un moteur Diesel à six cylindres fonctionnant suivant le cycle mixte à quatre temps. Le moteur à quatre temps effectue un cycle tous les deux tours. Il présente un rapport de compression volumétrique $a = \frac{V_A}{V_B} = 15$ et une cylindrée $\mathcal{C} = V_A - V_B = 15 \text{ L}$. V_A est le volume maximal d'air aspiré, V_B le volume disponible dans le cylindre au moment où commence l'injection du combustible, V_D le volume quand la combustion se termine. On note $c = \frac{V_D}{V_B}$ le rapport d'injection. On désigne par P_A la pression de l'air à l'aspiration, P_B celle au moment où commence l'injection du combustible et P_C la pression maximale atteinte dans le cylindre. On suppose que le cycle est décrit de manière quasistatique.

On admet que le gaz constitué essentiellement d'air est assimilable à un gaz parfait et que la masse de carburant est négligeable devant celle de l'air pour un cycle. On raisonnera sur un cycle fictif fermé sans se préoccuper des étapes consistant à évacuer les gaz issus de la combustion (échappement) pour les remplacer par de l'air frais (admission) qu'on modélisera par l'évolution isochore EA .

On note c_V et $c_P = 1,02 \text{ kJ.kg}^{-1}.\text{K}^{-1}$ les capacités thermiques massiques respectivement à volume et à pression constants, on admet qu'ils sont constants. Enfin $\gamma = \frac{c_P}{c_V} = 1,39$ et m désigne la masse d'air aspirée par cycle et par cylindre.

1. Représenter le cycle $ABCDEA$ dans un diagramme (P, V) . Vérifier graphiquement qu'il s'agit d'un cycle moteur.
2. Exprimer les transferts thermiques lors des transformations BC , CD et EA en fonction des températures, de m et de c_V et γ .
3. Justifier qu'on définit le rendement par $r_{th} = -\frac{W}{Q_{BC} + Q_{CD}}$.
4. Donner son expression en fonction des températures et de γ .
5. Exprimer les températures T_B , T_C , T_D et T_E en fonction de T_A , a , c , $d = \frac{P_D}{P_B}$ et γ .
6. En déduire l'expression du rendement en fonction de a , c , d et γ .
7. La combustion d'une masse m' de carburant produit un transfert thermique $\eta m'$ où $\eta = 42,2 \text{ MJ.kg}^{-1}$ désigne le pouvoir calorifique dont l'unité est J.kg^{-1} . On admet qu'il s'agit d'une constante indépendante des conditions durant la combustion. La masse de combustible pulvérisé dans un cylindre à chaque cycle est égale à αm avec $\alpha = 0,040$. Au début de la compression, la température est $T_A = 338 \text{ K}$ et la pression P_A . On désigne par μ la masse totale de combustible injectée par cycle et par cylindre et par μ_V celle qui brûle à volume

CHAPITRE 25 – MACHINES THERMIQUES

constant. Le cycle est caractérisé par $k_V = 100 \frac{\mu_V}{\mu}$. On note $M_{air} = 29 \text{ g.mol}^{-1}$ la masse molaire de l'air.

Exprimer la masse d'air contenue dans le cylindre m en fonction de a , $\mathcal{C}$, P_A , T_A , R et M_{air} .

8. Exprimer Q_{BC} en fonction de k_V , η , m et α .

9. En comparant les deux expressions obtenues pour Q_{BC} , établir l'expression de T_C et P_C en fonction de T_A , P_A , a , γ , k_V , η , α et c_V .

10. Exprimer Q_{CD} en fonction de k_V , η , m et α .

11. En déduire l'expression de T_D en fonction de T_A , a , γ , k_V , η , α et c_V .

12. Quelle valeur doit-on donner à k_V pour que la pression maximale dans le cylindre ne dépasse pas 65 bars sachant que $P_A = 1,0 \text{ bar}$?

13. Si k_V prend cette valeur, déterminer les valeurs de V_A , V_B , T_B , T_C , T_D , c , d , T_E et b .

14. Déterminer la valeur du rendement théorique.

15. Calculer la puissance du moteur sachant qu'il tourne à 2300 tours par minute.

16. La puissance réelle est-elle inférieure ou supérieure à cette valeur théorique ? Justifier votre réponse.

25.8 Moteur de Stirling, d'après concours PT (★)

On considère $n = 40.10^{-3} \text{ mol}$ d'air, considéré comme un gaz parfait de rapport $\gamma = \frac{C_P}{C_V}$ constant et égal à 1,4, subissant un cycle modélisé par les évolutions suivantes à partir de l'état A : $P_1 = 1 \text{ bar}$ et $T_1 = 300 \text{ K}$:

- compression isotherme réversible au contact de la source TH_1 à T_1 , jusqu'à l'état B, de volume $V_2 = \frac{V_1}{10}$,
- échauffement isochore au contact thermique de la source TH_2 à $T_2 = 600 \text{ K}$ jusqu'à l'état C de température T_2 ,
- détente isotherme réversible au contact de la source TH_2 jusqu'à l'état D de volume V_1 ,
- refroidissement isochore au contact thermique de la source TH_1 jusqu'à l'état A.

1. Calculer les valeurs numériques de P , V et T pour chacun des états A, B, C et D.

2. Représenter le cycle dans le diagramme de Clapeyron (P, V). Comment peut-on, sans calcul, savoir si le cycle proposé est celui d'un moteur ou d'un système mécaniquement récepteur ?

3. Calculer pour chaque étape le transfert thermique et le travail reçus par le fluide.

4. Commenter ces résultats. A-t-on bien un cycle moteur ?

5. Quel est, sur le plan énergétique, la production de ce système sur un cycle ? Quel est le coût sur le plan énergétique ? En déduire l'expression et la valeur du rendement.

6. Calculer l'entropie créée par irréversibilité au sein du système au cours du cycle. Quel type d'irréversibilité entre en jeu ici ?

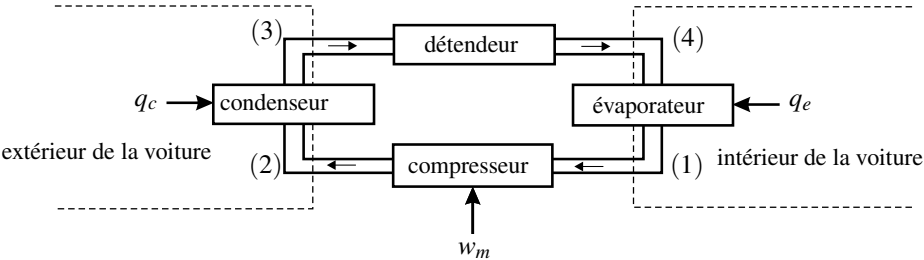
L'invention des frères Stirling (1816) a permis d'améliorer considérablement le rendement de la machine précédente. Leur idée est de faire en sorte que le gaz échange du transfert thermique au cours des transformation BC et DA, non pas avec TH_1 et TH_2 , mais avec un

système appelé régénérateur *RG* n’ayant aucun échange d’énergie avec l’extérieur autre que les échanges avec les gaz au cours des transformations *BC* et *DA*.

- 7. Justifier l’idée des frères Stirling.
- 8. Quelle est dans ces nouvelles conditions le rendement ? Ce rendement peut-il être encore amélioré sans changer les sources ?

25.9 Climatisation d’une voiture, d’après concours A.T.S. (MPSI) (★)

La quasi-totalité des véhicules neufs sont aujourd’hui équipés d’une climatisation. Pour refroidir l’air intérieur du véhicule, un fluide frigorigène, l’hydrofluorocarbone HFC connu sous le code R134a, effectue en continu des transferts énergétiques entre l’intérieur, l’extérieur du véhicule et le compresseur.



Sur le diagramme enthalpique (*P, h*) (voir figure page 912) de l’hydrofluorocarbone HFC, de masse molaire $M = 32 \text{ g}\cdot\text{mol}^{-1}$, sont représentés :

- la courbe de saturation de l’équilibre liquide-vapeur de l’hydrofluorocarbone HFC (en trait fort),
- les isothermes pour des températures comprises entre -40°C et 160°C par pas de 10°C ,
- les isentropiques pour des entropies massiques comprises entre $1,70 \text{ kJ}\cdot\text{K}^{-1}\cdot\text{kg}^{-1}$ et $2,25 \text{ kJ}\cdot\text{K}^{-1}\cdot\text{kg}^{-1}$, par pas de $0,05 \text{ kJ}\cdot\text{K}^{-1}\cdot\text{kg}^{-1}$,
- les isotitres en vapeur sous la courbe de saturation pour des titres massiques en vapeur x_G variant de 0 à 1 par pas de 0,1.

P est en bar et *h* en $\text{kJ}\cdot\text{kg}^{-1}$.

Lors de l’exploitation du diagramme, les mesures seront faites avec les incertitudes suivantes :

$\Delta h = \pm 5 \text{ kJ}\cdot\text{kg}^{-1}$, $\Delta s = \pm 50 \text{ J}\cdot\text{K}^{-1}\cdot\text{kg}^{-1}$, $\Delta x = \pm 0,05$, $\Delta T = \pm 5^\circ\text{C}$, $\frac{\Delta P}{P} = 5\%$.

- 1. Où sont sur le diagramme les domaines liquide, vapeur, équilibre liquide-vapeur du fluide ?
- 2. Dans quel domaine du diagramme le fluide à l’état gazeux peut-il être considéré comme un gaz parfait ?

On étudie dans la suite l’évolution du fluide au cours d’un cycle en régime permanent.

Le transfert thermique reçu par le fluide dans l’évaporateur permet la vaporisation isobare complète du fluide venant de (4) et conduit à de la vapeur à température $T_1 = 5^\circ\text{C}$ et pression $P_1 = 3 \text{ bar}$: point (1).

- 3. Placer le point (1) sur le diagramme. Relever la valeur de l’enthalpie massique h_1 et de l’entropie massique s_1 du fluide au point (1).

CHAPITRE 25 – MACHINES THERMIQUES

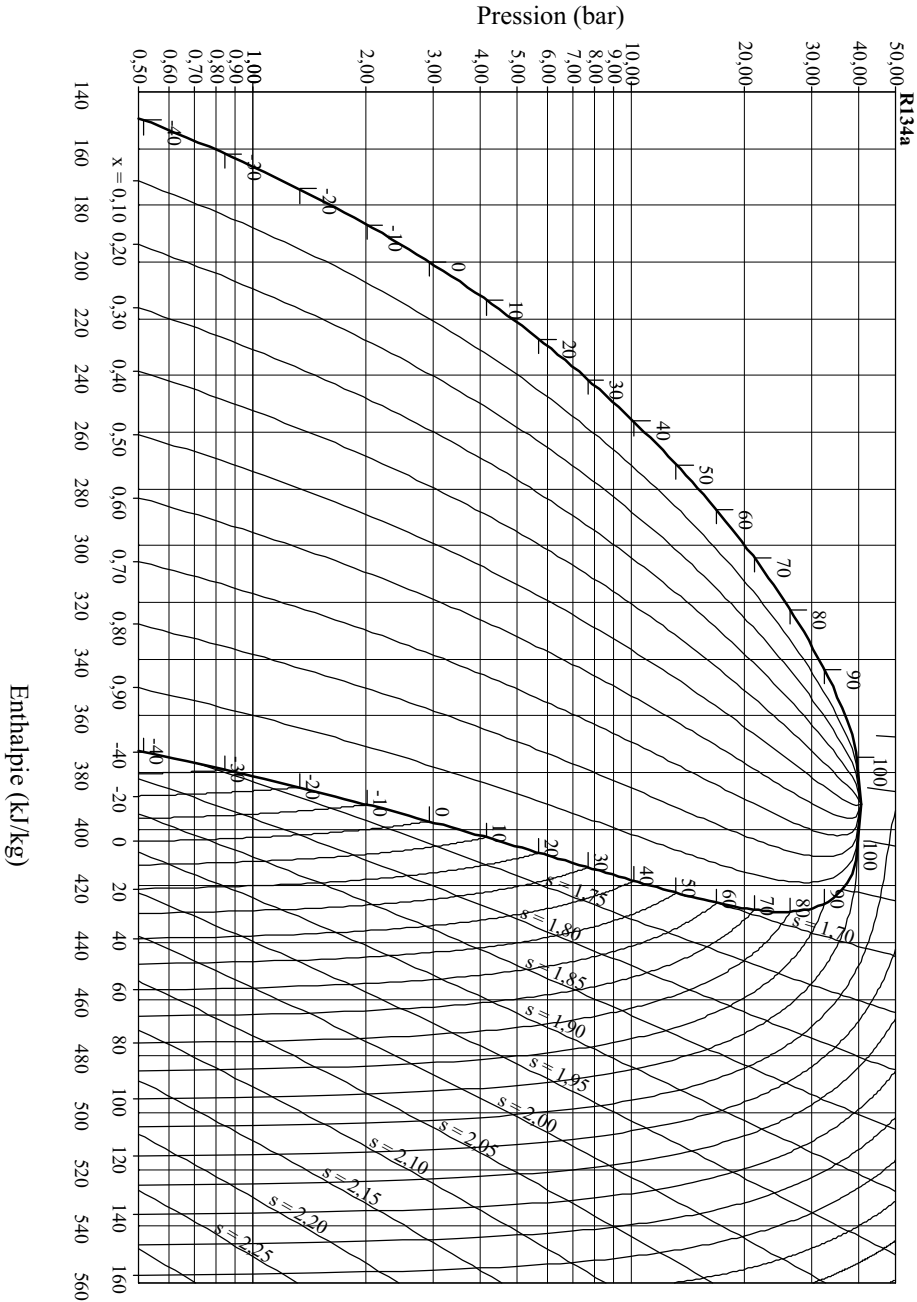


Diagramme enthalpique de l'hydrofluorocarbure HFC - R134a.

Le compresseur aspire la vapeur (1) et la comprime de façon isentropique avec un taux de compression $r = \frac{P_2}{P_1} = 6$.

4. Déterminer P_2 . Placer le point (2) sur le diagramme. Relever la valeur de la température T_2 et celle de l'enthalpie massique h_2 en sortie du compresseur.

5. Déterminer la valeur du travail mécanique massique w_m reçu par le fluide lors de son passage dans le compresseur. Commenter le signe de w_m .

Le fluide sortant du compresseur entre dans le condenseur dans lequel il est refroidi de manière isobare jusqu'à la température $T_3 = 60^\circ\text{C}$: point (3).

6. Placer le point (3) sur le diagramme. Relever la valeur de l'enthalpie massique h_3 en sortie du condenseur.

Le fluide sortant du condenseur est détendu dans le détendeur supposé adiabatique jusqu'à la pression de l'évaporateur P_1 : point (4).

7. Montrer que la transformation dans le détendeur est isenthalpique.

8. Placer le point (4) sur le diagramme et tracer le cycle complet. Relever la valeur de la température T_4 et le titre massique en vapeur x_4 en sortie du détendeur.

9. En déduire le transfert thermique massique q_e échangé par le fluide lors de son passage à travers l'évaporateur entre (4) et (1). L'air intérieur du véhicule est-il refroidi ?

10. Définir l'efficacité e , ou coefficient de performance, du climatiseur. Calculer sa valeur.

11. Comparer cette valeur à celle d'un climatiseur de Carnot fonctionnant entre la température de l'évaporateur et la température de liquéfaction du fluide sous la pression P_2 . Commenter le résultat obtenu.

12. Le débit massique du fluide est $D_m = 0,1 \text{ kg}\cdot\text{s}^{-1}$. Calculer la puissance thermique évacuée de l'intérieur du véhicule et la puissance mécanique consommée par le climatiseur.

25.10 Machine à vapeur : cycle de Rankine (MPSI) (★★)

Dans une machine à vapeur, l'eau décrit un cycle de Rankine :

- dans l'état A l'eau est à l'état de liquide saturant seul, dans les conditions de pression et température $P_1 = 0,2 \text{ bar}$ et $T_1 = 60^\circ\text{C}$;
- transformation AB : l'eau est comprimée de façon adiabatique et isentropique dans une pompe, jusqu'à la pression $P_2 = 15 \text{ bar}$,
- transformation BC : l'eau est injectée dans la chaudière et s'y réchauffe de manière isobare jusqu'à la température $T_2 = 200^\circ\text{C}$, telle que $P_{\text{sat}}(T_2) = P_2$;
- transformation CD : l'eau se vaporise entièrement à la température T_2 ;
- transformation DE : la vapeur est admise dans le cylindre à T_2 et P_2 et effectue une détente adiabatique et isentropique jusqu'à la température T_1 , on obtient un mélange liquide - vapeur ;
- transformation EA : le piston chasse le mélange liquide - vapeur dans le condenseur où il se liquéfie totalement.

1. Représenter le cycle précédent sur le diagramme (P, h) donné à la page 914.

2. Déduire de valeurs lues sur le diagramme le transfert thermique pour chaque transformation du cycle.

3. Calculer le rendement de ce moteur et le comparer au rendement de Carnot. Quelles sont les causes d'irréversibilité ?

CHAPITRE 25 – MACHINES THERMIQUES

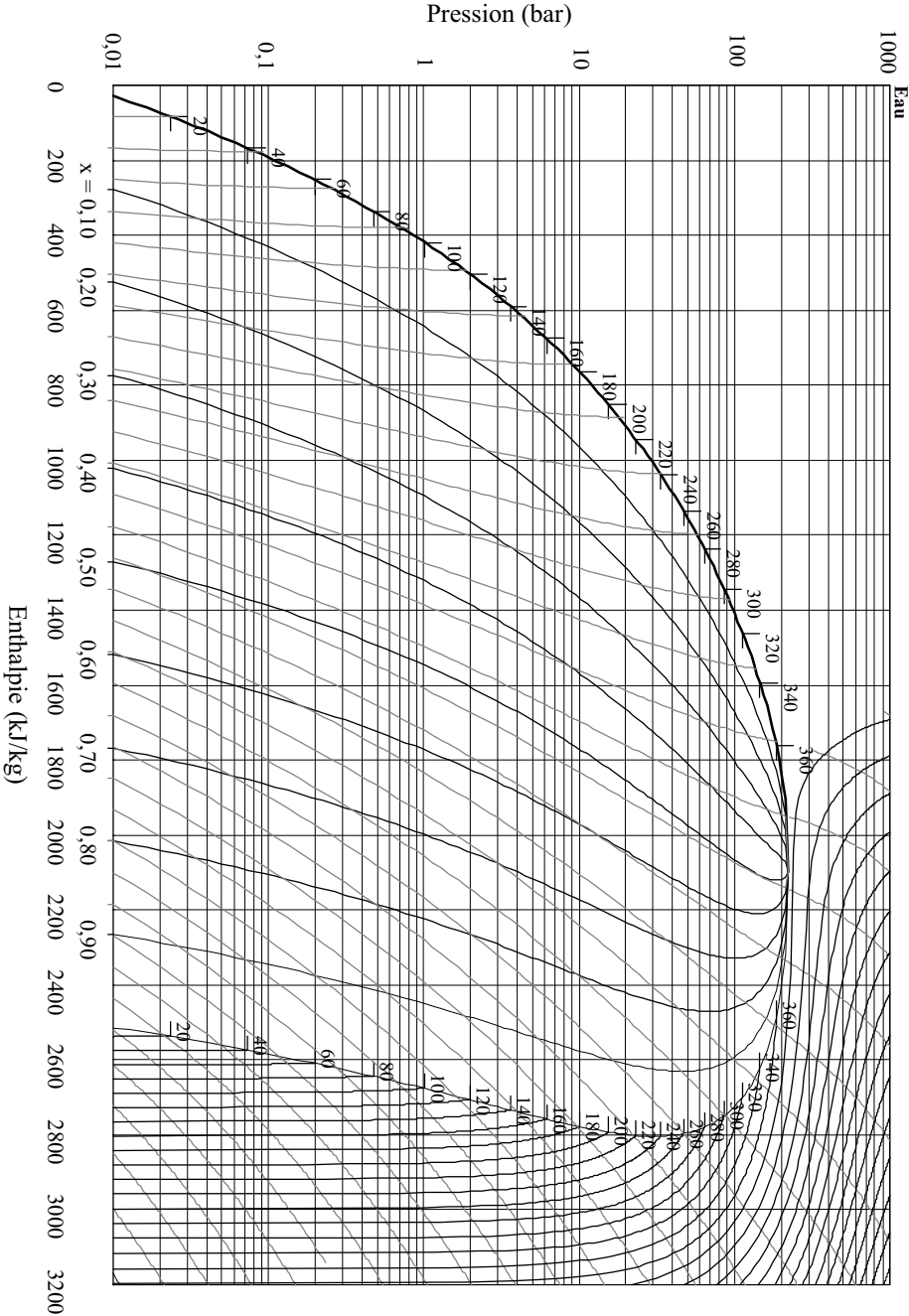


Diagramme enthalpique de l'eau (les courbes en gris sont les isentropes).

CORRIGÉS

25.1 Moteur réel

1. Le rendement du moteur étudié est $\rho = \frac{500}{1500} = \frac{1}{3}$. Le rendement du moteur de Carnot correspondant est $\rho_C = 1 - \frac{T_{fr}}{T_{ch}} = 1 - \frac{400}{650} = 0,385 > \rho$, ce qui normal.

2. Le second principe, appliqué au fluide sur un cycle, s'écrit :

$$\Delta S = 0 = S_{ech} + S_{créée} = \frac{Q_{fr}}{T_{fr}} + \frac{Q_{ch}}{T_{ch}} + S_{créée}$$

Or, $Q_{ch} = 1500\text{J}$ et $Q_{fr} = -W - Q_{ch}$ d'après le premier principe, avec $W = -500\text{ J}$, donc $Q_{fr} = -1000\text{ J}$. Finalement, $S_{créée} = 0,19\text{ J.K}^{-1}$ par cycle. L'entropie créée est strictement positive : le moteur est irréversible.

3. Si la dépense est la même (c'est-à-dire si Q_{ch} est la même), le premier principe appliqué à la machine réelle et à la machine de Carnot permet d'écrire : $W_{réel} - W_{rév} = Q_{fr,rév} - Q_{fr,réel}$.

Par le second principe, $Q_{fr,réel} = -\frac{T_{fr}}{T_{ch}} - T_{fr}S_{créée}$ et $Q_{fr,rév} = -\frac{T_{fr}}{T_{ch}}$. Il vient donc :

$$W_{réel} = W_{rév} + T_{fr}S_{cr}$$

Les travaux sont négatifs et l'entropie créée positive donc $|W_{réel}| < |W_{rév}|$: pour une dépense identique, le moteur réel fournit moins de travail que le moteur réversible.

25.2 Perte de performance d'un congélateur

1. Pour une machine réversible, $e_{rev} = \frac{T_{fr}}{T_{ch} - T_{fr}} = \frac{253}{292 - 253} = 6,5$. Donc : $\alpha = 0,31$.

2. Dans le cours, on a montré que : $\frac{1}{e} = \frac{T_{ch}}{T_{fr}} - 1 + \frac{T_{ch}S_{créée}}{Q_{fr}} = \frac{1}{e_{rev}} + \frac{T_{ch}S_{créée}}{Q_{fr}}$.

Si le terme $\frac{T_{ch}S_{créée}}{Q_{fr}}$ est doublé, l'efficacité devient e' avec :

$$\frac{1}{e'} - \frac{1}{e_{rev}} = 2 \frac{T_{ch}S_{créée}}{Q_{fr}} = 2 \left(\frac{1}{e} - \frac{1}{e_{rev}} \right) \text{ soit } e' = \frac{ee_{rev}}{2e_{rev} - e} \text{ soit } \alpha' = \frac{e'}{e_{rev}} = \frac{\alpha}{2 - \alpha}.$$

Numériquement : $\alpha' = 0,18$ et $e' = 1,2$.

25.3 Pompe à chaleur avec source de température variable

1. Les deux principes appliqués à la pompe à chaleur sur un cycle s'écrivent comme dans le cours, en remplaçant T_{ch} par T . Ainsi, le transfert thermique *cédé* par la pompe à la maison $Q_{cycle} = -Q_{ch}$ vérifie : $Q_{cycle} = -Q_{ch} = e_{pac}W \leq e_{pac,rev}W = \frac{T}{T - T_{fr}}W$.

D'autre part, d'après le premier principe appliqué à la maison : $Q_{cycle} = \Delta H = C\Delta T$.

Il vient donc : $W_{cycle} \leq C \frac{T}{T - T_{fr}} \Delta T$.

CHAPITRE 25 – MACHINES THERMIQUES

2. Le calcul de l'intégrale donne : $W_{\min} = C \left(T_2 - T_1 - T_{fr} \ln \frac{T_2}{T_1} \right)$. Cette valeur peut être atteinte, en théorie, avec une pompe à chaleur réversible.

3. Au cours du chauffage, la maison a reçu le transfert thermique : $Q = \Delta H = C(T_2 - T_1)$.

L'efficacité de la pompe sur l'opération de chauffage peut être définie par : $e = \frac{Q}{W}$.

Sa valeur maximale est : $e_{\max} = \frac{Q}{W_{\min}} = \frac{1}{1 - \frac{T_{fr}}{T_2 - T_1} \ln \frac{T_2}{T_1}}$.

Numériquement : $e_{\max} = 28,9$.

25.4 Pompe à chaleur avec un gaz parfait

1. Voir figure ci-contre.

2. Les relations de Laplace sont applicables si la transformation est adiabatique et réversible et elles ne concernent que les gaz parfaits. L'expression de la relation de Laplace en fonction de la pression et de la température est $P^{1-\gamma} T^\gamma$ constante.

3. On applique la relation de Laplace à la transformation AB et on obtient $T_B = T_0 a^\beta = 448 \text{ K}$ puis à la transformation CD ce qui donne $T_D = T_1 a^{-\beta} = 188 \text{ K}$.

4. L'efficacité est définie par $e = -\frac{Q_{BC}}{W} = \frac{Q_{BC}}{Q_{BC} + Q_{DA}}$ en utilisant l'expression du premier principe $W + Q_{BC} + Q_{DA} = \Delta U = 0$ sur un cycle.

5. Les transformations sont isobares donc les transferts thermiques sont égaux à la variation d'enthalpie soit, en tenant compte du caractère parfait des gaz $Q_{BC} = C_P (T_C - T_B)$ et $Q_{DA} = C_P (T_A - T_D)$. On reporte dans l'expression de l'efficacité et on en déduit $e = \frac{1}{1 - a^{-\beta}} = 2,71$.

6. Le cycle de Carnot est composé de deux transformations adiabatiques réversibles et de deux transformations isothermes au cours desquelles ont lieu les transferts thermiques.

7. Pour ce cycle réversible, le calcul du cours montre que : $e_C = \frac{T_1}{T_1 - T_0} = 19,9$.

8. On a donc $e < e_C$, parce que le cycle n'est pas réversible : dans les transformations BC et CD le gaz est mis en contact avec la source (chaude ou froide) alors que sa température n'est pas égale à celle de la source, il y a donc irréversibilité thermique.

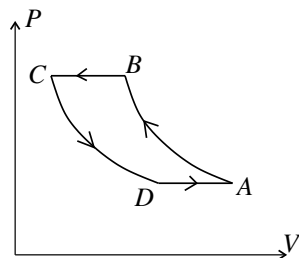
9. Le bilan d'entropie donne $\Delta S = S_{\text{éch}} + S_{\text{créée}} = 0$ sur le cycle. On a donc

$$S_{\text{créée}} = -S_{\text{éch}} = -\frac{Q_{BC}}{T_1} - \frac{Q_{DA}}{T_0} = \frac{nR}{\beta} \left(x + \frac{1}{x} - 2 \right) = \frac{nR}{\beta} \left(\sqrt{x} - \frac{1}{\sqrt{x}} \right)^2.$$

10. $S_{\text{créée}}$ est positive quel que soit x .

11. L'application numérique donne $S_{\text{cr}} = 4,84 \text{ J.K}^{-1}.\text{mol}^{-1}$.

12. Le régime est permanent donc ce qui est perdu est égal à ce qui est fourni. On en déduit que $\mathcal{P} = \frac{Q_f}{e} = 7,4 \text{ kW}$.



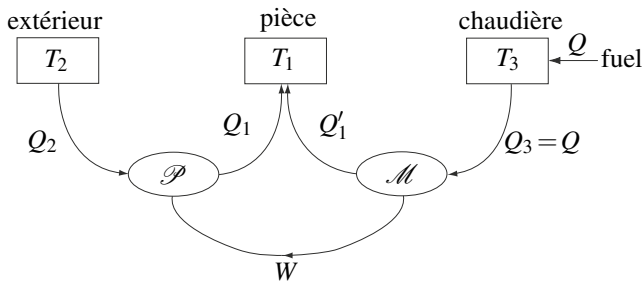
25.5 Optimisation du chauffage d'un local

1. Dans tout l'exercice, on note T la température en kelvins et θ la température en degrés Celsius.

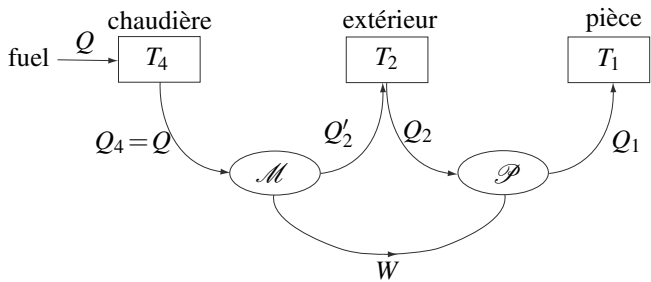
La pompe à chaleur réversible, fonctionnant entre le local de température T_1 et l'extérieur de température T_2 , a pour efficacité : $e_{\text{rev}} = \frac{T_1}{T_1 - T_2} = 13,3$. Ainsi, la puissance consommée est 13,3 fois plus petite que la puissance thermique fournie au local. La puissance du moteur faisant fonctionner cette pompe est donc : $P = \frac{32 \cdot 10^6}{3600 \times 13,3} = 668 \text{ W}$.

2. a. Les schémas de principe des deux dispositifs proposés sont les suivants ($\mathcal{P}$ désigne la pompe à chaleur et $\mathcal{M}$ le moteur) :

• Dispositif du premier conseiller :



• Dispositif du deuxième conseiller :



Dans chaque cas, on applique le premier principe et le second principe (sous forme de l'égalité de Clausius) au fluide qui décrit des cycles, pour la pompe à chaleur et pour le moteur. Attention, si on appelle W le travail reçu par le fluide au cours d'un cycle dans la pompe à chaleur, le travail que reçoit le fluide dans le moteur est $-W$.

• Dispositif du premier conseiller :

◦ Pompe à chaleur :

$$\begin{cases} W + Q_1 + Q_2 = 0 \\ \frac{Q_1}{T_1} + \frac{Q_2}{T_2} = 0 \end{cases}$$

CHAPITRE 25 – MACHINES THERMIQUES

◦ Moteur :

$$\begin{cases} -W + Q'_1 + Q = 0 \\ \frac{Q'_1}{T_1} + \frac{Q}{T_3} = 0 \end{cases}$$

Le local reçoit le transfert thermique :

$$Q_{\text{local}} = -Q_1 - Q'_1 = \frac{1 - \frac{T_1}{T_3}}{1 - \frac{T_1}{T_2}} Q + \frac{T_1}{T_3} Q = \frac{T_3 - T_2}{T_1 - T_2} Q.$$

Numériquement, $Q_{\text{local}} = 5,8$. Ainsi, avec a litre de fuel on chauffe le local pendant 5,8 jours.

• *Dispositif du deuxième conseiller :*

◦ Pompe à chaleur :

$$\begin{cases} W + Q_1 + Q_2 = 0 \\ \frac{Q_1}{T_1} + \frac{Q_2}{T_2} = 0 \end{cases}$$

◦ Moteur :

$$\begin{cases} -W + Q + Q'_2 = 0 \\ \frac{Q}{T_4} + \frac{Q'_2}{T_2} = 0 \end{cases}$$

Le local reçoit le transfert thermique :

$$Q_{\text{local}} = -Q_1 = \frac{1 - \frac{T_2}{T_4}}{1 - \frac{T_2}{T_1}} Q = \frac{T_4 - T_2}{T_1 - T_2} Q.$$

Numériquement, $Q_{\text{local}} = 6,5Q$. Ainsi, avec a litre de fuel on chauffe le local pendant 6,5 jours.

b. Le deuxième dispositif semble meilleur mais uniquement parce $T_4 > T_3$. Si $T_3 = T_4$, les deux systèmes sont équivalents. Le premier semble plus astucieux car l'énergie thermique fournie à la source froide du moteur est utilisée pour chauffer la salle mais le rendement du moteur du deuxième dispositif est meilleur puisqu'il fonctionne entre la source la plus froide et la source la plus chaude.

c. Augmenter la température de la chaudière auxiliaire augmente le rendement, mais une installation supportant une température plus élevée coûte aussi plus cher d'où un investissement plus grand au départ. D'autre part, plus cette température est élevée, plus il y a de pertes par fuites thermiques.

25.6 Étude de la turbine d'un réacteur à eau pressurisée (REP)

1. a. Lors d'une vaporisation totale d'une masse m d'un corps à température constante T , la variation d'enthalpie et la variation d'entropie sont :

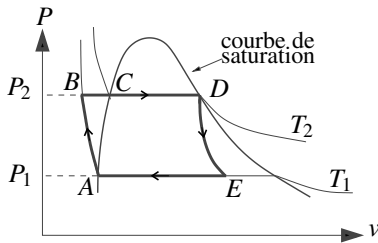
$$\Delta H = m(h_G - h_L) = m\Delta_{\text{vap}}h \quad \text{et} \quad \Delta S = m(s_V - s_L) = m\Delta_{\text{vap}}s = m \frac{\Delta_{\text{vap}}h}{T}.$$

b. D'après ce qui précède :

$$s_L(573) = s_G - \frac{h_G - h_L}{T} = 8,7 \cdot 10^3 = 6,0 - \frac{2890 - 1290}{573} = 3,2 \text{ kJ} \cdot \text{K}^{-1} \cdot \text{kg}^{-1}.$$

c. L'équation des gaz parfaits donne : $v_G = \frac{RT}{MP}$ soit $v_G(293) = 58,8 \text{ m}^3 \cdot \text{kg}^{-1}$.

2. On représente le cycle dans le diagramme de Clapeyron :



3. a. La turbine est calorifugée donc la transformation est adiabatique. De plus l'énoncé suppose qu'elle est réversible. La transformation est adiabatique et réversible donc isentropique.

b. On utilise le fait que la transformation DE est isentropique, on calcule l'entropie du mélange liquide-vapeur dans ces deux états :

- en D, il n'y a que de la vapeur à 573 K, donc : $s_D = ms_G(573)$,
- en E, il y a une fraction x de vapeur à 293 K, donc : $s_E = m(xs_G(293) + (1-x)s_L(293))$.

Ainsi la relation $s_D = s_E$, donne : $x = \frac{s_G(573) - s_L(293)}{s_G(293) - s_L(293)} = 0,68$.

4. Dans l'alternateur le fluide est en écoulement. D'après le premier principe pour un écoulement : $h_E - h_D = w_{DE}$, où w_{DE} est le travail massique reçu par le fluide dans la turbine (il n'y a pas de transfert thermique au cours de la transformation DE). Ce travail est l'opposé du travail reçu par l'alternateur, donc :

$$w_a = -w_{DE} = h_D - h_E.$$

Or : $h_D = h_G(573)$ et $h_E = (xh_G(293) + (1-x)h_L(293))$. Il vient : $w_a = 1,140 \cdot 10^6 \text{ J} \cdot \text{kg}^{-1}$.

5. a. Le premier principe pour un fluide en écoulement entre B et D s'écrit :

$$q_{BD} = h_D - h_B,$$

car il n'y a pas d'échange de travail autre que celui des forces de pression dans la transformation BD. La variation d'enthalpie entre B et D est la somme de la variation d'enthalpie de l'échauffement isobare BC puis de la vaporisation complète CD :

$$h_D - h_B = (h_D - h_C) + (h_C - h_B) = c(T_2 - T_1) + h_G(T_2) - h_L(T_2).$$

On calcule ainsi : $q_{BD} = 2,78 \cdot 10^6 \text{ J} \cdot \text{kg}^{-1}$.

b. Le rendement est : $e = \frac{1,14}{2,86} = 0,40$.

Dans cette définition, on omet de tenir compte du travail nécessaire pour faire fonctionner la pompe (transformation AB), le rendement est donc légèrement plus faible mais ce travail est négligeable devant w_a .

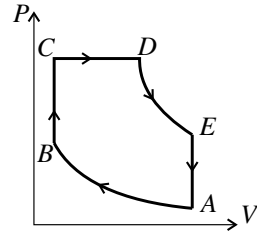
CHAPITRE 25 – MACHINES THERMIQUES

c. Le rendement maximal qu'on aurait pu avoir est celui de Carnot pour un moteur réversible : $e_c = 1 - \frac{T_1}{T_2} = 0,49$. Le rendement obtenu est inférieur, donc certaines transformations ne sont pas réversibles, ce sont les transformations AB et BC .

25.7 Moteur Diesel

1. Voir figure ci-contre.

On a un cycle moteur s'il fournit du travail donc avec la convention usuelle de compter positivement ce qui est reçu, on doit avoir $W < 0$ pour un moteur. Dans ces conditions, l'interprétation graphique du travail impose d'avoir un sens horaire pour parcourir le cycle, ce qui est bien le cas.



2. La transformation BC est isochore donc le travail est nul $W_{BC} = 0$ et le transfert thermique est $Q_{BC} = \Delta U_{BC} = mc_V (T_C - T_B)$.

La transformation CD est isobare donc le calcul du travail donne $W_{CD} = \Delta(PV)_{CD}$ et on en déduit celui du transfert thermique $Q_{CD} = \Delta H_{CD} = m\gamma c_V (T_D - T_C)$.

La transformation EA est isochore donc le travail est nul $W_{EA} = 0$ et le calcul du transfert thermique donne $Q_{EA} = \Delta U_{EA} = mc_V (T_A - T_E)$.

3. Les transferts thermiques sont tels que $Q_{BC} > 0$ et $Q_{CD} > 0$ donc ils correspondent à une dépense. Le travail W est utile. Par conséquent, le rendement est défini par $r_{th} = -\frac{W}{Q_{BC} + Q_{CD}}$.

4. L'écriture du premier principe sur le cycle donne $\Delta U = 0$ soit en explicitant la variation d'énergie interne $\Delta U = W + Q_{AB} + Q_{BC} + Q_{CD} + Q_{DE} + Q_{EA}$.

Par ailleurs, on a $Q_{AB} = Q_{DE} = 0$ car ces transformations sont adiabatiques. On en déduit :

$$r_{th} = 1 + \frac{Q_{EA}}{Q_{BC} + Q_{CD}} = 1 + \frac{T_A - T_E}{\gamma(T_D - T_C) + T_C - T_B}.$$

5. La transformation AB est adiabatique et réversible. De plus, elle concerne un gaz parfait donc on peut appliquer les relations de Laplace par exemple $TV^{\gamma-1}$ constant. On en déduit $T_B = a^{\gamma-1}T_A$.

La transformation BC est isochore donc en utilisant la loi des gaz parfaits et la relation précédente, on a $T_C = \frac{P_C}{P_B}T_B = da^{\gamma-1}T_A$.

La transformation CD est isobare donc en utilisant la loi des gaz parfaits et la relation précédente, on obtient $T_D = \frac{V_D}{V_C}T_C = cda^{\gamma-1}T_A$.

La transformation DE est adiabatique et réversible et elle s'applique à un gaz parfait donc les relations de Laplace sont valables comme $TV^{\gamma-1}$ constant, donc $T_E = c^\gamma dT_A$.

6. En reportant les expressions de la question précédente, on obtient

$$r_{th} = 1 + \frac{1 - c^\gamma d}{(d-1)\gamma + (c-1)d} a^{1-\gamma}.$$

7. La relation des gaz parfaits avec $n = \frac{m}{M_{\text{air}}}$ permet d'écrire $m = \frac{M_{\text{air}} P_A V_A}{RT_A}$ ainsi que $\mathcal{C} = V_A - V_B = V_A \frac{a-1}{a}$. On en déduit $m = \frac{a}{a-1} \frac{M_{\text{air}} P_A \mathcal{C}}{RT_A}$.

8. Par définition, $Q_{BC} = \eta \mu_V = \frac{k_V}{100} \eta \alpha m$.

9. Comme $Q_{BC} = m c_V (T_C - T_B)$, il vient : $T_C = a^{\gamma-1} T_A + \frac{k_V \eta \alpha}{100 c_V}$.

La relation des gaz parfaits donne $P_C = \frac{V_A}{V_C} \frac{T_C}{T_A} P_A$ donc : $P_C = a \left(a^{\gamma-1} + \frac{k_V \eta \alpha}{100 c_V T_A} \right) P_A$.

10. Maintenant on a $Q_{CD} = \eta (\mu - \mu_V) = \left(1 - \frac{k_V}{100} \right) \eta \alpha m$.

11. Comme $Q_{CD} = m C_P (T_D - T_C)$, il vient : $T_D = a^{\gamma-1} T_A + \frac{\eta \alpha}{\gamma c_V} \left(\frac{k_V}{100} (\gamma - 1) + 1 \right)$.

12. On veut $P_C < P_{\text{max}}$. Il faut pour cela que : $k_V < \frac{100 C_P T_A}{\eta \alpha a \gamma} \left(\frac{P_{\text{max}}}{P_A} - a^\gamma \right)$ soit, numériquement : $k_V < 21,4$.

13. Les applications numériques donnent :

- $V_A = \frac{a}{a-1} \mathcal{C} = 3,21 \text{ L}$,
- $V_B = \frac{V_A}{a} = 0,214 \text{ L}$,
- $T_B = a^{\gamma-1} T_A = 972 \text{ K}$,
- $T_C = a^{\gamma-1} T_A + \frac{k_V \eta \alpha}{100 c_V} = 1464 \text{ K}$,
- $T_D = a^{\gamma-1} T_A + \frac{\eta \alpha}{\gamma c_V} \left(\frac{k_V}{100} (\gamma - 1) + 1 \right) = 2765 \text{ K}$,
- $c = \frac{T_D}{T_C} = 1,89$,
- $d = \frac{T_C}{T_B} = 1,51$,
- $T_E = c^\gamma d T_A = 1236 \text{ K}$,

14. On en déduit le rendement théorique $r_{th} = 1 + \frac{1 - c^\gamma d}{(d-1)\gamma + (c-1)d} a^{1-\gamma} = 0,55$.

15. De même, $m = \frac{a}{a-1} \frac{M_{\text{air}} P_A \mathcal{C}}{RT_A} = 3,31 \text{ g}$, ce qui permet de calculer $Q_{BC} = \frac{k_V}{100} \eta \alpha m = 1136 \text{ J}$ et $Q_{CD} = \left(1 - \frac{k_V}{100} \right) \eta \alpha m$. On en déduit $W = r_{th} (Q_{BC} + Q_{CD}) = 3073 \text{ K}$ par cycle et par cylindre. La puissance $\mathcal{P}$ est le produit du nombre de cylindres 6, du nombre de cycles par tour 0,5, du nombre de tours par seconde $\frac{2300}{60}$ et du travail soit $\mathcal{P} = 353 \text{ kW}$.

16. Le cycle est réel et non réversible donc la puissance réelle est inférieure.

CHAPITRE 25 – MACHINES THERMIQUES

25.8 Moteur de Stirling, d'après concours PT

1. En utilisant les informations de l'énoncé et l'équation d'état du gaz parfait on obtient :

état A	état B	état C	état D
$P_1 = 1 \text{ bar}$	$P_2 = 10P_1 = 10 \text{ bar}$	$P_3 = P_2 \frac{T_2}{T_1} = 20 \text{ bar}$	$P_4 = P_3 \frac{V_2}{V_1} = 2 \text{ bar}$
$V_1 = \frac{nRT_1}{P_1} = 0,98 \text{ L}$	$V_2 = \frac{V_1}{10} = 0,098 \text{ L}$	V_2	V_1
$T_1 = 300 \text{ K}$	T_1	$T_2 = 600 \text{ K}$	T_2

2. Le cycle est décrit dans le sens horaire : il s'agit d'un cycle moteur.

3. La transformation AB est mécaniquement réversible donc : $W_{AB} = \int_A^B -PdV = \int_A^B \frac{nRT_1}{V} dV$
 $= -nRT_1 \ln \frac{V_2}{V_1} = nRT_1 \ln 10 = 230 \text{ J}.$

Les transformations BC et DA sont isochores donc $W_{BC} = W_{DA} = 0.$

Par le même calcul que W_{AB} , on trouve : $W_{CD} = -nRT_2 \ln 10 = -459 \text{ J}.$

D'après le premier principe appliqué au système formé par le gaz sur la transformation AB on trouve : $Q_{AB} = \Delta U_{AB} - W_{AB}$; la transformation AB étant isotherme $\Delta U_{AB} = 0$ car le gaz parfait suit la première loi de Joule. Finalement : $Q_{AB} = -nRT_1 \ln 10 = -230 \text{ J}.$

En appliquant le premier principe au même système sur la transformation BC on trouve : $Q_{BC} = \Delta U_{BC} = \frac{nR}{\gamma-1} (T_2 - T_1) = 249 \text{ J}.$

Il vient de même : $Q_{CD} = nRT_2 \ln 10 = 460 \text{ J}$ et $Q_{DA} = -\frac{nR}{\gamma-1} (T_2 - T_1) = -249 \text{ J}.$

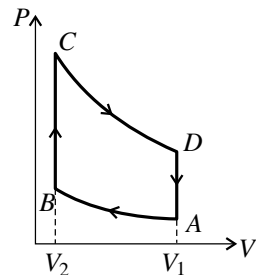
4. Le travail total échangé par le gaz au cours du cycle est : $W = W_{AB} + W_{CD} = -nR(T_2 - T_1) \ln(10) < 0.$ Le gaz cède donc du travail, c'est bien un cycle moteur.

5. L'énergie utile est $-W$ et l'énergie coûteuse $Q_{BC} + Q_{CD}$, transfert thermique reçu de la source chaude TH_2 . Le rendement est donc :

$$\rho = -\frac{W}{Q_{BC} + Q_{CD}} = \frac{(T_2 - T_1) \ln 10}{\frac{1}{\gamma-1} (T_2 - T_1) + T_2 \ln 10} = 0,32.$$

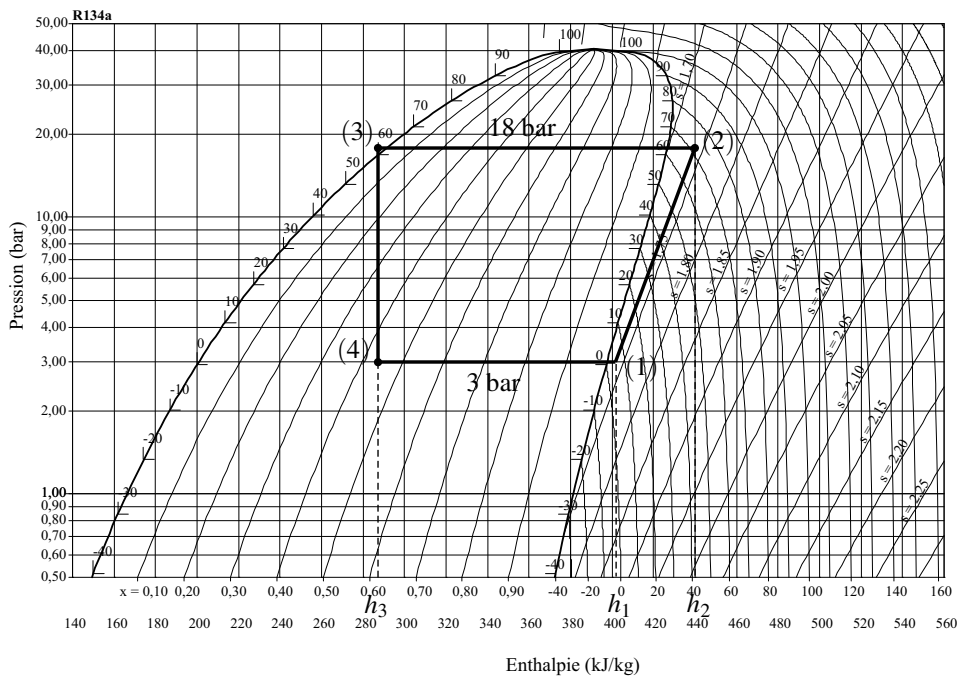
6. Le bilan d'entropie pour le système formé par le gaz sur le cycle complet s'écrit : $0 = S_{\text{éch}} + S_{\text{créée}}$, donc : $S_{\text{créée}} = -S_{\text{éch}} = -\frac{Q_{DA} + Q_{AB}}{T_1} - \frac{Q_{BC} + Q_{CD}}{T_2} = \frac{nR(T_2 - T_1)^2}{(\gamma-1)T_1T_2} > 0.$

Numériquement : $S_{\text{créée}} = 0,42 \text{ J} \cdot \text{K}^{-1}$. La création d'entropie est due aux transferts thermiques lors des transformations BC et DA au cours desquelles le système n'a pas la même température que la source avec laquelle il est contact (irréversibilité thermique).



7. On a $Q_{BC} = -Q_{DA}$. L'idée des frères Stirling est que le régénérateur reçoit du gaz ce transfert thermique au cours de l'étape DA et le lui rend au cours de l'étape BC.
8. Alors TH_2 ne doit plus fournir que Q_{CD} . Le rendement devient : $\rho = -\frac{W}{Q_{CD}} = 1 - \frac{T_2}{T_1}$. C'est le rendement de Carnot, donc le rendement maximal possible avec ces deux sources.

25.9 Climatisation d'une voiture, d'après concours A.T.S. (MPSI)



1. Le domaine de l'équilibre liquide-vapeur se trouve entre la courbe de saturation et l'axe des abscisses ; le domaine de la vapeur est à droite, du côté des plus grandes enthalpies massiques (donc des plus grandes températures) ; le domaine du liquide est à gauche, du côté des plus petites enthalpies massiques (donc des plus petites températures).
2. Le gaz parfait suit la deuxième loi de Joule : son enthalpie massique est fonction uniquement de la température. Donc, pour un gaz parfait, si $T = \text{constante}$ alors $h = \text{constante}$ et les isothermes sont des isenthalpes soit des droites verticales. Sur le diagramme c'est le cas dans la zone $P < 0,8 \text{ bar}$ et $h > 50 \text{ kJ}\cdot\text{kg}^{-1}$, en bas à droite du diagramme. Dans cette zone le fluide réel se comporte comme un gaz parfait.
3. On place le point (1) sur le diagramme, sur l'isobare $P_1 = 3 \text{ bar}$ et entre les isothermes 0°C et 10° , à peu près au milieu.
On lit à l'abscisse de ce point : $h_1 = 405 \text{ kJ}\cdot\text{kg}^{-1}$.
Ce point se trouve entre les isentropes $1,70$ et $1,75 \text{ kJ}\cdot\text{K}^{-1}\cdot\text{kg}^{-1}$, plus près de la seconde ; la réponse : $s_1 = 1,75 \text{ kJ}\cdot\text{K}^{-1}\cdot\text{kg}^{-1}$ est suffisante pour la précision demandée par l'énoncé.

CHAPITRE 25 – MACHINES THERMIQUES

4. $P_2 = 3P_1 = 18$ bar. La compression étant isentropique, le point (2) se trouve à l'intersection de l'isobare 18 bar, droite horizontale, et de l'isentropie $1,75 \text{ kJ} \cdot \text{K}^{-1} \cdot \text{kg}^{-1}$.

On lit à l'abscisse de ce point : $h_2 = 440 \text{ kJ} \cdot \text{kg}^{-1}$.

Et ce point se trouve pratiquement sur l'isotherme 70°C ; avec la précision demandée par l'énoncé, $T_2 \simeq 70^\circ\text{C}$.

5. Le premier principe pour un fluide en écoulement stationnaire, appliqué entre l'entrée et la sortie du compresseur, s'écrit : $\Delta h = h_2 - h_1 = w_m + 0$ car la compression est adiabatique, soit : $w_m = 440 - 405 = 35 \text{ kJ} \cdot \text{kg}^{-1}$.

6. Le point (3) se trouve à l'intersection de l'isobare 18 bar et de l'isotherme 60°C . Le point se trouve dans la zone du liquide où l'isotherme n'est pas tracée. On sait que c'est une droite verticale que l'on peut compléter pour trouver le point (3).

On lit à l'abscisse de ce point : $h_3 = 285 \text{ kJ} \cdot \text{kg}^{-1}$.

7. Dans le détendeur, le fluide ne reçoit pas de transfert thermique, ni de travail autre que celui des forces de pression. Le premier principe pour un fluide en écoulement stationnaire s'écrit donc : $\Delta h = h_4 - h_3 = q + w = 0$, soit $h_4 = h_3$. La transformation est isenthalpique.

8. Le point (4) se trouve à l'intersection de l'isenthalpe (droite verticale) passant par le point (3) et de l'isobare 3 bar (droite horizontale).

Ce point se trouve pratiquement sur l'isotherme 0°C donc $T_4 = 0^\circ\text{C}$.

Il se trouve entre les isotitres 0,40 et 0,50. Le titre massique en vapeur en ce point est : $x_4 = 0,45$.

9. Le premier principe pour un fluide en écoulement stationnaire, appliqué entre l'entrée et la sortie de l'évaporateur, s'écrit : $h_1 - h_4 = q_e$ car il n'y a pas de travail autre que celui des forces de pression, soit : $q_e = 405 - 285 = 120 \text{ kJ} \cdot \text{kg}^{-1}$.

L'air intérieur à la voiture est bien refroidi car $q_e > 0$: le transfert thermique est reçu par le fluide.

10. L'efficacité du climatiseur est le rapport de l'énergie utile, q_e , divisée par l'énergie coûteuse, w_m , soit : $e = \frac{q_e}{w_m} = \frac{120}{35} \simeq 3$.

11. L'efficacité d'un climatiseur réversible fonctionnant entre la température de l'évaporateur T_4 et la température d'équilibre liquide - vapeur pour 18 bar, soit environ T_3 serait :

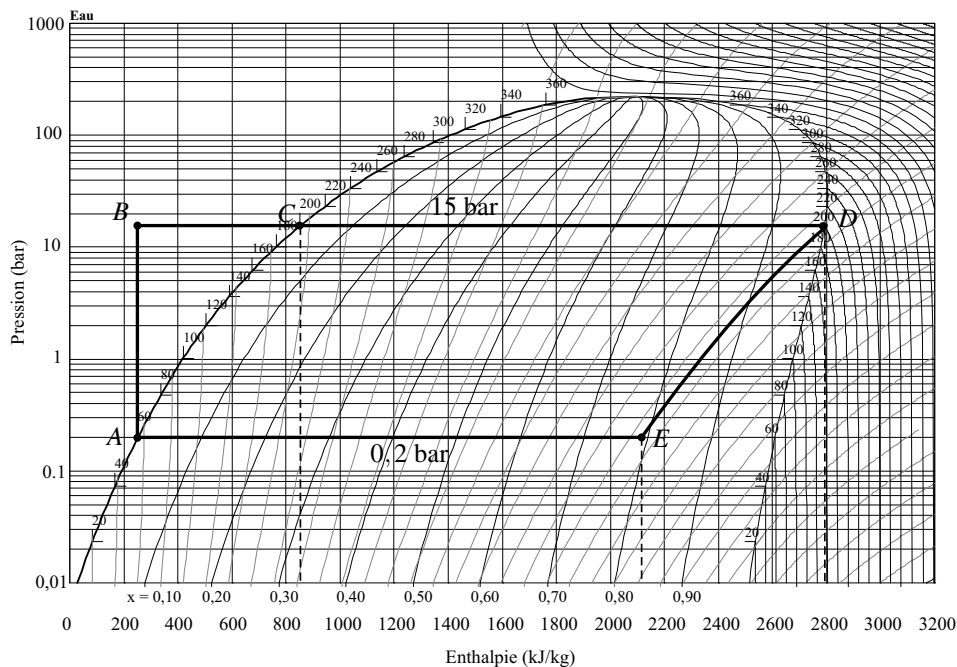
$$e_{\text{rev}} = \frac{T_4}{T_3 - T_4} = \frac{273}{60 - 0} = 4,6. \text{ Elle est plus grande que l'efficacité de la machine réelle.}$$

C'est le signe que la machine réelle n'est pas réversible. La transformation dans le détendeur est irréversible.

12. Pendant une durée Δt , une masse de fluide $D_m \Delta t$ passe dans l'évaporateur. L'énergie thermique prise à l'intérieur de la voiture est donc $Q_e = m q_e = D_m q_e \Delta t$. La puissance thermique évacuée de l'intérieur de la voiture est donc : $\mathcal{P}_e = D_m q_e = 12 \text{ kW}$.

25.10 Machine à vapeur : cycle de Rankine (MPSI)

1. Voir figure ci-dessous. La transformation AB est le long de l'isentropie passant par A qui est, dans le domaine du liquide, quasiment verticale. Les points B , C et D sont sur la même isobare dont le niveau est donné par le palier de liquéfaction à la température $T_2 = 200^\circ\text{C}$. La transformation DE est le long d'une isentropie que l'on dessine par analogie avec les isentropes voisines.



2. On calculera des transferts thermiques massiques.

$q_{AB} = 0$ car la transformation est adiabatique et isentropique. D'après le premier principe pour un fluide en écoulement stationnaire, $q_{BC} = h_C - h_B - w_{BC}$ où w_{BC} est le travail des forces autres que les forces de pression, nul dans cette transformation. Donc : $q_{BC} = h_C - h_B = 860 - 240 = 620 \text{ kJ} \cdot \text{kg}^{-1}$, en lisant les valeurs des enthalpies massiques à l'abscisse des points sur le diagramme.

De même : $q_{CD} = h_D - h_C = 2800 - 860 = 1940 \text{ kJ} \cdot \text{kg}^{-1}$.

La transformation DE étant isentropique : $h_{DE} = 0$.

De même que pour BC et CD : $q_{EA} = h_A - h_E = 240 - 2080 = -1840 \text{ kJ} \cdot \text{kg}^{-1}$.

3. Le rendement d'un moteur thermique est : $\rho = -\frac{W}{Q_{ch}} = 1 + \frac{Q_{fr}}{Q_{ch}}$ où Q_{ch} est le transfert thermique reçu de la source chaude et Q_{fr} le transfert thermique reçu de la source froide au cours d'un cycle. Ici : $Q_{ch} = m(q_{BC} + q_{CD})$, où m est la masse d'eau, et $Q_{fr} = mq_{EA}$. On a donc :

$$\rho = 1 + \frac{q_{EA}}{q_{BC} + q_{CD}} = 1 - \frac{1840}{620 + 1940} = 0,28.$$

Le rendement de Carnot pour les températures des sources est : $\rho_C = 1 - \frac{T_1}{T_2} = 1 - \frac{273 + 60}{273 + 200} = 0,30$. Il est plus grand que le rendement calculé à cause des irréversibilités.

4. La transformation BC est irréversible parce que l'eau liquide est mise en contact avec la source à la température T_2 alors que sa propre température est plus faible, proche de T_1 dans l'état B .

CHAPITRE 25 – MACHINES THERMIQUES

Cinquième partie

Induction et forces de Laplace

Trouver la bibliothèque des livres ici : <https://1024terabox.com/s/1mg0wxl22G8NjM8MA2RzgFw>

Le champ magnétique

De la relativité générale à la physique quantique, en passant par les moteurs électriques, la physique moderne travaille sur des champs. Ils sont omniprésents. Ce chapitre en introduit un, le champ magnétique, qui permet de comprendre certains principes de fonctionnement des appareils de haute technologie actuelle.

1 Carte de champ

1.1 Les champs en physique

Un **champ** est une grandeur physique qui dépend de la position et du temps. Par exemple, à la surface du globe, on définit le champ de température : en tout point M , $T(M,t)$ est la température à la date t . On définit aussi le champ de pression $P(M,t)$. Ces deux champs sont des champs scalaires, car la température et la pression sont des scalaires, c'est-à-dire qu'ils ne sont définis que par un unique nombre.

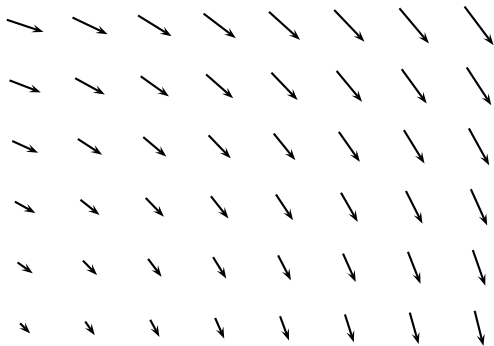


Figure 26.1 – Exemple de champ vectoriel

Un champ vectoriel est, quant à lui, défini par trois nombres, les coordonnées d'un vecteur dans un espace à trois dimensions. Toujours dans le cas météorologique, on définit le champ des vitesses $\vec{v}(M,t)$, qui représente la vitesse de l'atmosphère au point M , à la date t . On

CHAPITRE 26 – LE CHAMP MAGNÉTIQUE

représente graphiquement ce champ vectoriel en traçant, en de nombreux points de l'espace, le vecteur vitesse des particules.

On passe aux **lignes de champ** en traçant les courbes qui sont en tout point tangentes aux vecteurs. Si une ligne de champ montre bien la direction que prend le champ dans l'espace, l'information relative à la norme du vecteur semble perdue. Toutefois, on peut prendre la convention que plus la norme est importante, plus on dessine des lignes de champ serrées, convention qui est appliquée sur la figure 26.2 pour le champ de la figure 26.1.

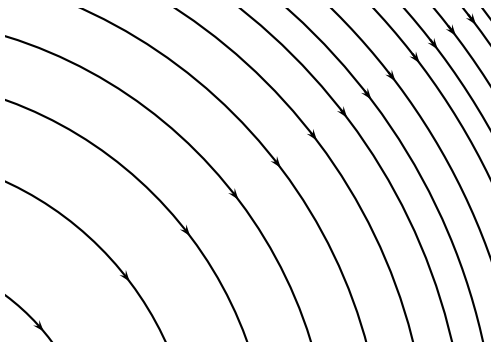


Figure 26.2 – Exemple de lignes de champ vectoriel.

1.2 Un champ vectoriel permet de décrire une interaction à distance

Lorsqu'on approche un aimant d'une pièce de 1, 2 ou 5 centimes d'euro, la pièce est attirée par l'aimant. Mais comment ? Il n'y a rien qui relie les deux objets, aucun fil visible. On utilise le concept de champ pour interpréter cette action à distance : l'aimant crée un **champ magnétique** en tout point de l'espace autour de lui et c'est ce champ qui a une influence sur la pièce métallique.

On interprète de la même l'attraction gravitationnelle avec un champ gravitationnel (par exemple, le Soleil crée dans tout le système solaire un champ gravitationnel qui agit sur les planètes) ou l'interaction électrostatique de deux corps électrisés avec la notion de champ électrique. Dans ces théories les champs (électrique, magnétique, de gravitation) sont les intermédiaires qui transmettent les actions mécaniquement observables. Ils contiennent une certaine énergie, fournie à l'objet qui se met en mouvement. Le calcul de cette énergie relève du programme des classes de seconde année.

1.3 Unités et ordres de grandeur

Quelle est la dimension du champ magnétique ? Raisonnons sur des équations aux unités, avec l'expression de la force magnétique et le lien entre l'intensité du courant et la charge :

$$\left\{ \begin{array}{l} \vec{f} = q \vec{v} \wedge \vec{B} = m \frac{d\vec{v}}{dt} \\ i = \frac{dQ}{dt} \end{array} \right. \quad \text{donc} \quad \left\{ \begin{array}{l} \text{N} = \text{C.m.s}^{-1} [B] = \text{kg.m.s}^{-2} \\ \text{A} = \text{C.s}^{-1} \end{array} \right.$$

Au final, l'unité du champ magnétique est le $\text{kg.s}^{-2}.\text{A}^{-1}$, qu'on nomme usuellement **Tesla**¹, symbolisé par un T majuscule.

L'unité du champ magnétique est le **Tesla**, de symbole T.

Les champs magnétiques existent tout autour de nous, mais leur valeur dépend de leur utilisation et de leur source.

	ordre de grandeur (en T)
champ terrestre	$4,7.10^{-5}$
aimant usuel	0,1 à 1
appareil d'IRM	3

1.4 Topographie du champ magnétique

La topographie du champ magnétique est l'étude de ses lignes de champ. On admettra ici deux propriétés importantes qui sont illustrées par les exemples qui suivent et qui seront justifiées dans le cours de deuxième année.

a) Première propriété

Une manière de créer un champ magnétique est de faire passer un courant, continu ou variable, dans un circuit. Le champ magnétique a des lignes de champ fermées qui tournent autour des fils parcourus par du courant, toujours dans le même sens par rapport au sens du courant dans les fils.

Les lignes de champ du champ magnétique sont, dans la plupart des cas, des courbes fermées qui entourent les fils dans lesquels le courant électrique qui crée le champ passe. L'orientation de la ligne de champ et le sens du courant dans les fils sont liés par la **règle de la main droite** : si l'on met la main droite le long de la ligne de champ orientée de la base des doigts vers le bout des doigts, le pouce donne le sens du courant.



Figure 26.3 – Règle de la main droite.

Cette propriété permet de savoir immédiatement, à l'observation d'une carte des lignes de champ magnétique, en quel lieu il y a un courant qui passe et dans quel sens ce courant est.

1. En l'honneur de Nikola Tesla, 1856 – 1943, ingénieur serbe naturalisé américain, dont les travaux portèrent sur l'électromagnétisme et ses applications.

CHAPITRE 26 – LE CHAMP MAGNÉTIQUE

b) Deuxième propriété

Dans une carte de champ magnétique, si l'on se déplace le long d'une ligne de champ, l'évolution de l'écartement de cette ligne avec les lignes de champ voisines est liée à l'évolution de la norme du champ magnétique. Si la norme du champ magnétique augmente, les lignes de champ voisines se rapprochent et inversement, si la norme du champ magnétique diminue, elles s'écartent. Ainsi, le champ magnétique est le plus intense, là où les lignes de champ sont les plus serrées.

Sur une carte de ligne de champ magnétique, le champ magnétique est le plus intense là où les lignes de champ se resserrent.

1.5 Quelques cartes de champ magnétique

Les champ magnétiques pris pour exemples dans ce paragraphe sont créés par un courant électrique qui circule ou bien un aimant.

a) Spire circulaire de courant

On examine les lignes de champ créé par le circuit le plus simple : une spire circulaire, parcourue par un courant d'intensité constante, dû à un générateur non représenté.

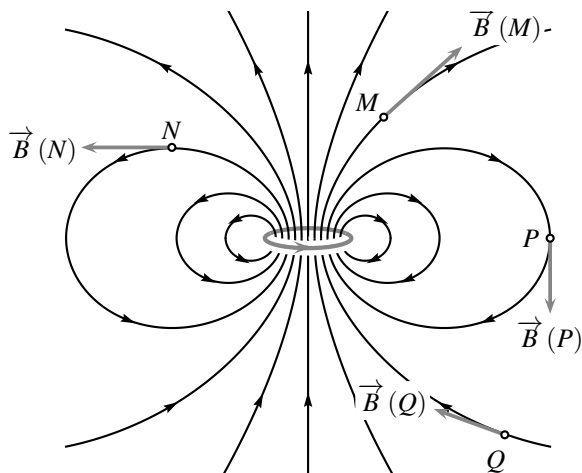


Figure 26.4 – Lignes de champ magnétique (en noir) d'une spire circulaire (en gris) parcourue par un courant dans le sens de la flèche.

On observe que le champ magnétique n'est pas uniforme en tout point de l'espace, les lignes de champ sont resserrées sur la surface de la spire : c'est là que la norme du champ est la plus importante.

D'autre part, on constate que les lignes de champ s'enroulent autour du courant et sont orientées selon la règle de la main droite donnée ci-dessus.

On observe de plus que le champ traverse la spire en respectant une autre **règle de la main droite** :

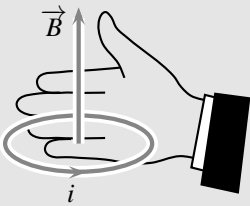


Figure 26.5 – Règle de la main droite.

Si on appose la main sur la spire avec les doigts qui épousent la courbe, le courant allant de la base vers le bout des doigts, le pouce indique alors le sens du vecteur champ magnétique.

Cette propriété est générale et elle permet d’orienter systématiquement le champ magnétique par rapport au courant qui le crée.
D’autre part, la mesure de la norme du champ magnétique en un point pour différentes valeurs de l’intensité i permet d’établir que :

La norme du champ magnétique est proportionnelle à l’intensité du courant qui traverse la spire.

b) Bobine longue

Une bobine est constituée de N de spires jointives, de rayon r , de même axe, alignées sur une distance ℓ , toutes parcourues par le même courant d’intensité i . On définit le nombre de spires par unité de longueur n par : $n = \frac{N}{\ell}$.

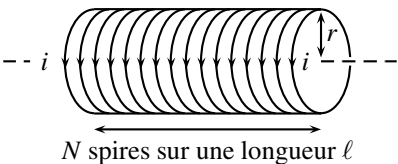


Figure 26.6 – Bobine (chaque spire est parcourue par le même courant).

Les lignes de champ magnétique, créé par la bobine, sont représentées sur la figure 26.7. Elles vérifient la règle de la main droite quant à leur orientation par rapport aux spires de la bobine.

CHAPITRE 26 – LE CHAMP MAGNÉTIQUE

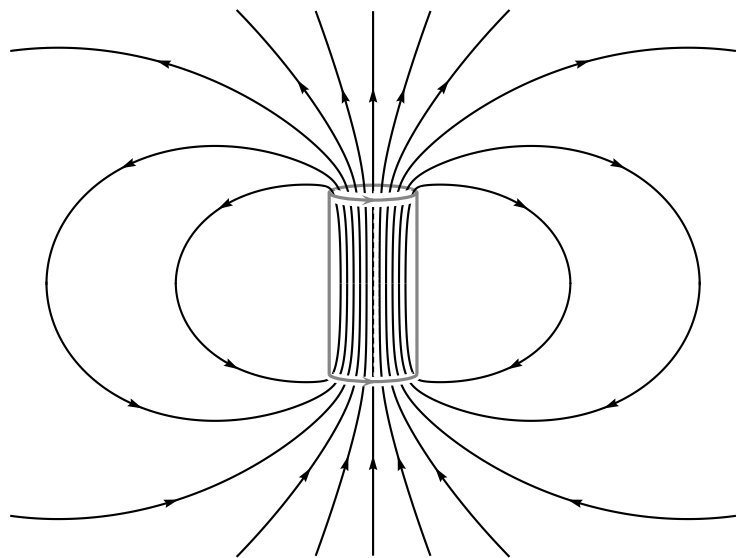


Figure 26.7 – Lignes de champ magnétique, en noir, d'une bobine, en gris (seules les spires extrêmes sont représentées, afin de voir les lignes de champ dans la bobine).

On observe que le champ magnétique est quasi-uniforme à l'intérieur de la bobine. En effet, les lignes de champ restent parallèles, sans se resserrer ni s'éloigner.

À l'intérieur d'une bobine longue, le champ magnétique est quasiment uniforme, parallèle à l'axe de la bobine et sa direction est liée au sens du courant par la règle de la main droite. Sa norme est donnée par :

$$B = \mu_0 n i,$$

où n est le nombre de spires par unité de longueur et $\mu_0 = 4\pi \cdot 10^{-7} \text{ H.m}^{-1}$ la perméabilité du vide.

Par exemple, pour une bobine formée de $N = 1000$ spires sur une longueur $\ell = 10 \text{ cm}$, parcourue par un courant d'intensité $i = 0,5 \text{ A}$, la norme du champ magnétique intérieur est :

$$B = 4\pi \cdot 10^{-7} \times \frac{1000}{10 \cdot 10^{-2}} \times 0,5 = 6,3 \cdot 10^{-3} \text{ T}.$$

Une bobine longue est aussi nommée **solénoïde**, du grec $\sigma\omega\lambda\eta\nu$, *solen*, qui signifie tuyau.

Effets de bords (PTSI) Il faut noter toutefois, que la zone où la norme du champ magnétique est uniforme ne couvre pas l'intégralité de la bobine. Elle n'est valable que loin des bords, là où la courbure des lignes de champ reste imperceptible.

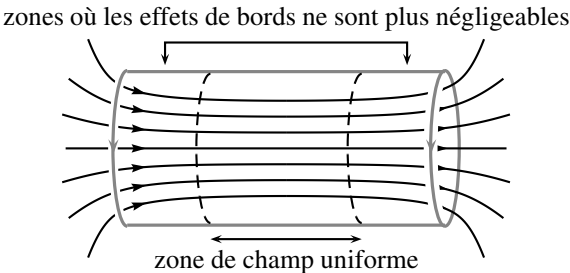


Figure 26.8 – Zoom des lignes de champ magnétique, en noir, d’une bobine, en gris.

On parle d’**effet de bord** pour nommer le phénomène où les lignes de champ s’évasent à proximité des bords de la bobine. On conçoit que plus la bobine est longue, par rapport à son rayon, plus la zone touchée par les effets de bord est étroite. Cet effet est visible sur la figure 26.9, où est tracé la norme du champ magnétique, sur l’axe de la bobine, pour des bobines de plus en plus longues par rapport à leur rayon. Le nombre n de spires par unité de longueur, le rayon r , ainsi que l’intensité i du courant, sont identiques dans chaque bobine ; seule la longueur ℓ varie. $B(x)$ est la norme du champ magnétique à l’abscisse x , repérée par rapport au centre de la bobine.

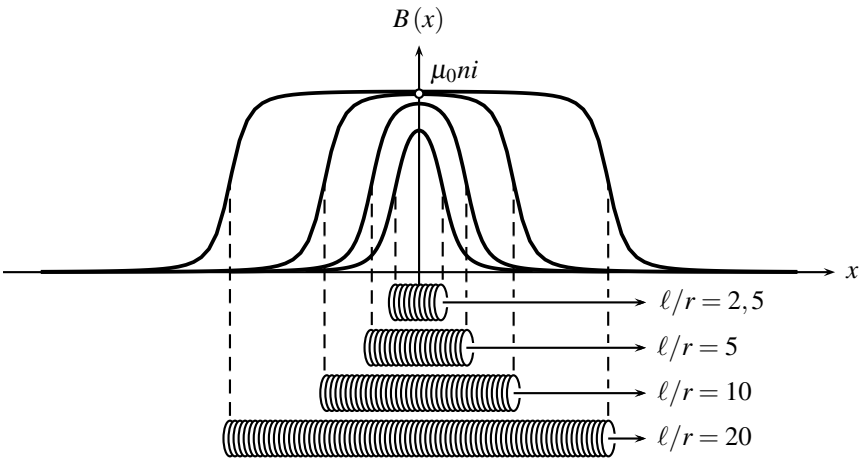


Figure 26.9 – Norme du champ magnétique sur l’axe de diverses bobines : effets de bord.

Pour les deux bobines les plus longues, $\ell/R = 10$ ou 20 , la norme du champ magnétique vaut $\mu_0 n i$ presque partout, sauf proche des bords. Mais pour les deux bobines les plus courtes, $\ell/R = 2,5$ ou 5 , la norme du champ est inférieure à $\mu_0 n i$; la zone où les effets de bord ne sont pas négligeables occupe l’intégralité de la bobine.

CHAPITRE 26 – LE CHAMP MAGNÉTIQUE

Remarque

Le lecteur pourra voir, dans la suite de ses études, que la norme du champ magnétique sur l’axe, avec les notations de ce paragraphe, est donnée par la formule :

$$B(x) = \frac{\mu_0 n i}{2} \left(\frac{x + \frac{\ell}{2}}{\sqrt{r^2 + \left(x + \frac{\ell}{2}\right)^2}} - \frac{x - \frac{\ell}{2}}{\sqrt{r^2 + \left(x - \frac{\ell}{2}\right)^2}} \right).$$

Finalement, on observe sur le graphe de la figure 26.9 que la norme du champ magnétique sur l’axe d’une bobine décroît fortement à l’extérieur de celle-ci. Ce résultat est cohérent avec les lignes de champ, représentées sur la figure 26.7, qui s’écartent nettement au sortir de la bobine.

c) Aimant

La carte de champ d’un aimant est similaire à celle d’une spire ou d’une bobine. Par convention, on appelle pôle nord la face de l’aimant d’où les lignes de champ sortent et pôle sud la face de l’aimant dans laquelle les lignes de champ entrent.

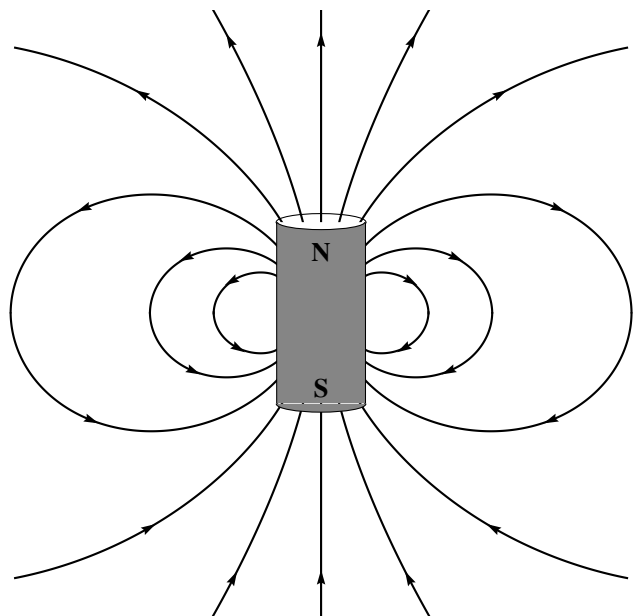


Figure 26.10 – Lignes de champ magnétique d’un aimant. **N** est le pôle nord et **S** le pôle sud.

2 Moment magnétique

Une spire, une bobine longue, un aimant, admettent tous trois des lignes de champ magnétique de même allure à grande distance, c'est-à-dire en un point M tel que, si on note O le centre de la spire, de la bobine ou de l'aimant, la distance OM est très supérieure à la taille de la spire, de la bobine ou de l'aimant. Afin de comparer leur effets magnétiques, on introduit le **moment magnétique**.

2.1 Vecteur surface

Avant de caractériser le moment magnétique, il convient de définir le **vecteur surface** associé à une spire de courant plane. C'est un vecteur normal à la spire, de norme égale à sa surface S . Si la direction du vecteur surface est orthogonale à la spire, il y a deux possibilités pour son sens. Laquelle choisir ? Une fois le sens de parcours de l'intensité du courant explicité, on appose la main sur la courbe, avec les doigts qui épousent la courbe, dans le sens de l'orientation. Le pouce indique alors le sens du vecteur surface. C'est encore une **règle de la main droite**.

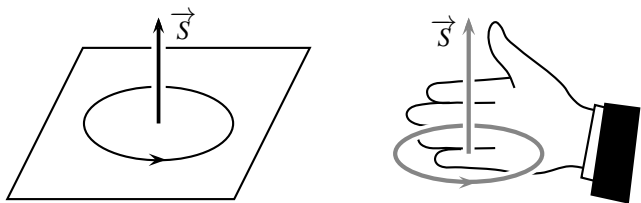


Figure 26.11 – Surface orientée et vecteur surface.

Cette méthode est applicable pour des boucles de courant de toute forme.

Pour une surface orientée, de surface S , le vecteur surface $\vec{S}$ est un vecteur de norme S , orthogonal à la surface, de sens donné par la règle de la main droite.

Remarque

On définit aussi le vecteur normal à la surface $\vec{n}$, par $\vec{S} = S\vec{n}$. Ce vecteur normal est unitaire, c'est à dire que sa norme vaut 1.

2.2 Définition du moment magnétique

Le **moment magnétique** d'une boucle de courant plane, de surface S , parcourue par un courant d'intensité i , est le vecteur $\vec{\mathcal{M}} = i\vec{S}$. Il s'exprime en $A.m^2$.

2.3 Moment magnétique d'un aimant

Attendu que les lignes de champ d'une boucle de courant et d'un aimant sont identiques à grande distance, on étend la notion de moment magnétique aux aimants. Toutefois, ceux-

CHAPITRE 26 – LE CHAMP MAGNÉTIQUE

ci ne sont parcourus par aucun courant interne. Le moment magnétique d'un aimant, bien qu'exprimé en $A.m^2$, ne représente plus le produit de deux termes.

Le moment magnétique d'un aimant dépend de sa taille. L'ordre de grandeur du moment magnétique d'un aimant usuel est d'environ $10 A.m^2$. Le moment magnétique de la Terre vaut $7,9.10^{22} A.m^2$.

2.4 Lignes de champ d'un moment magnétique

Un moment magnétique représente tout aussi bien une boucle de courant qu'un aimant. Il présente de manière unifié le champ magnétique créé à grande distance.

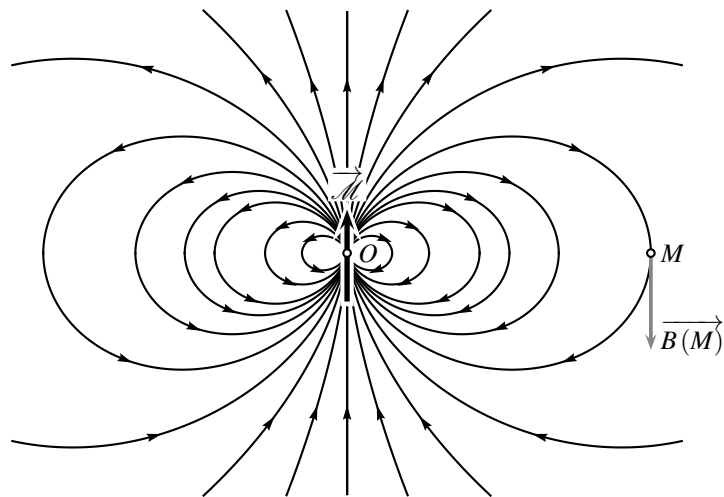


Figure 26.12 – Lignes de champ magnétique d'un moment magnétique.

La distance OM est très supérieure à la largeur caractéristique de la source modélisée par le moment magnétique $\vec{M}$, comme le rayon d'une spire, le rayon ou la longueur d'une bobine longue ou d'un aimant. Ainsi, la source n'est pas représentée, car totalement concentrée au point O .

SYNTHÈSE

SAVOIRS

- allure les cartes de champ d'une spire, d'une bobine longue, d'un aimant
- dispositif réalisant un champ magnétique quasi uniforme
- ordres de grandeurs de champs magnétiques (aimant, IRM, Terre)
- moment magnétique d'une boucle de courant
- ordre de grandeur du moment magnétique d'un aimant

SAVOIR-FAIRE

- identifier sur une carte de champ les zones de champ uniforme, fort ou faible
- identifier sur une carte de champ l'emplacement des sources
- orienter le champ magnétique d'une bobine longue

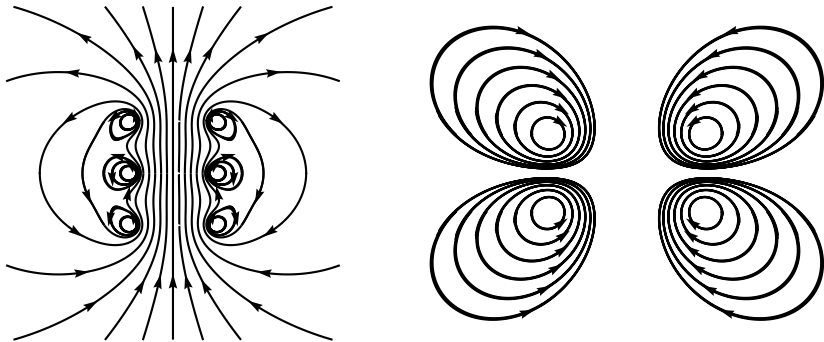
MOTS-CLÉS

- | | | |
|--------------------|------------------------------|---------------------|
| • champ vectoriel | • spire ou boucle de courant | • moment magnétique |
| • ligne de champ | • bobine longue | |
| • champ magnétique | • aimant | |

S'ENTRAÎNER

26.1 Cartes de champ (★)

Dans les cartes de champs magnétique suivantes, où le champ est-il le plus intense ? Où sont placées les sources ? Le courant sort-il ou rentre-t-il du plan de la figure ?



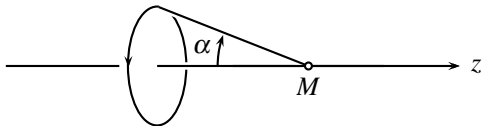
26.2 Champ créé par une bobine longue (★)

On considère une bobine de longueur $L = 60$ cm, de rayon $R = 4$ cm, parcouru par un courant d'intensité $i = 0,6$ A.

1. La formule du champ dans un solénoïde est-elle valable ?
2. Déterminer le nombre de spires nécessaires pour obtenir un champ magnétique de $0,1 \cdot 10^{-2}$ T.
3. La bobine est réalisée en enroulant un fil de 1,5 mm de diamètre autour d'un cylindre en carton. Combien de couches faut-il bobiner pour obtenir le champ précédent ?

26.3 Champ créé par une spire sur son axe (★)

Le champ créé par une spire de courant, parcourue par un courant d'intensité i , de rayon R , est donné, en un point M qui appartient à l'axe de la spire, par la formule : $\vec{B}(M) = \frac{\mu_0 i}{2R} \sin^3 \alpha (\pm \vec{u}_z)$. α est l'angle sous lequel on voit la spire depuis le point M .

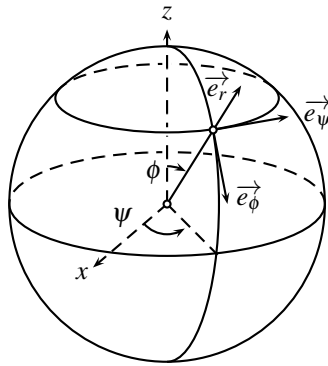


1. Le champ est-il dirigé suivant $\oplus \vec{u}_z$ ou suivant $\ominus \vec{u}_z$?
2. Calculer la norme de $\vec{B}$ en un point de l'axe distant de $L = 10$ cm du centre de la spire. On prendra $R = 2$ cm et $i = 0,5$ A.

26.4 Champ magnétique terrestre (*)

Le champ magnétique terrestre est décrit en première approximation par le champ magnétique d'un dipôle magnétique situé au centre de la terre O , de moment $\vec{\mathcal{M}} = -\mathcal{M}\vec{u}_z$ ($\mathcal{M} = 7,9 \cdot 10^{22}$ A.m² et $\vec{u}_z$ désigne le vecteur unitaire de l'axe géomagnétique de la Terre, qui est légèrement incliné par rapport à l'axe de rotation terrestre). Un point de l'espace est repéré par ses coordonnées sphériques (r, ϕ, ψ) par rapport à l'axe géomagnétique. En un point suffisamment éloigné de O , les composantes de $\vec{B}$ s'écrivent :

$$B_r = -\frac{\mu_0}{4\pi} \mathcal{M} \frac{2 \cos \phi}{r^3}, \quad B_\phi = -\frac{\mu_0}{4\pi} \mathcal{M} \frac{\sin \phi}{r^3} \quad \text{et} \quad B_\psi = 0.$$



Calculer la norme du champ magnétique vers le centre de la France métropolitaine, où $r = 6300$ km et $\phi = 42^\circ$.

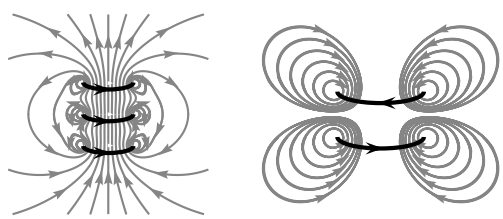
CHAPITRE 26 – LE CHAMP MAGNÉTIQUE

CORRIGÉS

26.1 Cartes de champ

Le champ est le plus intense là où les lignes sont les plus serrées, c'est à dire dans la zone centrale, où les lignes sont verticales, dans le cas de gauche, et proche des sources, dans le cas de droite.

Les lignes de champ tournent autour des sources. Sur la figure de gauche : on a donc trois sources (trois spires). Sur la figure de droite, on a deux spires.



26.2 Champ créé par une bobine longue

1. $\frac{L}{R} = 15$, la formule $B = \mu_0 ni$ est valable.
2. $B = \mu_0 ni = \mu_0 \frac{N}{L} i$ donc $N = 796$ spires.
3. Il s'agit d'un problème d'encombrement spatial. Combien de spires de diamètre $d = 1,5$ mm peut-on enrouler sur une distance $L = 60$ cm ? $L/d = 400$ spires. Il faudra donc 2 couches pour caser les 796 spires.

26.3 Champ créé par une spire sur son axe

1. D'après la règle de la main droite, le champ en M est dirigé suivant $\oplus \vec{u}_z$.
2. Il faut calculer le sinus de l'angle α : $\sin \alpha = \frac{R}{\sqrt{L^2 + R^2}} = 1,96 \cdot 10^{-1}$. Ainsi, en M :
$$B = \frac{4\pi 10^{-7} \times 0,5}{2 \times 2 \cdot 10^{-2}} \times (1,96 \cdot 10^{-1})^3 = 1,2 \cdot 10^{-7} \text{ T.}$$

26.4 Champ magnétique terrestre

La norme du vecteur est $\|\vec{B}\| = \sqrt{B_r^2 + B_\phi^2 + B_\psi^2}$. On la note aussi B pour simplifier. Pour le calcul du sinus et du cosinus, on passe en mode degrés de la calculatrice, ou alors on reste en mode radians, en convertissant les degrés en radians : $\phi = 42^\circ = 42 \times \frac{\pi}{180}$ radian. Ainsi :

$$B = \sqrt{\left(\frac{4\pi 10^{-7}}{4\pi} \times 7,9 \cdot 10^{22} \times \frac{2 \cos(42^\circ)}{(6300 \cdot 10^3)^3}\right)^2 + \left(\frac{4\pi 10^{-7}}{4\pi} \times 7,9 \cdot 10^{22} \times \frac{\sin(42^\circ)}{(6300 \cdot 10^3)^3}\right)^2},$$

soit : $B = 5,1 \cdot 10^{-5} \text{ T.}$

Actions d'un champ magnétique



Un champ magnétique, créé par un circuit parcouru par courant ou un aimant, exerce une force sur un autre circuit ou aimant. Dans ce chapitre on apprendra à calculer la force et ou le couple exercé(e) par un champ magnétique sur un circuit parcouru par un courant.

1 Force de Laplace

1.1 Force de Laplace sur une tige en translation

Une tige $\mathcal{T}$, conductrice, est posée sur deux rails, eux aussi conducteurs, nommés rails de Laplace¹. L'ensemble forme un circuit électrique fermé, parcouru par un courant i , créé par un générateur non représenté. L'ensemble est plongé dans un champ magnétique uniforme et stationnaire $\vec{B} = B\vec{u}_y$, orthogonal au plan des rails.

La force de Laplace $\vec{f}_L$ qui s'exerce sur la tige $\mathcal{T}$ est due à la présence simultanée du courant i et du champ magnétique $\vec{B}$. Pour la calculer, on représente sur un schéma clair le sens positif choisi pour l'intensité i du courant dans le circuit.

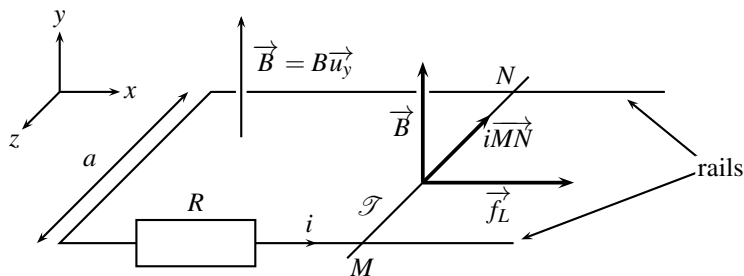


Figure 27.1 – Force de Laplace sur la tige $\mathcal{T}$.

1. En l'honneur de Pierre Simon de Laplace, 1749 – 1827, astronome, physicien et mathématicien français. Il est l'auteur d'une équation célèbre sur le champ de potentiel gravitationnel, de la méthode d'intégration par variation de la constante et d'avancées en calcul de probabilités.

CHAPITRE 27 – ACTIONS D’UN CHAMP MAGNÉTIQUE

La force de Laplace s’exerce sur la portion de la tige $\mathcal{T}$ comprise entre les points M et N , de longueur $a = MN$. Avec le sens positif choisi, c’est-à-dire une tige orientée comme le courant, de M vers N , on admet que la force de Laplace se met sous la forme : $\vec{f}_L = i\overrightarrow{MN} \wedge \vec{B}$.

La force de Laplace exercée sur un tronçon de conducteur rectiligne MN , plongé dans un champ magnétique uniforme extérieur $\vec{B}$ et parcouru par un courant d’intensité i allant de M vers N est :

$$\vec{f}_L = i\overrightarrow{MN} \wedge \vec{B}.$$

1.2 Puissance de la force Laplace

Si la tige $\mathcal{T}$ a un mouvement de translation de vitesse $\vec{v} = v\vec{u}_x$, alors la puissance de la force de Laplace est :

$$\mathcal{P}_L = \vec{f}_L \cdot \vec{v} = iaBv.$$

2 Couple magnétique

2.1 Expression du Couple

Un circuit ou un aimant de moment magnétique $\vec{\mathcal{M}}$ plongé dans une champ magnétique extérieur uniforme $\vec{B}$ subit un **couple magnétique** de moment :

$$\vec{\Gamma}_L = \vec{\mathcal{M}} \wedge \vec{B}.$$

2.2 Puissance de l’action de Laplace

On suppose que la spire $MNPQ$ tourne à la vitesse angulaire $\omega = -\dot{\alpha}$ (le signe $-$ provient du fait que l’angle α part de la direction liée à la spire au lieu d’arriver sur la direction liée à la spire).

Le moment du couple de Laplace par rapport à l’axe (Oy) est $\Gamma_{Ly} = \vec{\Gamma}_L \cdot \vec{u}_y = \mathcal{M}B \sin \alpha$. Alors la puissance de ce couple magnétique est :

$$\mathcal{P}_L = \Gamma_{Ly} \omega = -\mathcal{M}B \dot{\alpha} \sin \alpha.$$

2.3 Complément : établissement du couple

On étudie une spire rectangulaire $MNPQ$ parcourue par un courant i , et plongée dans un champ magnétique uniforme et stationnaire $\vec{B}$, créé par un environnement extérieur. On suppose que la spire peut tourner autour de l’axe (Oy) et on s’intéresse donc au moment par rapport au point O des forces de Laplace exercées par le champ magnétique sur la spire. On appelle $a = NP$ et $b = PQ$ les longueurs des côtés de la spire qui a donc une surface $S = ab$.

Pour minimiser les calculs, on observe sur un schéma les directions des forces de Laplace exercées sur les quatre côtés de la spire (voir figure 27.2).

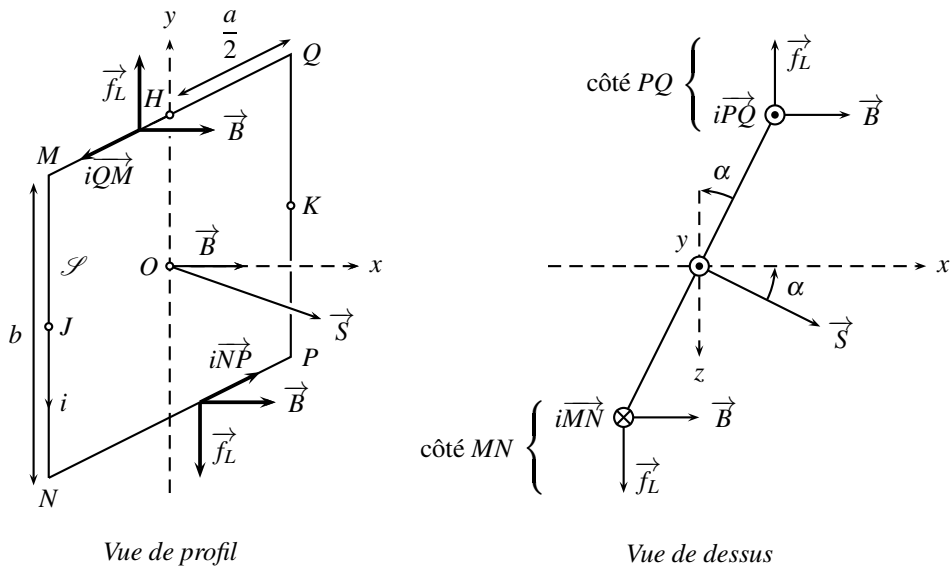


Figure 27.2 – Schéma des moments élémentaires des forces de Laplace.

Le schéma montre que la somme des forces de Laplace sur le tour de la spire est nulle. En effet, les forces s’annulent 2 à 2 pour deux côtés parallèles. Ainsi la spire est soumise à un **couple magnétique**.

Le force de Laplace $\vec{F}_{QM}$ appliquée sur le côté horizontal QM est dirigée suivant $\vec{u}_y$ (voir la vue de profil) et elle s’applique au milieu H de QM . Le moment en O de cette force est, par suite : $\vec{M}_{O,QM} = \vec{OH} \wedge \vec{F}_{QM} = \vec{0}$, car c’est le produit vectoriel de deux vecteurs colinéaires (tous les deux verticaux). Il en est de même pour la force de Laplace $\vec{F}_{PN}$ appliquée sur l’autre côté horizontal PN .

Les forces de Laplace sur les côtés verticaux MN et PQ ont un moment en O non nul et un effet sur la rotation autour de (Oy) . Sur la figure vue de dessus, on observe qu’elles tendent à faire tourner la spire dans le sens trigonométrique, c’est-à-dire qu’elles tendent à aligner le vecteur surface $\vec{S}$ sur le vecteur champ magnétique $\vec{B}$.

Le champ magnétique étant uniforme, la force de Laplace sur le brin PQ est :

$$\vec{F}_{PQ} = i\vec{PQ} \wedge \vec{B} = -ibB\vec{u}_z.$$

Le moment par rapport à O de $\vec{F}_{PQ}$, force s’applique en K , milieu de PQ , est :

$$\vec{M}_{O,PQ} = \vec{OK} \wedge \vec{F}_{PQ} = \frac{a}{2}(\sin \alpha \vec{u}_x - \cos \alpha \vec{u}_z) \wedge (-ibB\vec{u}_z) = \frac{a}{2}ibB \sin \alpha \vec{u}_y,$$

car $\vec{u}_x \wedge \vec{u}_z = -\vec{u}_y$ et $\vec{u}_z \wedge \vec{u}_z = \vec{0}$. On mène le même calcul pour le côté MN :

$$\vec{F}_{MN} = i\vec{MN} \wedge \vec{B} = ibB\vec{u}_z,$$

CHAPITRE 27 – ACTIONS D’UN CHAMP MAGNÉTIQUE

et :

$$\vec{M}_{O,MN} = \vec{OJ} \wedge \vec{F}_{MN} = \frac{a}{2}(-\sin \alpha \vec{u}_x + \cos \alpha \vec{u}_z) \wedge (ibB \vec{u}_z) = \frac{a}{2}ibB \sin \alpha \vec{u}_y.$$

Le moment résultant total de Laplace est donc :

$$\vec{\Gamma}_L = \vec{M}_{O,MN} + \vec{M}_{O,PQ} = aibB \sin \alpha \vec{u}_y.$$

Ce couple est proportionnel à l’intensité i et à la norme du champ magnétique. Il change de sens si le courant i change de sens ou si le champ magnétique change de sens.

L’expression du couple fait apparaître la combinaison iab , produit de la surface $S = ab$ de la spire par l’intensité du courant qui est, au signe près, la norme du moment magnétique de la spire. En observant la figure on peut écrire : $\vec{\mathcal{M}} = i \vec{S} = iab(\cos \alpha \vec{u}_x + \sin \alpha \vec{u}_z)$. Et , et ainsi :

$$\vec{\mathcal{M}} \wedge \vec{B} = iab(\cos \alpha \vec{u}_x + \sin \alpha \vec{u}_z) \wedge (B \vec{u}_x) = iabB \sin \alpha \vec{u}_y = \vec{\Gamma}_L.$$

Remarque

On a bien (voir appendice mathématique) : $\|\vec{\Gamma}_L\| = \|\vec{\mathcal{M}}\| \cdot \|\vec{B}\| \cdot |\sin \alpha|$, où α est l’angle entre le vecteur $\vec{\mathcal{M}}$ et le vecteur $\vec{B}$.

3 Action d’un champ magnétique sur un aimant

3.1 Orientation d’un aimant

Le couple magnétique exercé par un champ magnétique $\vec{B}$ sur un aimant de moment magnétique $\vec{\mathcal{M}}$ tend à aligner le vecteur $\vec{\mathcal{M}}$ sur le vecteur $\vec{B}$. On peut s’en convaincre en considérant le cas d’un aimant pouvant tourner autour d’un pivot d’axe (Oz) représenté sur la figure :

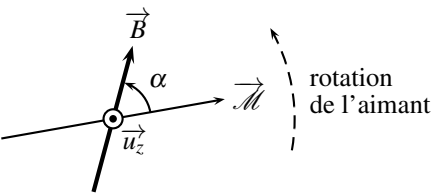


Figure 27.3 – Rotation d’un aimant dans un champ magnétique.

Dans la configuration de la figure, le couple exercée par le champ magnétique sur l’aimant est $\vec{\Gamma}_L = \mathcal{M}B \sin \alpha \vec{u}_z$. Ce couple est porté par $\oplus \vec{u}_z$, il tend donc à faire tourner l’aimant dans le sens positif lié à $\vec{u}_z$, indiqué par la règle de la main droite. L’influence du couple est donc d’aligner l’aimant sur le champ magnétique. Lorsque l’aimant est parallèle au champ magnétique, l’angle α est nul, le couple s’annule, l’aimant ne tourne plus et reste dans cette position.

Le couple magnétique exercé par un champ magnétique extérieur $\vec{B}$ sur un aimant de moment magnétique $\vec{\mathcal{M}}$ tend à aligner le vecteur $\vec{\mathcal{M}}$ sur le vecteur $\vec{B}$.

3.2 Positions d'équilibre

On sait que le produit vectoriel de deux vecteurs est nul si et seulement si les deux vecteurs sont colinéaires. Ainsi, le couple magnétique s'annule pour $\vec{\mathcal{M}}$ parallèle à $\vec{B}$ et de même sens, ou bien $\vec{\mathcal{M}}$ parallèle à $\vec{B}$ et de sens contraire. Dans le premier cas $\alpha = 0$ et dans le deuxième $\alpha = \pi$.

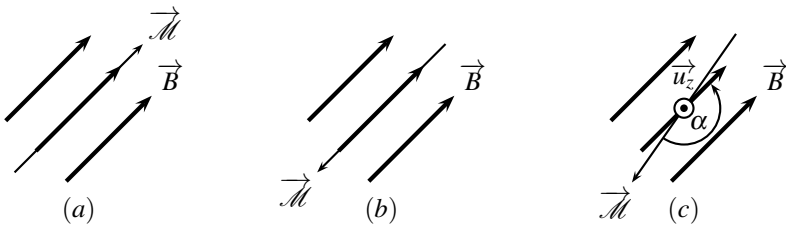


Figure 27.4 – Stabilité de la position d'un aimant dans un champ magnétique, cas (a) parallèle (b) antiparallèle (c) antiparallèle perturbé.

La position antiparallèle n'est pas stable : lorsque l'aimant tourne d'un angle faible, alors il ne revient pas dans sa position d'origine, mais effectue un demi-tour pour arriver en position parallèle. En effet, le couple exercé par le champ magnétique sur l'aimant est $\vec{\Gamma}_L = \mathcal{M} B \sin \alpha \vec{u}_z$. Ainsi, dans le cas du schéma, l'aimant est entraîné en rotation autour de (Oz) , dans le sens positif, jusqu'à ce que le couple s'annule, dans la position parallèle où α est nul.

3.3 Application : la boussole

Le champ magnétique de la Terre se modélise, en première approximation, par un moment magnétique placé au centre de la planète. Il sort de la Terre par le pôle nord magnétique, situé au pôle sud géographique, et entre par le pôle sud, situé au pôle nord géographique (voir figure 27.5).

L'aiguille d'une boussole, constitué d'un petit aimant de moment magnétique $\vec{\mathcal{M}}$ en surface de la planète, s'oriente spontanément sur le champ magnétique terrestre $\vec{B}_T$. Ainsi, le pôle nord de l'aimant indique la direction du sud magnétique de la Terre, c'est à dire le pôle nord géographique. Les géographes parlent de pôle nord magnétique pour indiquer la position du pôle sud de l'aimant qu'est la Terre.

CHAPITRE 27 – ACTIONS D’UN CHAMP MAGNÉTIQUE

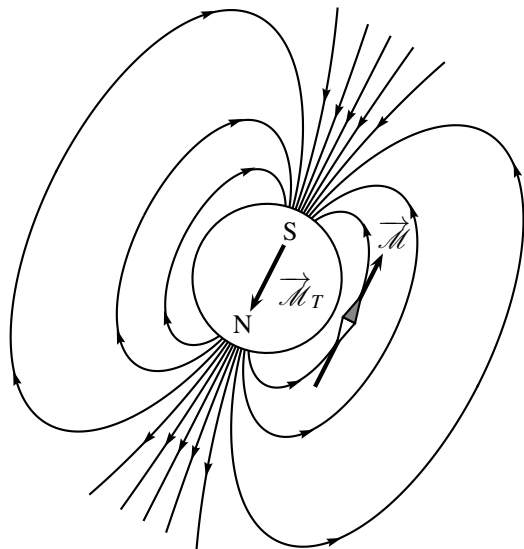


Figure 27.5 – Orientation d’une boussole dans le champ magnétique terrestre de surface.

3.4 Effet moteur d’un champ magnétique tournant

a) Création d’un champ magnétique tournant

On peut obtenir un champ magnétique tournant avec le dispositif représenté sur la figure 27.6. Deux bobines d’axes orthogonaux entre eux, $\mathcal{B}_x$ et $\mathcal{B}_y$, créent chacune un champ magnétique en O . $\mathcal{B}_x$ crée le champ $\vec{B}_{Ox} = Ki_x\vec{u}_x$ et $\mathcal{B}_y$ le champ $\vec{B}_{Oy} = Ki_y\vec{u}_y$, où K est un facteur qui dépend de la géométrie des bobines et de constantes fondamentales. Les champs des deux bobines s’ajoutent.

On fait en sorte que les intensités dans les deux bobines varient sinusoïdalement dans le temps, avec la même amplitude i_0 et la même pulsation ω_0 , et soient en quadrature de phase : $i_x(t) = i_0 \cos(\omega_0 t)$ et $i_y(t) = i_0 \sin(\omega_0 t)$. Le champ total créé par les deux bobines en O est alors :

$$\vec{B}_O = \begin{pmatrix} Ki_x \\ Ki_y \\ 0 \end{pmatrix} = K \begin{pmatrix} i_0 \cos(\omega_0 t) \\ i_0 \sin(\omega_0 t) \\ 0 \end{pmatrix}.$$

C’est un vecteur tournant autour de l’axe (Oz) à la vitesse angulaire ω_0 , comme on le voit en le traçant en fonction du temps.

On crée un champ magnétique tournant avec deux bobines d’axes perpendiculaires, alimentées par des courants en quadrature de phase.

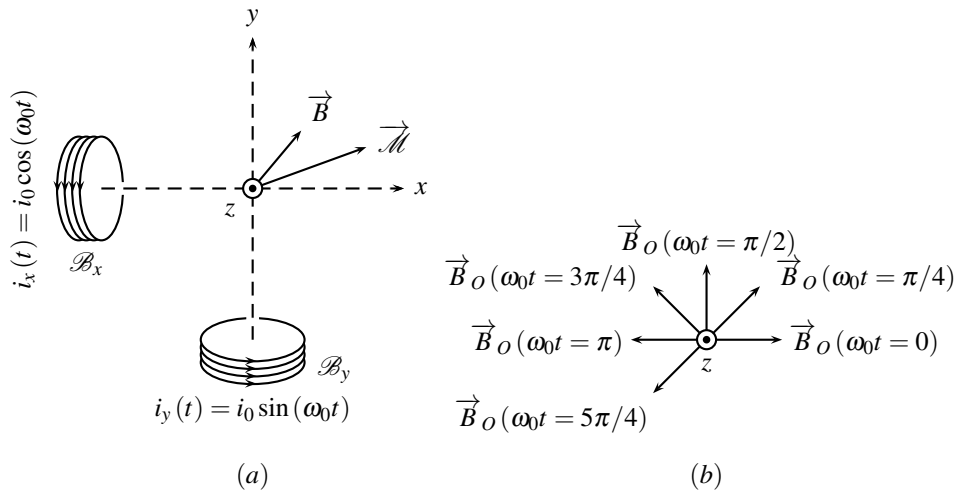


Figure 27.6 – Champ magnétique tournant et moment dipolaire : (a) dispositif expérimental, (b) champ tournant.

b) Action sur un aimant

Qualitativement, lorsqu'on place un aimant en O , mobile en rotation autour de l'axe (Oz) , le champ magnétique exerce un couple sur l'aimant de moment magnétique $\vec{\mathcal{M}}$. Ainsi l'aimant tourne pour s'orienter parallèlement au champ magnétique. Quand celui-ci tourne, celui-là en fait de même. L'aimant tourne donc à la vitesse angulaire ω_0 autour de (Oz) . C'est le principe des moteurs synchrones, dans lesquels le champ magnétique et l'aimant tournent à la même vitesse angulaire.

CHAPITRE 27 – ACTIONS D’UN CHAMP MAGNÉTIQUE

SYNTHÈSE

SAVOIRS

- résultante des force de Laplace pour une tige dans la configuration des rails de Laplace
- couple de Laplace pour une spire de courant en rotation
- action d’un champ magnétique extérieur sur un aimant : positions d’équilibre et de stabilité
- effet moteur d’un champ tournant

SAVOIR-FAIRE

- étudier les positions d’équilibre et de stabilité d’un aimant dans un champ magnétique uniforme
- mettre en œuvre un dispositif expérimental pour étudier l’action d’un champ magnétique uniforme sur une boussole
- créer un champ magnétique tournant à l’aide de 2 ou 3 bobines et mettre en rotation une aiguille aimantée

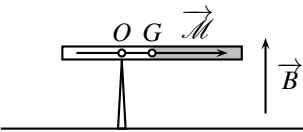
MOTS-CLÉS

- | | | |
|---------------------|--------------------|----------|
| • force de Laplace | • rails de Laplace | • aimant |
| • couple de Laplace | • champ tournant | |

S'ENTRAÎNER

27.1 Aimant en équilibre (★)

Un aimant très fin, de moment magnétique $\vec{\mathcal{M}}$, de masse m , repose en équilibre sur une pointe en O . Il est soumis à l'action d'un champ magnétique uniforme $\vec{B}$ et à la gravité, de direction opposée au champ magnétique.



Évaluer la distance $d = OG$ pour que l'aimant reste en équilibre horizontal.

27.2 Petites oscillations d'un aimant (★)

Un aimant homogène, de moment magnétique $\vec{\mathcal{M}}$, de moment d'inertie J par rapport à son centre de gravité G , est libre de tourner autour de G dans un plan horizontal. Il est soumis à l'action d'un champ magnétique $\vec{B}$ uniforme.

1. L'aimant est légèrement tourné par rapport à sa position d'équilibre, tout en restant de le plan horizontal, puis lâché. Quelle est la période des petites oscillations ultérieures ?
2. Afin d'en déduire la valeur du champ magnétique $\vec{B}$, sans connaître ni le moment d'inertie, ni le moment magnétique de l'aimant, on ajoute au champ $\vec{B}$ un champ magnétique $\vec{B}'$ créé par une bobine longue. On place d'abord la bobine telle que $\vec{B}'$ et le champ $\vec{B}$ soient parallèles et de même sens et on mesure la période τ_1 des petites oscillations de l'aimant. On change ensuite le sens du courant dans la bobine et on mesure la nouvelle valeur τ_2 de la période des petites oscillations.

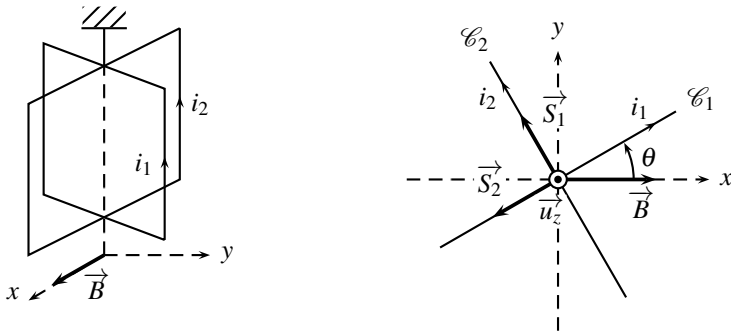
En déduire B en fonction de l'intensité B' du champ créé par la bobine et du rapport τ_1/τ_2 sachant : $B < B'$.

27.3 Deux cadres croisés (★)

Deux cadres rectangulaires $\mathcal{C}_1$ et $\mathcal{C}_2$, identiques et solidaires, de surface S , dont les plans forment un angle droit, sont suspendus au bout d'un fil attaché au bâti qui constitue l'axe Oz . Ils sont mobiles en rotation autour de l'axe vertical (Oz). Les cadres sont parcourus par des courants d'intensités constantes i_1 et i_2 . Il n'y a aucun contact électrique entre les cadres, leur courants ne se mélangent pas.

Ils sont placés dans un champ magnétique uniforme et constant $\vec{B} = B\vec{u}_x$ horizontal.

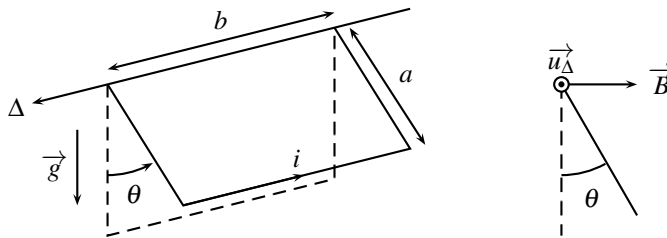
CHAPITRE 27 – ACTIONS D'UN CHAMP MAGNÉTIQUE



Établir l'expression du rapport i_1/i_2 en fonction de l'angle θ , angle entre le plan du cadre parcouru par i_1 et le plan (xOy) .

27.4 Action magnétique sur un cadre (1) (★)

Un cadre conducteur tourne sans frottement autour de l'axe Δ . Il est composé de 4 segments, 2 de longueur a , 2 de longueur b . La masse totale du cadre est m , son moment d'inertie par rapport à Δ est J . Un dispositif, non représenté sur la figure, impose une intensité du courant i constante dans le cadre.



Le cadre est placé dans un champ de pesanteur et un champ magnétique. Le champ magnétique est horizontal, placé dans un plan perpendiculaire à l'axe Δ .

1. Quelle est la position d'équilibre θ_0 ?
2. On écarte légèrement le cadre de sa position d'équilibre. Quelle est la pulsation des petites oscillations alors observées ? On répondra en fonction de J, a, b, i, B, m et g .

APPROFONDIR

27.5 Action magnétique sur un cadre (2) (★★)

On reprend le cadre de l'exercice précédent. Mais cette fois, le champ magnétique est vertical, de sens opposé à celui de $\vec{g}$.

1. Quelle est la position d'équilibre θ_0 ?
2. On écarte légèrement le cadre de sa position d'équilibre. Quelle est la pulsation des petites oscillations alors observées autour de θ_0 ? On répondra en fonction de $\sin(\theta_0), \cos(\theta_0), J, a, b, m, g, i$ et B .

CORRIGÉS

27.1 Aimant en équilibre

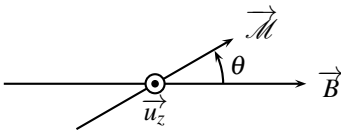
L'aimant est soumis à trois forces : la force magnétique de moment $\vec{\Gamma}_L = \vec{\mathcal{M}} \wedge \vec{B}$, la pesanteur de moment en O $\vec{\Gamma}_0 = \vec{OG} \wedge m \vec{g}$ et la force de réaction $\vec{R}$ de la pointe, de moment nul en O .

On construit l'axe Ox , orthogonal au plan de la feuille, dirigé vers le lecteur. Alors, en position horizontale : $\vec{\Gamma}_L = \mathcal{M} B \vec{u}_x$ et $\vec{\Gamma}_0 = -dm g \vec{u}_x$.

D'après le théorème du moment cinétique, l'aimant reste immobile si : $\vec{\Gamma}_L + \vec{\Gamma}_0 = \vec{0}$ soit $d = \frac{\mathcal{M} B}{mg}$.

27.2 Petites oscillations d'un aimant

1. On note θ l'angle que fait le moment magnétique de l'aimant par rapport au champ magnétique et (Oz) l'axe vertical autour duquel l'aimant tourne.



L'aimant est soumis à la réaction du support, qui s'exerce en O , donc de moment nul en O , au poids, de moment nul pour la même raison, et à la force magnétique de moment $\vec{\Gamma} = \vec{\mathcal{M}} \wedge \vec{B}$:

$$\vec{\Gamma} = \mathcal{M} B \sin(-\theta) \vec{u}_z = -\mathcal{M} B \sin \theta \vec{u}_z.$$

Attention au signe de l'angle dans le sinus : pour calculer le produit vectoriel, on part de $\vec{\mathcal{M}}$ pour aller vers $\vec{B}$, on parcourt l'angle $-\theta$.

Le théorème du moment cinétique, appliqué à l'aimant en O , en projection sur l'axe (Oz) , mène à : $J \frac{d^2 \theta}{dt^2} = -\mathcal{M} B \sin \theta$. Dans le cas de petites oscillations, on linéarise le sinus. Au deuxième ordre en θ , $\sin(\theta) = \theta$ et donc : $\frac{J}{\mathcal{M} B} \frac{d^2 \theta}{dt^2} + \theta(t) = 0$. La pulsation caractéristique du système est telle que $\frac{1}{\omega_0^2} = \frac{J}{\mathcal{M} B}$. D'où la période $T : T = \frac{2\pi}{\omega_0} = 2\pi \sqrt{\frac{J}{\mathcal{M} B}}$.

2. Dans la première expérience, on ajoute $\vec{B}'$ à $\vec{B}$: $\tau_1 = 2\pi \sqrt{\frac{J}{\mathcal{M} (B+B')}} ;$ dans la seconde, le courant est opposé, on retranche donc $\vec{B}'$ à $\vec{B}$. Comme $B < B' : \tau_2 = 2\pi \sqrt{\frac{J}{\mathcal{M} (B'-B)}}$.

CHAPITRE 27 – ACTIONS D'UN CHAMP MAGNÉTIQUE

De ces deux équations : $B = B' \frac{1 - \left(\frac{\tau_1}{\tau_2}\right)^2}{1 + \left(\frac{\tau_1}{\tau_2}\right)^2}$.

27.3 Deux cadres croisés

La figure indique clairement les orientations des deux cadres et les vecteurs surface, orientés d'après la règle de la main droite. L'ensemble des deux cadres est soumis aux deux couples de Laplace qui s'exercent sur chacun des cadres. Les moments magnétiques des cadres sont :

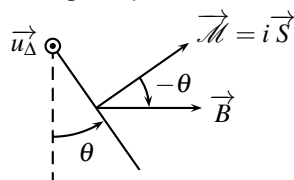
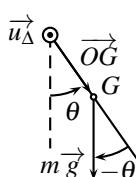
$$\vec{\mathcal{M}}_1 = i_1 S \begin{pmatrix} -\sin \theta \\ \cos \theta \\ 0 \end{pmatrix} \quad \text{et} \quad \vec{\mathcal{M}}_2 = i_2 S \begin{pmatrix} -\cos \theta \\ -\sin \theta \\ 0 \end{pmatrix}.$$

Les couples de Laplace sont : $\vec{\Gamma}_1 = \vec{\mathcal{M}}_1 \wedge \vec{B} = i_1 S B \begin{pmatrix} 0 \\ 0 \\ -\cos \theta \end{pmatrix}$ et $\vec{\Gamma}_2 = \vec{\mathcal{M}}_2 \wedge \vec{B} = i_2 S B \begin{pmatrix} 0 \\ 0 \\ \sin \theta \end{pmatrix}$.

À l'équilibre, le théorème du moment cinétique appliqué à l'ensemble de cadres mène à : $\vec{\Gamma}_1 + \vec{\Gamma}_2 = \vec{0}$ soit $\tan \theta = \frac{i_1}{i_2}$.

27.4 Action magnétique sur un cadre (1)

1. Le cadre est soumis à son poids $m \vec{g}$, à la force de Laplace $\vec{f}_L$ et la réaction du support.



Le moment des forces de Laplace, par rapport à l'axe de rotation Δ , est : $\vec{\Gamma}_L = \vec{\mathcal{M}} \wedge \vec{B} = i \vec{S} \wedge \vec{B} = iabB \sin(-\theta) \vec{u}_\Delta = -iabB \sin \theta \vec{u}_\Delta$.

Soit G le centre de gravité du cadre, en son centre. Le moment du poids par rapport à O est : $\vec{\Gamma}_p = \vec{OG} \wedge m \vec{g} = \frac{a}{2} mg \sin(-\theta) \vec{u}_\Delta = -\frac{a}{2} mg \sin \theta \vec{u}_\Delta$.

La liaison pivot est réputée parfaite, aucun moment n'est à considérer.

Dans le cas de l'équilibre, le théorème du moment cinétique, projeté sur l'axe Δ , mène à : $-iabB \sin \theta_0 - \frac{a}{2} mg \sin \theta_0 = 0$, d'où $\theta_0 = 0$.

2. Le théorème du moment cinétique, projeté sur l'axe Δ , mène à :

$$J \frac{d^2 \theta}{dt^2} = -iabB \sin \theta - \frac{a}{2} mg \sin \theta.$$

Dans le cas d'oscillations de faible amplitude, on linéarise l'équation. Au deuxième ordre en θ , $\sin \theta = \theta$ et l'équation devient : $J \frac{d^2 \theta}{dt^2} + a \left(ibB + \frac{mg}{2} \right) \theta(t) = 0$.

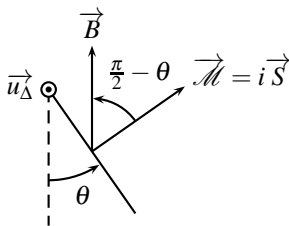
On l'écrit sous forme canonique afin d'identifier la pulsation ω_0 :

$$\frac{J}{a \left(ibB + \frac{mg}{2} \right)} \frac{d^2 \theta}{dt^2} + \theta(t) = 0 \quad \text{d'où} \quad \frac{1}{\omega_0^2} = \frac{J}{a \left(ibB + \frac{mg}{2} \right)}.$$

On en déduit la période des petites oscillations : $T = \frac{2\pi}{\omega_0} = 2\pi \sqrt{\frac{J}{a \left(ibB + \frac{mg}{2} \right)}}.$

27.5 Action magnétique sur un cadre (2)

1. Par rapport à l'exercice précédent, seul change le moment des actions de Laplace par rapport à l'axe de rotation Δ : $\vec{\Gamma}_L = \vec{\mathcal{M}} \wedge \vec{B} = i \vec{S} \wedge \vec{B} = iabB \sin \left(\frac{\pi}{2} - \theta \right) \vec{u}_\Delta = iabB \cos \theta \vec{u}_\Delta$.



Soit G le centre de gravité du cadre, en son centre. Le moment du poids par rapport à O est : $\vec{\Gamma}_p = \vec{OG} \wedge m \vec{g} = \frac{a}{2} mg \sin(-\theta) \vec{u}_\Delta = -\frac{a}{2} mg \sin \theta \vec{u}_\Delta$.

Dans le cas de l'équilibre le théorème du moment cinétique, projeté sur l'axe Δ , mène à :

$$iabB \cos \theta_0 - \frac{a}{2} mg \sin \theta_0 = 0 \quad \text{d'où} \quad \tan \theta_0 = \frac{2ibB}{mg}.$$

2. Le théorème du moment cinétique, appliqué au cadre, projeté sur l'axe Δ , mène à :

$$J \frac{d^2 \theta}{dt^2} = iabB \cos \theta - \frac{a}{2} mg \sin \theta.$$

Le cadre oscille autour de sa position d'équilibre θ_0 . Soit ε l'angle que fait le cadre par rapport à θ_0 , c'est à dire : $\theta = \theta_0 + \varepsilon$, où $|\varepsilon| \ll \theta_0$.

Effectuons un développement limité de $\cos(\theta_0 + \varepsilon)$ et $\sin(\theta_0 + \varepsilon)$, autour de θ_0 , au premier ordre en ε . Attendu que $f(x_0 + \varepsilon) = f(x_0) + \varepsilon f'(x_0) + o(\varepsilon)$: $\cos(\theta_0 + \varepsilon) = \cos(\theta_0) - \varepsilon \sin(\theta_0) + o(\varepsilon)$ et $\sin(\theta_0 + \varepsilon) = \sin(\theta_0) + \varepsilon \cos(\theta_0) + o(\varepsilon)$.

Avec $\frac{d^2 \theta}{dt^2} = \frac{d^2}{dt^2} (\theta_0 + \varepsilon) = \frac{d^2 \varepsilon}{dt^2}$, le théorème du moment dynamique devient :

$$J \frac{d^2 \varepsilon}{dt^2} = iabB \cos \theta_0 - \frac{a}{2} mg \sin \theta_0 - iabB \sin \theta_0 \varepsilon - \frac{a}{2} mg \cos \theta_0 \varepsilon.$$

CHAPITRE 27 – ACTIONS D'UN CHAMP MAGNÉTIQUE

Les termes constants s'annulent (ils définissent la position d'équilibre) et il reste :

$$J \frac{d^2 \varepsilon}{dt^2} + a \left(ibB \sin \theta_0 + \frac{mg}{2} \cos \theta_0 \right) \varepsilon = 0, \text{ qui s'écrit sous forme canonique :}$$

$$\frac{J}{a \left(ibB \sin \theta_0 + \frac{mg}{2} \cos \theta_0 \right)} \frac{d^2 \varepsilon}{dt^2} + \varepsilon = 0.$$

D'où $\frac{1}{\omega_0^2} = \frac{J}{a \left(ibB \sin \theta_0 + \frac{mg}{2} \cos \theta_0 \right)}$ puis la période des petites oscillations :

$$T = \frac{2\pi}{\omega_0} = 2\pi \sqrt{\frac{J}{a \left(ibB \sin \theta_0 + \frac{mg}{2} \cos \theta_0 \right)}}.$$

Lois de l'induction

28

La technologie actuelle des machines électriques, des moteurs électriques aux microphones, repose principalement la **loi de Faraday**. Le but de ce chapitre est d'expliquer cette loi expérimentale, afin d'en étudier les applications dans les deux chapitres ultérieurs.

1 Flux magnétique

1.1 Définition du flux magnétique

On considère un champ magnétique $\vec{B}$ uniforme et une spire rectangulaire, de surface S , située dans un plan orthogonal au champ. Quelle est la « quantité » de champ magnétique qui traverse la spire ? On choisit de nommer cette grandeur flux du champ magnétique, ou plus simplement **flux magnétique** ; on lui donne pour valeur $B \times S$, produit de la norme du champ par la surface de la spire offerte à ce champ.

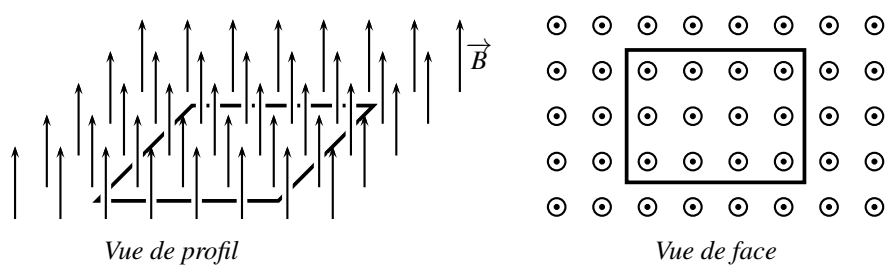


Figure 28.1 – Flux du champ magnétique à travers une surface.

La situation se complique lorsque la spire n'est plus orthogonale au champ, mais inclinée, comme dans le cas de la surface grisée dans le schéma de la page suivante. On voit immédiatement que si l'angle β est nul, alors le champ magnétique passe « au-dessus » et « en-dessous » de la surface grisée, mais pas à travers. Le flux magnétique est nul car la surface offerte au champ est nulle.

On est donc amené à évaluer la surface offerte au champ magnétique. Sur le schéma précédent, cette surface offerte est celle qui est orthogonale au champ, elle mesure $S \times \sin \beta$ ou

CHAPITRE 28 – LOIS DE L'INDUCTION

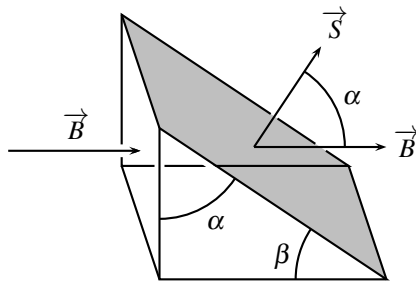


Figure 28.2 – Flux du champ magnétique à travers une surface inclinée (seuls deux vecteurs champ magnétique sont représentés, afin d'alléger la figure).

$S \times \cos \alpha$. Le flux magnétique est alors $B \times S \cos \alpha$.

Pour exprimer le flux par une formule générale on introduit un **vecteur surface** $\vec{S}$ qui est :

- normal à la surface à travers laquelle on calcule le flux magnétique,
- de norme égale à l'aire S de cette surface.

Il faut remarquer qu'il y a deux possibilités pour le vecteur $\vec{S}$, c'est-à-dire deux orientations possibles de la surface. Sur la figure, on a choisi celle pour laquelle les vecteur $\vec{B}$ et $\vec{S}$ sont dans le même sens. L'angle entre $\vec{S}$ et $\vec{B}$ est α de sorte que : $B \times S \cos \alpha = \vec{B} \cdot \vec{S}$ (voir appendice mathématique).

La définition générale du flux magnétique est donc la suivante :

Le **flux magnétique** φ , à travers une surface orientée de vecteur surface $\vec{S}$, est :

$$\varphi = \vec{B} \cdot \vec{S}.$$

Le flux magnétique peut être positif ou négatif. C'est une grandeur algébrique dont le signe dépend du choix d'orientation de la surface.

1.2 Orientation d'une surface

Le choix d'un vecteur surface, dans un sens ou l'autre, est un choix arbitraire, qu'il faut toujours préciser sur un schéma avant de commencer les calculs.

Dans le cas d'une surface délimitée par un contour plan, ce choix est lié au choix d'un sens de parcours le long du contour. Prenons une surface plane et traçons une courbe fermée dessus, par exemple un cercle, comme sur la figure 28.3. Orienter la surface revient à choisir un sens de parcours positif sur cette courbe fermée. Une fois le sens de parcours choisi, on utilise la **règle de la main droite**. Le pouce indique alors le sens du vecteur surface quand le sens positif sur le contour va de la base des doigts vers leur extrémité.

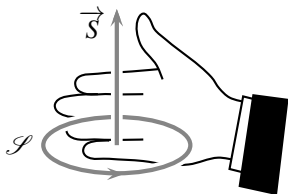


Figure 28.3 – Surface orientée, vue de profil, et vecteur surface.

Remarque

L’orientation de la surface, donc l’orientation du vecteur surface, permet de compter positivement le flux magnétique quand le champ magnétique traverse la surface dans la même direction que $\vec{S}$, négativement dans le cas contraire.

1.3 Unité de flux magnétique

L’unité de flux champ magnétique est le **weber**¹, de symbole Wb. Comment exprimer les webers dans les unités de base du Système International ? Dans le chapitre sur le champ magnétique, on montre que le tesla, unité du champ magnétique, est identique à des $\text{kg.s}^{-2}.\text{A}^{-1}$. Ainsi :

$$\varphi = \vec{B} \cdot \vec{S} \quad \text{donc} \quad \text{Wb} = \text{T.m}^2 = \text{m}^2.\text{kg.s}^{-2}.\text{A}^{-1}.$$

2 Expériences d’induction électromagnétique

2.1 Expérience historique de Faraday

En 1831, Faraday², se demanda si un champ magnétique pouvait être à l’origine d’un courant électrique, puisque la réciproque avait été mise en évidence quelques années plus tôt. Il réalisa l’expérience suivante :

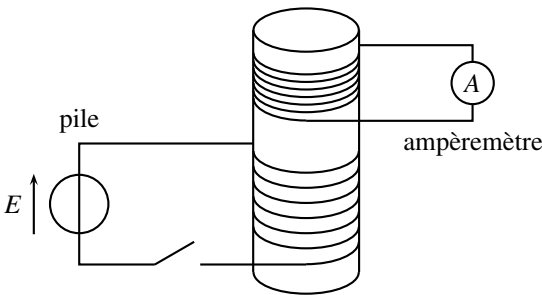


Figure 28.4 – Expérience de Faraday.

1. En l’honneur de Wilhelm Eduard Weber, 1804 – 1891, physicien allemand, professeur et directeur de l’observatoire de l’université de Göttingen, connu pour ses travaux sur l’électrodynamique.
2. Michael Faraday, 1791 – 1867, physicien anglais, célèbre pour ses travaux fondateurs en électromagnétisme. Il fut membre de la Royal Society, où il fonda les leçons du vendredi, toujours en vigueur. Il contribua aussi à l’étude de l’électrochimie, et fut professeur de chimie à l’académie militaire de Wollwich.

CHAPITRE 28 – LOIS DE L'INDUCTION

Deux enroulements de fils sur un cylindre en bois étaient reliés l'un à une pile par l'intermédiaire d'un interrupteur, l'autre à un ampèremètre. Dans un premier temps, Faraday espérait mesurer un courant avec l'ampèremètre en alimentant le premier enroulement pour créer un champ magnétique. Quand l'interrupteur était fermé, il n'observait rien. Il se rendit compte cependant que l'aiguille de l'ampèremètre déviait dans un sens, de manière très brève, à la fermeture de l'interrupteur et dans l'autre sens à l'ouverture. Pour amplifier le phénomène, il utilisa une batterie de plusieurs piles et remplaça le cylindre en bois par un cylindre en fer doux. Il arriva à la conclusion que c'était la **variation** du courant dans le premier circuit qui était à l'origine du courant électrique détecté dans le second. Il publia ses résultats en 1831. Pour l'anecdote, le physicien américain, Joseph Henry, avait fait presque la même expérience un an plus tôt mais n'avait pas publié ses résultats.

Faraday utilisa cette découverte pour construire le premier générateur électrique de courant alternatif, qui fut l'une des plus importantes innovations de l'époque.

2.2 Expériences avec un aimant et une bobine

On effectue l'expérience suivante : une bobine est reliée à un ampèremètre. On approche un aimant, le pôle nord de celui-ci avançant vers la bobine. On constate l'apparition d'un courant dans la bobine dans le sens indiqué sur la figure suivante ; ce courant est appelé courant induit. Quand on retire l'aimant, le sens du courant s'inverse. Si maintenant on avance le pôle sud de l'aimant vers la bobine, le sens du courant est le même que si on éloigne le pôle nord. On observe que l'intensité du courant qui apparaît dans la bobine est d'autant plus grande que l'aimant avance vite.

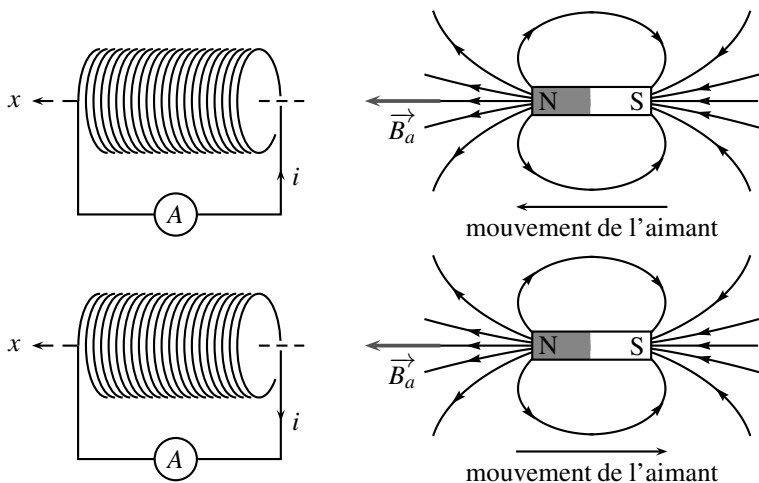


Figure 28.5 – Courant induit par un aimant en déplacement.

Si on déplace la bobine au lieu de l'aimant, les résultats sont identiques : le sens du courant est le même que l'on approche la bobine de l'aimant fixe, ou que l'on approche l'aimant de la bobine fixe.

Si on remplace l'aimant par une bobine alimentée en courant continu, on observe encore un courant dans l'autre bobine non alimentée. Le courant induit dans la deuxième bobine est de sens opposé au courant dans la première bobine quand on rapproche les deux bobines et de même sens quand on les éloigne.

Remarque

On peut réaliser toutes ces expériences en mettant la bobine en circuit ouvert. Dans ce cas, on la relie à un voltmètre ou à un oscilloscope. On mesure alors au lieu d'un courant induit une tension induite entre les bornes de la bobine. Le signe de cette tension dépend de l'orientation de l'aimant et de son sens de déplacement de la même manière que le sens du courant induit.

2.3 Le phénomène d'induction électromagnétique

Ces expériences mettent en évidence le phénomène d'**induction électromagnétique** qui se manifeste par l'apparition d'un courant dans un circuit fermé ou d'une tension aux bornes d'un circuit ouvert, sans qu'il y ait de générateur à l'intérieur de ces circuits. Une condition pour voir ce phénomène est que *le champ magnétique « traversant le circuit » varie*. Cette variation peut provenir, soit d'une variation du champ magnétique, soit d'un déplacement du circuit dans un champ magnétique non uniforme. Les deux cas sont illustrés dans les expériences précédentes.

Il peut y avoir un phénomène d'induction dans les deux cas suivants :

1. Le circuit est fixe dans un champ magnétique qui dépend du temps ;
2. Le circuit est en mouvement dans un champ magnétique stationnaire.

2.4 Loi de Lenz

On sait que le champ magnétique d'un aimant sort par le pôle nord, comme le montre la figure 28.5. Il est plus intense près des pôles de l'aimant, là où les lignes de champ se resserrent. Lorsque le pôle nord de l'aimant s'approche de la bobine, le champ magnétique vu par celle-ci augmente suivant $\oplus \vec{u}_x$. En effet, on approche les zones de fort champ de la bobine. Un courant induit apparaît dans le sens indiqué sur la figure ; ce courant crée, d'après la règle de la main droite, un champ magnétique porté par $\ominus \vec{u}_x$, afin de diminuer le champ dans la bobine qui augmente.

Lorsque le pôle nord de l'aimant s'éloigne de la bobine, le champ magnétique vu par celle-ci suivant $\oplus \vec{u}_x$ diminue. Le courant induit crée, d'après la règle de la main droite, un champ magnétique porté par $\oplus \vec{u}_x$, afin de renforcer le champ dans la bobine qui diminue. Dans chaque cas, le courant induit crée un champ magnétique qui s'oppose à la variation de celui de l'aimant. Cette loi expérimentale est connue sous le nom de loi de Lenz³.

D'après la **loi de Lenz**, les phénomènes d'induction s'opposent, par leurs effets, aux causes qui leur ont donné naissance.

3. En l'honneur de son découvreur, Heinrich Friedrich Emil Lenz, 1804 – 1865, physicien balte d'origine allemande, sujet russe, professeur puis recteur à l'université de Saint Pétersbourg.

CHAPITRE 28 – LOIS DE L'INDUCTION

3 Loi de Faraday

3.1 Règle du flux

Dans toutes les expériences décrites plus haut, une *variation du flux magnétique* φ à travers le circuit provoque l'apparition d'un courant dans celui-ci.

En 1831, Faraday déduisit de ses expériences la **loi de Faraday** : le courant induit dans le circuit est égal à celui que produirait un générateur fictif, dit *générateur induit*, dont la force électromotrice e est donnée par la formule :

$$e = - \frac{d\varphi}{dt}.$$

e est appelée **force électromotrice induite**, en abrégé f.é.m. induite.

3.2 Convention d'algébrisation

Si l'on veut étudier un circuit électrique qui est le siège d'un phénomène d'induction électromagnétique, il faut ajouter dans le schéma électrocinétique le générateur induit. Dans quel sens doit-on placer la flèche de ce générateur ? La réponse est donnée par le schéma ci-dessous : la flèche du générateur induit doit être mise dans le sens positif conventionnel pour le courant i dans le circuit.

La f.é.m. induite e est comptée positive dans le sens conventionnel positif du courant.

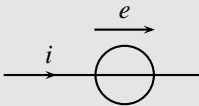


Figure 28.6 – Convention d'algébrisation générateur induit.

Ceci correspond à une **convention générateur**.

3.3 Exceptions à la règle du flux

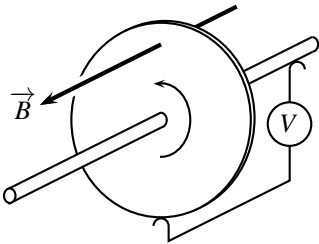


Figure 28.7 – Roue de Barlow.

On considère tout d’abord le cas d’un circuit dans lequel une f.é.m. est induite, mais sans variation de flux : la roue de Barlow. Un disque de surface S tourne autour de son axe. L’ensemble est en métal conducteur et plongé dans un champ magnétique uniforme et stationnaire. Deux contacts glissants relient le centre du disque à sa périphérie pour mesurer, avec un voltmètre, la f.é.m. induite.

Le flux du champ magnétique $\vec{B}$ à travers la roue de surface S est simplement BS . Ce flux reste constant lors du mouvement de la roue, alors qu’on observe expérimentalement une f.é.m. différente de 0. La loi de Faraday n’est ici pas applicable.

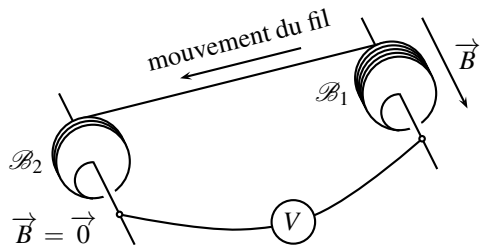


Figure 28.8 – Expérience d’André Blondel.

On considère maintenant le cas, assez rare, d’un circuit où nulle f.é.m. n’est induite alors que le flux du champ magnétique varie au cours du temps : l’expérience d’André Blondel⁴. Deux bobines, B_1 et B_2 , tournent autour de leurs axes, qui restent parallèles. On déroule une bobine afin d’enrouler le fil électrique sur l’autre, comme le montre le schéma. Deux contacts, placés sur les axes des spires, permettent de mesurer la f.é.m. induite au moyen d’un voltmètre. La bobine B_1 est placée dans une région où règne un champ magnétique uniforme et stationnaire, parallèle à son axe. La bobine B_2 est dans une zone de l’espace sans aucun champ magnétique.

Si l’on note n_1 le nombre de spires de la bobine B_1 , alors n_1 diminue au cours du temps lorsqu’on la déroule. Le flux du champ magnétique $\vec{B}$ à travers les n_1 spires de la bobine B_1 , chacune de surface S , vaut n_1BS . Bien que ce flux varie au cours du temps, nulle f.é.m. n’est expérimentalement induite.

Dans ce cas, le circuit électrique ne coupe aucune ligne de champ magnétique. Lors de son mouvement, il passe « entre les lignes de $\vec{B}$ ».

3.4 Loi de Faraday

Compte tenu des exceptions développées dans le précédent paragraphe, on admet la formulation générale de la loi de Faraday.

4. 1863 – 1938, ancien élève de l’École polytechnique, corps des Ponts, il est l’auteur d’avancées fondamentales dans le fonctionnement des machines synchrones (MS) et leur couplage au réseau. Le diagramme de Blondel, méthode d’étude de l’induit d’une MS à pôles saillants avec les vecteurs de Fresnel, est toujours couramment utilisé.

CHAPITRE 28 – LOIS DE L'INDUCTION

Dans le cas d'un **circuit mobile dans un champ magnétique stationnaire**, si, à la fois :

- on peut définir un flux variable $\varphi(t)$ à travers le circuit,
- le circuit coupe des lignes de champ magnétique dans son déplacement,

alors la f.é.m. induite dans le circuit est $e = -\frac{d\varphi}{dt}$.

Dans le cas d'un **circuit fixe dans un champ magnétique variable**, la f.é.m. induite dans le circuit est $e = -\frac{d\varphi}{dt}$.

SYNTHÈSE

SAVOIRS

- établir l'équation différentielle du circuit
- loi de Faraday
- algébrisation de la loi de Faraday

SAVOIR-FAIRE

- évaluer un flux magnétique
- décrire, mettre en œuvre et interpréter des expériences illustrant les lois de Lenz et de Faraday

MOTS-CLÉS

- flux magnétique
- f.é.m. induite
- loi de Lenz
- loi de Faraday

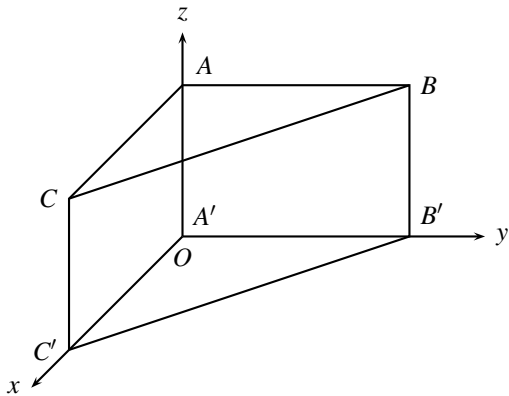
S'ENTRAÎNER

28.1 Flux magnétique (★)

L'espace est décrit par un repère cartésien $Oxyz$ et baigné par un champ magnétique uniforme et stationnaire $\vec{B} = \begin{pmatrix} -1 \\ 2 \\ 4 \end{pmatrix}$ (chaque coordonnée est exprimée en teslas). Quel est le flux magnétique à travers une surface d'aire $S = 5,0 \cdot 10^{-1} \text{ m}^2$ dessinée sur le plan d'équation $x - 2y + z = 2$?

28.2 Flux magnétique à travers un prisme (★)

On considère le prisme droit suivant, où $AB = 2 \text{ m}$, $AC = 2 \text{ m}$ et $AA' = 1,5 \text{ m}$.



Évaluer le flux magnétique à travers les surfaces ABC , $ABB'A'$, $ACC'A'$ et $CBB'C'$ (les surfaces sont orientées dans l'ordre des points, $A \rightarrow B \rightarrow C$ pour ABC), dans le cas d'un champ magnétique :

1. $\vec{B} = 0,2\vec{u}_x$ (T),
2. $\vec{B} = -0,2\vec{u}_z$ (T),
3. dans le plan yOz , de norme $B = 0,2 \text{ T}$ et qui fait un angle $(\vec{u}_y, \vec{B}) = \pi/4$,
4. dans le plan xOy , de norme $B = 0,2 \text{ T}$ et qui fait un angle $(\vec{u}_x, \vec{B}) = \pi/4$.

28.3 Spire en rotation (★)

Une spire circulaire de surface S est en rotation, à la vitesse angulaire constante ω , autour d'un de ses diamètres, qui constitue l'axe Δ . Elle est placée dans un champ magnétique uniforme et stationnaire $\vec{B}$, orthogonal à Δ .

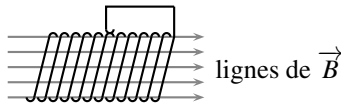
1. Établir l'expression de la f.é.m. induite e dans la spire.
2. Sachant que le courant induit vaut $i = e/R$, où R est la résistance électrique de la spire, établir la valeur du moment magnétique de la spire.

CHAPITRE 28 – LOIS DE L'INDUCTION

3. En déduire le couple de Laplace instantané puis moyen qui s'exerce sur la spire.

28.4 Bobine de constitution variable (★)

Dans la bobine suivante, un curseur fait varier le nombre de spires de la bobine. Elle est placée dans un champ magnétique dont les lignes de champ sont représentées sur la figure suivante.

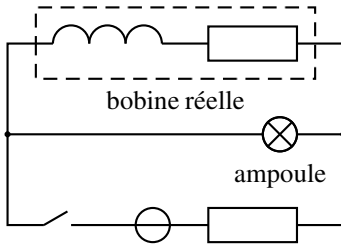


1. Que dire de la norme du champ magnétique dans la bobine ?
2. Expliquer pourquoi le flux varie à travers la portion de bobine comprise entre le curseur et l'extrémité de droite, mais qu'on n'observe aucune f.é.m. induite.

28.5 Circuit électrique (★)

Un générateur et une résistance, une bobine réelle et une ampoule sont branchées en parallèle. Interpréter, sans équation, les observations suivantes :

1. à la fermeture du circuit, l'ampoule brille plus intensément un instant, puis voit son éclat pâlir et rester constant,
2. à l'ouverture du circuit, l'ampoule brille beaucoup plus intensément puis s'éteint.



28.6 Influence du champ terrestre sur un téléphone portable (★★)

Un expérimentateur tient son téléphone portable dans sa main. Son bras passe rapidement d'une position horizontale à une position verticale afin d'entrer en communication. On tient compte de la composante horizontale du champ magnétique terrestre, d'environ 2.10^{-5} T.

1. Pourquoi le circuit électronique du téléphone est-il un circuit fermé ?
2. Évaluer l'ordre de grandeur de la f.é.m. induite dans le téléphone lors de son déplacement. Commenter. Pour répondre à cette question, le lecteur est libre de modéliser le phénomène de manière personnelle, mais crédible.

CORRIGÉS

28.1 Flux magnétique

Le vecteur $\vec{v} = \begin{pmatrix} 1 \\ -2 \\ 1 \end{pmatrix}$ est orthogonal au plan d'équation $x - 2y + z = 2$. Le vecteur unitaire normal $\vec{n}$ au plan est : $\vec{n} = \frac{\vec{v}}{\|\vec{v}\|} = \frac{1}{\sqrt{6}} \begin{pmatrix} 1 \\ -2 \\ 1 \end{pmatrix}$. Le flux magnétique à travers la surface S est donc : $\varphi = \vec{B} \cdot \vec{S} = \pm \vec{B} \cdot S\vec{n} = \pm \frac{S}{\sqrt{6}} (-1) = \mp 2,0 \cdot 10^{-1} \text{ Wb}$, où le $\pm$ est dû au manque d'information sur l'orientation de la surface ; on ne sait pas si $\vec{S} = S\vec{n}$ ou si $\vec{S} = -S\vec{n}$.

28.2 Flux magnétique à travers un prisme

Les vecteurs surfaces sont dans chaque cas, compte tenu de la règle de la main droite pour l'orientation, des distances pour les aires, $\vec{S}_{ABC} = -2\vec{u}_z \text{ m}^2$, $\vec{S}_{ABB'A'} = -3\vec{u}_x \text{ m}^2$ et $\vec{S}_{ACC'A'} = 3\vec{u}_y \text{ m}^2$.

Quant à $CBB'C'$, il faut déjà trouver un vecteur $\vec{v}$, normal au plan. On peut passer par l'équation cartésienne du plan, de la forme $ax + by = \text{constante}$, indépendante de z car on observe que lorsqu'un point appartient au plan, si sa cote z varie, il appartient toujours au plan. On utilise les coordonnées des points B et C pour trouver a et b : $\begin{cases} a \times 0 + b \times 2 = 1 \\ a \times 2 + b \times 0 = 1, \end{cases}$ (la

constante est arbitrairement choisie à 1) donc $\begin{cases} a = 1/2 \\ b = 1/2. \end{cases}$

L'équation cartésienne du plan $CBB'C'$ est $\frac{x}{2} + \frac{y}{2} = 1$, ou $x + y = 2$; il admet comme vecteur

normal le vecteur $\vec{v} = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}$ et comme vecteur normal unitaire le vecteur $\vec{n} = \frac{\vec{v}}{\|\vec{v}\|} =$

$\frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}$. Avec $BC = \sqrt{AB^2 + AC^2} = \sqrt{8} \text{ m}$ et $BB' = 1,5 \text{ m}$, le vecteur surface est donc,

en m^2 : $\vec{S}_{CBB'C'} = \begin{pmatrix} 3 \\ 3 \\ 0 \end{pmatrix}$. Il ne reste plus qu'à effectuer les produits scalaires :

1. $\varphi_{ABC} = \vec{B} \cdot \vec{S}_{ABC} = 0$, $\varphi_{ABB'A'} = \vec{B} \cdot \vec{S}_{ABB'A'} = -0,6 \text{ Wb}$, $\varphi_{ACC'A'} = \vec{B} \cdot \vec{S}_{ACC'A'} = 0$ et $\varphi_{CBB'C'} = \vec{B} \cdot \vec{S}_{CBB'C'} = 0,6 \text{ Wb}$.

2. $\varphi_{ABC} = -0,4 \text{ Wb}$, $\varphi_{ABB'A'} = 0$, $\varphi_{ACC'A'} = 0$ et $\varphi_{CBB'C'} = 0$.

3. $\vec{B} = 0,2 \begin{pmatrix} 0 \\ \cos(\pi/4) \\ \sin(\pi/4) \end{pmatrix}$ donc $\varphi_{ABC} = -0,28 \text{ Wb}$, $\varphi_{ABB'A'} = 0$, $\varphi_{ACC'A'} = 0,42 \text{ Wb}$ et

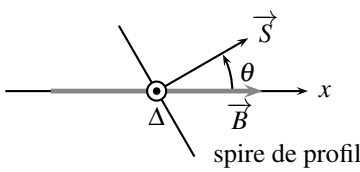
CHAPITRE 28 – LOIS DE L'INDUCTION

$\varphi_{CBB'C'} = 0,42 \text{ Wb}.$

4. $\vec{B} = 0,2 \begin{pmatrix} \cos(\pi/4) \\ \sin(\pi/4) \\ 0 \end{pmatrix}$ donc $\varphi_{ABC} = 0$, $\varphi_{ABB'A'} = -0,42 \text{ Wb}$, $\varphi_{ACC'A'} = -0,42 \text{ Wb}$ et $\varphi_{CBB'C'} = 0,85 \text{ Wb}.$

28.3 Spire en rotation

On construit l'axe (Ox) afin d'avoir $\vec{B} = B\vec{u}_x$. La spire est repérée par l'angle $\theta = \omega t$ (l'origine des temps est arbitrairement choisie à la date où $\vec{S}$ et $\vec{B}$ sont colinéaires et de même sens).



1. Le flux magnétique à travers la spire est : $\varphi = \vec{B} \cdot \vec{S} = BS \cos(\omega t)$. D'où la f.é.m. induite :
$$e = -\frac{d\varphi}{dt} = BS \frac{d\theta}{dt} \sin(\omega t) = BS\omega \sin(\omega t).$$

2. $i = \frac{e}{R} = \frac{BS}{R} \omega \sin(\omega t)$. D'où le moment magnétique de la spire : $\vec{\mathcal{M}} = i \vec{S}$ soit $\vec{\mathcal{M}} = \frac{BS}{R} \omega \sin(\omega t) \vec{S}.$

3. Le couple de Laplace qui s'exerce sur la spire est : $\vec{\Gamma} = \vec{\mathcal{M}} \wedge \vec{B} = \mathcal{M} B \sin(-\theta) \vec{u}_\Delta.$
Attention à l'angle orienté : pour le produit vectoriel, on part de $\vec{\mathcal{M}}$ pour arriver à $\vec{B}$, c'est à dire qu'on tourne d'un angle $-\theta$.

Finalement : $\vec{\Gamma} = -\frac{B^2 S^2}{R} \omega \sin^2(\omega t) \vec{e}_\Delta$ et $\langle \vec{\Gamma} \rangle = -\frac{B^2 S^2 \omega}{2R} \vec{u}_\Delta.$ Ce couple, de projection négative sur $\vec{u}_\Delta$, freine la spire.

28.4 Bobine de constitution variable

1. Les lignes de champ magnétique sont régulièrement espacées, $\vec{B}$ est uniforme (sa norme est en tout point identique).

2. Le mouvement du curseur fait varier le nombre de spires offertes au champ magnétique, le flux varie donc.

Toutefois, le curseur, dans son mouvement, ne coupe aucune ligne de champ, il n'y a donc aucune f.é.m. induite. C'est une exception à la règle du flux.

28.5 Circuit électrique

1. D'après la loi de Lenz, le f.é.m. induite dans la bobine s'oppose électriquement aux causes qui lui ont donné naissance, c'est à dire à l'augmentation de l'intensité qui la traverse. L'intensité du courant est donc principalement détournée vers l'ampoule. Lorsque la f.é.m. induite dans la bobine diminue, cet effet est de moins en moins marqué, un courant plus important circule dans la bobine au détriment de l'ampoule, qui brille moins.

2. À l'ouverture, l'ampoule se retrouve en série avec la bobine. Toute l'énergie emmagasinée dans la bobine se décharge dans l'ampoule, qui brille alors intensément. L'ampoule s'étend lorsque toute l'énergie de la bobine est dissipée, transformée en rayonnement lumineux et en effet Joule résistif.

28.6 Influence du champ terrestre sur un téléphone portable

1. Un circuit électrique est toujours un circuit fermé. Un générateur y fait circuler le courant ; c'est ici la pile du générateur, ou bien la f.é.m. induite engendrée par l'onde électromagnétique reçue.

2. Il n'y a pas de réponse unique à cette question. L'auteur propose d'évaluer la surface du circuit du téléphone à celle d'un téléphone, soit environ $S \approx 5 \text{ cm} \times 10 \text{ cm} = 50 \text{ cm}^2$.

De plus, le flux magnétique est initialement nul à travers le téléphone quand il est posé à plat. On propose de considérer le cas le plus défavorable pour son arrivée près de l'oreille, avec toute la surface offerte au champ magnétique. Alors le flux final vaut $\varphi_f \approx BS$. La variation de flux entre l'instant initial et final est $\Delta\varphi = \varphi_f - \varphi_i \approx BS$.

La durée du mouvement du téléphone est d'environ $\Delta t \approx 0,5 \text{ s}$.

D'où une f.é.m. induite d'environ $e = -\frac{\Delta\varphi}{\Delta t} \approx -2 \cdot 10^{-7} \text{ V}$.

Cette f.é.m., très faible comparée aux tensions manipulées dans un téléphone, de l'ordre du mV, ne perturbe pas son fonctionnement.

CHAPITRE 28 – LOIS DE L’INDUCTION

Circuit fixe dans un champ magnétique variable

29

Dans ce chapitre, on étudie des phénomènes d'induction dans des circuits fixes mettant en jeu des champs magnétiques variant dans le temps.

1 Auto-induction

1.1 Inductance propre

Imaginons un circuit filiforme, parcouru par un courant d'intensité i , une spire par exemple, bien que ce puisse être n'importe quel autre circuit plus compliqué. Il crée un champ magnétique $\vec{B}$, dont la norme est proportionnelle à l'intensité i du courant qui le traverse. Attendu que les lignes de champ magnétique s'enroulent autour de leurs sources, ce champ magnétique traverse le circuit qui lui a donné naissance, il crée donc un flux magnétique à travers ce circuit, nommé **flux propre**.

L'expression de ce flux magnétique propre, noté φ_p , est compliquée car le champ n'est pas nécessairement uniforme sur toute la surface du circuit. Toutefois, φ_p est proportionnel à la norme de $\vec{B}$, qui est elle même proportionnelle à l'intensité i du courant ; φ_p est finalement proportionnel à i .

Le flux, à travers un circuit, de son propre champ magnétique s'écrit :

$$\varphi_p = Li,$$

où i est l'intensité dans le circuit et L le **coefficient d'auto-inductance** du circuit.

Le coefficient d'auto-inductance est noté L ¹. On l'appelle aussi **inductance propre**.

1. En l'honneur de Heinrich Lenz.

CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE

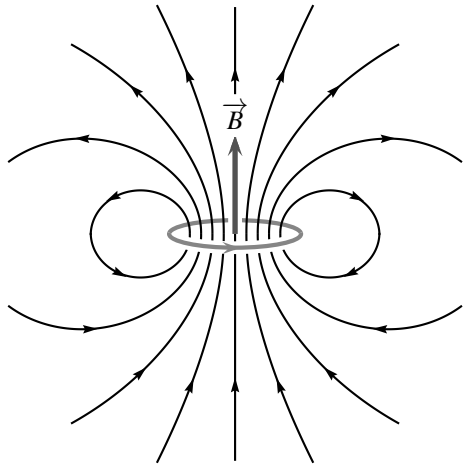


Figure 29.1 – Lignes de champ magnétique, en noir, d'une spire circulaire, en gris.

L'unité d'inductance est le **henry**², noté H.

Comment exprimer les henrys dans les unités de base du système international ? On montre, dans le chapitre sur le champ magnétique, que les teslas T, unités du champ magnétique, sont des $\text{kg.s}^{-2}.\text{A}^{-1}$. De plus, le flux magnétique est en webers Wb :

$$\varphi = \vec{B} \cdot \vec{S} \quad \text{donc} \quad \text{Wb} = \text{T.m}^2 = \text{kg.s}^{-2}.\text{A}^{-1}.\text{m}^2.$$

Et :

$$\varphi_p = Li \quad \text{donc} \quad \text{Wb} = \text{H.A} \quad \text{soit} \quad \text{H} = \text{m}^2.\text{kg.s}^{-2}.\text{A}^{-2}.$$

Ce résultat n'est pas à retenir, mais la méthode suivie doit être connue. Il montre l'intérêt à utiliser une unité dérivée afin d'alléger les notations.

L'unité d'inductance est le henry, dont le symbole est H.

L'auto-inductance d'un circuit est toujours *positive*. Il est possible de le vérifier en considérant la figure 29.1 (sur cette figure la flèche indique le sens réel du courant) si on choisit pour sens positif le long de la spire ce même sens, alors $i > 0$ et, en appliquant la règle de la main droite, on constate que le champ magnétique traverse la surface de la spire dans le sens positif donc que $\varphi_p > 0$. Les deux sont négatifs si on fait le choix inverse. Le signe du flux propre est le même que celui de l'intensité. Cela provient du fait que la règle de la main droite donne à la fois le sens positif de traversée de la surface du circuit à partir du sens conventionnel du courant et la direction du champ magnétique à partir du sens réel du courant. Ainsi l'autoinductance est positive, quel que soit le sens positif du courant choisi.

2. En l'honneur de Joseph Henry, 1797 – 1878, physicien américain, successivement professeur de « natural philosophy » à Princeton, directeur de la Smithsonian Institution et président de la National Academy of Sciences. Ses travaux abordèrent avec éclectisme de nombreux thèmes, dont la météorologie. Il reste célèbre pour sa découverte des phénomènes d'auto-induction et de mutuelle induction, exposés dans son article de 1838, *On Electro-Dynamic Induction*. Faraday découvrit en même temps ces phénomènes au Royaume-Uni, mais Henry publia en premier.

Remarque

Dans le cas d'un circuit comportant N spires identiques, le champ magnétique est proportionnel à N . De plus, le flux à travers le circuit est égal à N fois le flux à travers une spire. Le coefficient d'inductance propre est donc proportionnel à N^2 .

1.2 Calcul d'une inductance propre

Le but est d'établir l'expression de l'inductance propre L d'une bobine longue de longueur ℓ , nommée solénoïde, constituée de N spires, chacune parcourue par un courant d'intensité i , de surface S .

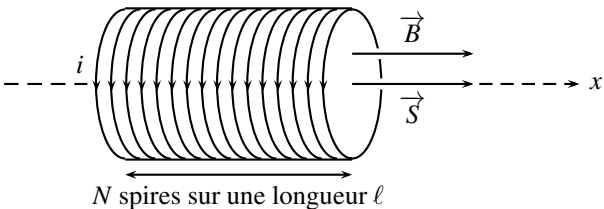


Figure 29.2 – Solénoïde dont on calcule l'inductance propre.

Soit $n = \frac{N}{\ell}$ le nombre de spires par unité de longueur de la bobine. Alors le champ magnétique créé dans la bobine a pour expression :

$$\vec{B} = \mu_0 n i \vec{u}_x = \mu_0 \frac{N}{\ell} i \vec{u}_x.$$

Le flux de ce champ à travers une seule spire est :

$$\varphi_{1 \text{ spire}} = \vec{B} \cdot \vec{S} = \mu_0 \frac{N}{\ell} i S.$$

Le flux total à travers les N spires, c'est à dire le flux propre à travers la bobine, est :

$$\varphi_p = N \varphi_{1 \text{ spire}} = \mu_0 \frac{N^2}{\ell} i S.$$

On identifie alors l'expression de l'autoinductance :

$$\varphi_p = Li = \mu_0 \frac{N^2}{\ell} i S \quad \text{donc} \quad L = \mu_0 \frac{N^2}{\ell} S.$$

Par exemple, pour une bobine de 1000 spires de rayon $a = 3 \text{ cm}$, réparties sur une longueur $\ell = 10 \text{ cm}$:

$$L = 4\pi \cdot 10^{-7} \times \frac{1000^2}{10 \cdot 10^{-2}} \times \pi (3 \cdot 10^{-2})^2 = 35 \text{ mH}.$$

CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE

1.3 Circuit électrique équivalent

Si l'intensité du courant traversant le circuit varie dans le temps, le flux propre varie et il apparaît donc une force électromotrice induite, nommée force électromotrice autoinduite, qui s'exprime avec la loi de Faraday :

$$e(t) = -\frac{d\varphi_p}{dt} = -L\frac{di}{dt},$$

si l'auto-inductance L est une constante, c'est à dire si les caractéristiques géométriques du circuit ne varient pas.

En prenant soin de bien orienter bien la f.é.m. induite dans le même sens que l'intensité du courant (voir chapitre précédent), on peut dessiner le schéma électrique équivalent :

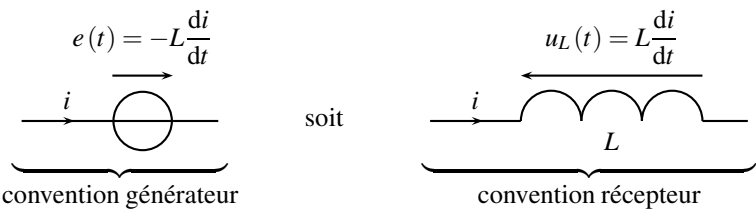


Figure 29.3 – Schéma électrique équivalent pour l'autoinduction.

On voit donc que la f.é.m. auto-induite peut être représentée soit par un générateur de f.é.m. $e(t)$, orienté en convention générateur, soit par le symbole normalisé d'une inductance L , avec une tension $u_L(t)$ à ses bornes, orientée en convention récepteur.

Remarque

Dans le cas où le circuit est plongé dans un champ magnétique extérieur, par exemple le champ terrestre, un flux supplémentaire à travers le circuit s'ajoute au flux propre, $\varphi_{tot} = \varphi_p + \varphi_{ext}$. Dans ce cas, la f.é.m. induite est :

$$e(t) = -L\frac{di}{dt} - \frac{d\varphi_{ext}}{dt}.$$

Cette f.é.m. additionnelle due au champ magnétique extérieur forme du bruit électrique, par exemple dans le cas d'un téléphone portable dont la position bouge, même faiblement, dans le champ terrestre. On minimise ce bruit en éliminant toutes les inductances des circuits électroniques, sachant qu'au besoin, une inductance est simulée avec des résistances, des condensateurs et des amplificateurs linéaires intégrés.

1.4 Loi de Lenz

Les formules trouvées sont cohérentes avec la loi de modération de Lenz. Examinons par exemple une bobine, branchée sur un générateur extérieur qui fait croître le courant dans la bobine : $\frac{di}{dt} > 0$. Alors, $e = -L\frac{di}{dt} < 0$ et d'après le schéma électrique de la figure 29.3, cette f.é.m. induite négative s'oppose à l'augmentation du courant.

1.5 Mesure d’une inductance propre

On mesure expérimentalement la valeur de l’autoinductance L en l’insérant dans un circuit série comportant un générateur de force électromotrice $e_g(t)$ et une résistance R_0 :

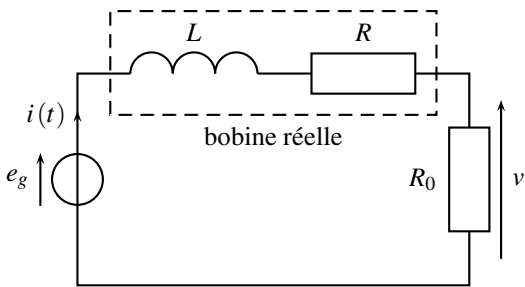


Figure 29.4 – Montage électrique pour la mesure de l’autoinductance.

Les fils de cuivre qui constituent la bobine présentent une résistance R . On ajoute au circuit une résistance connue R_0 très supérieure à R . Ainsi, la résistance de l’ensemble est-elle $R_0 + R \simeq R_0$. L’ensemble est alors régi par l’équation différentielle du premier ordre, écrite sous forme canonique :

$$\begin{cases} e_g(t) = L \frac{di}{dt} + R_0 i(t) \\ v(t) = R_0 i(t) \end{cases} \quad \text{donc} \quad \frac{L}{R_0} \frac{dv}{dt} + v(t) = e_g(t).$$

On mesure la constante de temps $\tau = \frac{L}{R_0}$ grâce au temps de montée à 63%, dans le cas où la tension d’alimentation $e_g(t)$ est un signal créneau de grande période par rapport à τ . Cette méthode est détaillée dans le chapitre d’électricité sur les circuits linéaires du premier ordre. La valeur de τ mène à celle de L , attendu que celle de R_0 est connue.

1.6 Étude énergétique

On reprend le montage électrique de la figure 29.4. Afin de mener un bilan de puissance, on multiplie la loi des mailles par $i(t)$:

$$e_g(t) \times i(t) = L \frac{di}{dt} \times i(t) + R_0 i(t)^2.$$

Attendu que $\frac{di}{dt} \times i(t) = \frac{d}{dt} \left(\frac{i^2}{2} \right)$:

$$e_g i = \frac{d}{dt} \left(\frac{1}{2} L i^2 \right) + R_0 i^2.$$

Le terme $e_g i$ représente la puissance délivrée par le générateur, qui se répartie en deux contributions :

CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE

- une puissance $R_0 i^2$, dissipée par effet Joule dans la résistance R_0 ,
- une puissance $\frac{d}{dt} \left(\frac{1}{2} Li^2 \right)$, stockée dans la bobine.

Il apparaît dès lors que de l'énergie est stockée dans la bobine, et que cette énergie s'écrit $\frac{1}{2} Li^2 + \text{constante}$. Il s'agit de l'énergie du champ magnétique créé par la bobine. Si l'intensité s'annule, le champ magnétique s'annule aussi et l'énergie magnétique aussi, ce qui montre que la constante est égale à zéro.

L'énergie magnétique d'un circuit d'auto-inductance L parcouru par un courant d'intensité i est :

$$\mathcal{E}_{\text{magné}} = \frac{1}{2} Li^2.$$

2 Deux circuits en interaction

2.1 Inductance mutuelle

Imaginons un circuit filiforme $\mathcal{C}_1$, parcouru par un courant d'intensité i_1 , une spire par exemple, bien que ce puisse être n'importe quel autre circuit plus compliqué. Il crée un champ magnétique $\vec{B}_1$, dont la norme est proportionnelle à l'intensité i_1 du courant qui le traverse. Un circuit $\mathcal{C}_2$ est disposé dans le champ $\vec{B}_1$ dont il intercepte des lignes de champ. Ainsi $\vec{B}_1$ a-t-il un flux à travers $\mathcal{C}_2$ que l'on note $\varphi_{1 \rightarrow 2}$. C'est le flux envoyé par le circuit $\mathcal{C}_1$ à travers le circuit $\mathcal{C}_2$.

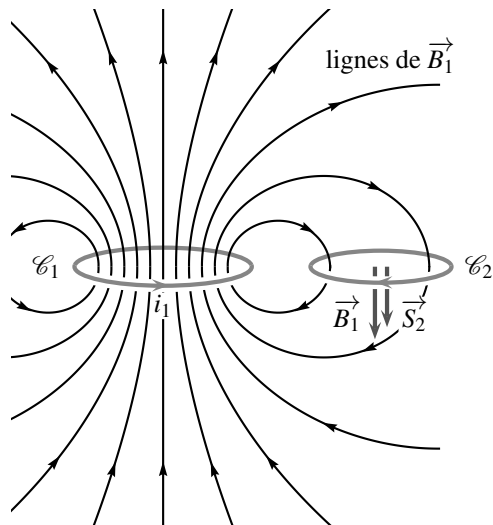


Figure 29.5 – Deux circuits magnétiquement couplés.

L'expression de ce flux magnétique $\varphi_{1 \rightarrow 2}$ est compliquée car $\vec{B}_1$ n'est pas nécessairement uniforme sur toute la surface du circuit $\mathcal{C}_2$. Toutefois, $\varphi_{1 \rightarrow 2}$ est proportionnel à la norme de

$\vec{B}_1$, qui est elle-même proportionnelle à l'intensité i_1 du courant ; $\varphi_{1 \rightarrow 2}$ est donc finalement proportionnel à i_1 .

De même, si $\mathcal{C}_2$ est parcouru par le courant i_2 il envoie un flux magnétique $\varphi_{2 \rightarrow 1}$ à travers $\mathcal{C}_1$ qui est proportionnel à i_2 . Les coefficients de proportionnalité entre $\varphi_{1 \rightarrow 2}$ et i_1 et entre $\varphi_{2 \rightarrow 1}$ et i_2 sont identiques : leur valeur commune est notée M . Ce résultat constitue le théorème de **Neumann**³ que l'on admettra ici.

M est appelé **inductance mutuelle** entre les deux circuits $\mathcal{C}_1$ et $\mathcal{C}_2$. M s'exprime en **henry**.

Les flux magnétiques envoyés réciproquement par un circuit $\mathcal{C}_1$ et un circuit $\mathcal{C}_2$ l'un à travers l'autre sont donnés par les formules :

$$\varphi_{1 \rightarrow 2} = M i_1 \quad \text{et} \quad \varphi_{2 \rightarrow 1} = M i_2,$$

où M est appelé **inductance mutuelle** entre les deux circuits $\mathcal{C}_1$ et $\mathcal{C}_2$.

L'inductance mutuelle se mesure en henry, de symbole H.



L'intensité du courant qui apparaît dans l'expression de ces flux est à chaque fois l'intensité dans le circuit qui crée le champ magnétique.

Le signe de l'inductance mutuelle est arbitraire. Il dépend de l'orientation relative des deux circuits.

2.2 Circuits électriques équivalents

On considère l'ensemble de deux circuits, $\mathcal{C}_1$ et $\mathcal{C}_2$, de la figure 29.5.

$\mathcal{C}_1$ est parcouru par le courant d'intensité i_1 . Ce courant crée un champ magnétique $\vec{B}_1$, donc un flux propre $\varphi_{1,p}$ à travers $\mathcal{C}_1$: $\varphi_{1,p} = L_1 i_1$.

$\mathcal{C}_2$ est parcouru par i_2 , qui crée le champ $\vec{B}_2$. $\vec{B}_2$ envoie un flux $\varphi_{2 \rightarrow 1}$ à travers $\mathcal{C}_1$: $\varphi_{2 \rightarrow 1} = M i_2$.

Le flux total à travers $\mathcal{C}_1$ est donc :

$$\varphi_1 = \varphi_{1,p} + \varphi_{2 \rightarrow 1} = L_1 i_1 + M i_2.$$

La f.é.m. induite dans $\mathcal{C}_1$ est alors :

$$e_1(t) = -\frac{d\varphi_1}{dt} = -L_1 \frac{di_1}{dt} - M \frac{di_2}{dt}.$$

On procède de même pour la f.é.m. induite dans $\mathcal{C}_2$:

$$\varphi_2 = \varphi_{2,p} + \varphi_{1 \rightarrow 2} = L_2 i_2 + M i_1 \quad \text{donc} \quad e_2(t) = -L_2 \frac{di_2}{dt} - M \frac{di_1}{dt}.$$

3. nommé d'après son découvreur, Franz Ernst Neumann, 1798 – 1895, physicien allemand, professeur à l'université de Königsberg, dont les recherches portèrent sur la cristallographie, la thermodynamique, les ondes lumineuses et l'induction.

CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE

On en déduit les circuits électriques équivalents aux deux spires :

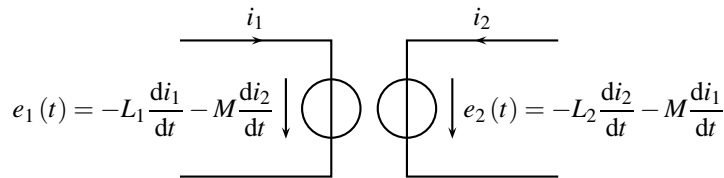


Figure 29.6 – Circuits électriques équivalents (convention générateur).

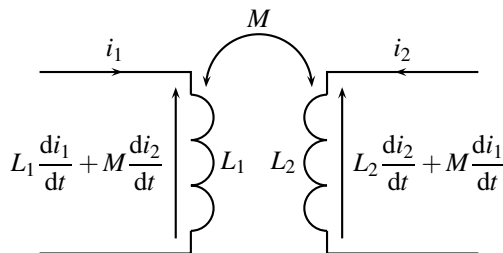


Figure 29.7 – Circuits électriques équivalents (convention récepteur).

2.3 Étude harmonique

On considère l'ensemble de deux circuits, $\mathcal{C}_1$ et $\mathcal{C}_2$, de la figure 29.5, couplés par mutuelle induction. $\mathcal{C}_1$, de résistance électrique R_1 , est alimenté par un générateur qui impose la tension $v_1(t) = V_0 \cos(\omega t)$. Quant à $\mathcal{C}_2$, de résistance électrique R_2 , il est court-circuité.

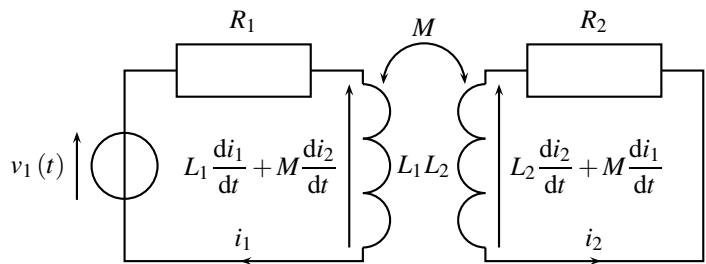


Figure 29.8 – Circuits couplés étudiés.

La loi des mailles mène aux équations couplées :

$$\begin{cases} v_1(t) = L_1 \frac{di_1}{dt} + M \frac{di_2}{dt} + R_1 i_1 \\ 0 = L_2 \frac{di_2}{dt} + M \frac{di_1}{dt} + R_2 i_2 \end{cases}$$

En régime sinusoïdal permanent à la pulsation ω , imposée par $v_1(t)$:

$$\begin{cases} \underline{v}_1(t) = (R_1 + jL_1\omega) \underline{i}_1(t) + jM\omega \underline{i}_2(t) \\ 0 = (R_2 + jL_2\omega) \underline{i}_2(t) + jM\omega \underline{i}_1(t) \end{cases}$$

On peut dès lors, par exemple, établir le schéma électrique équivalent vu du circuit $\mathcal{C}_1$. La seconde équation mène à :

$$\underline{i}_2(t) = -\frac{jM\omega}{R_2 + jL_2\omega}\underline{i}_1(t),$$

qu'on injecte dans la première :

$$\begin{aligned}\underline{v}_1(t) &= (R_1 + jL_1\omega)\underline{i}_1(t) - \frac{(jM\omega)^2}{R_2 + jL_2\omega}\underline{i}_1(t) \\ \underline{v}_1(t) &= \underbrace{\left(R_1 + jL_1\omega - \frac{(jM\omega)^2}{R_2 + jL_2\omega}\right)}_{\underline{Z}}\underline{i}_1(t)\end{aligned}$$

Le couplage entre les deux circuits est équivalent, vu du circuit $\mathcal{C}_1$, à un unique dipôle d'impédance $\underline{Z}$, dans lequel intervient les caractéristiques de $\mathcal{C}_2$ via la couplage inductif.

2.4 Étude énergétique

Pour les circuits couplés de la figure 29.8, multiplions par i_1 et i_2 les équations électriques :

$$\begin{cases} v_1(t) i_1(t) = L_1 \frac{di_1}{dt} i_1(t) + M \frac{di_2}{dt} i_1(t) + R_1 i_1^2(t) \\ 0 = L_2 \frac{di_2}{dt} i_2(t) + M \frac{di_1}{dt} i_2(t) + R_2 i_2^2(t) \end{cases}$$

Additionnons ces équations :

$$\begin{aligned}v_1 i_1 &= \frac{d}{dt} \left(\frac{1}{2} L_1 i_1^2 + \frac{1}{2} L_2 i_2^2 \right) + \underbrace{M \frac{di_1}{dt} i_2 + M \frac{di_2}{dt} i_1}_{= \frac{d}{dt} (M i_1 i_2)} + R_1 i_1^2 + R_2 i_2^2, \\ &= \frac{d}{dt} (M i_1 i_2)\end{aligned}$$

soit :

$$v_1 i_1 = \frac{d}{dt} \left(\frac{1}{2} L_1 i_1^2 + \frac{1}{2} L_2 i_2^2 + M i_1 i_2 \right) + R_1 i_1^2 + R_2 i_2^2.$$

La puissance $v_1 i_1$ délivrée par le générateur extérieur se répartit en deux termes :

- une puissance $R_1 i_1^2 + R_2 i_2^2$, perdue par effet Joule dans les deux circuits,
- une puissance magnétique qui dérive d'une énergie magnétique du circuit $\mathcal{C}_1$ $\left(\frac{1}{2} L_1 i_1^2 \right)$, du circuit $\mathcal{C}_2$ $\left(\frac{1}{2} L_2 i_2^2 \right)$, et du couplage entre les deux circuits $(M i_1 i_2)$.

L'énergie magnétique de deux circuits, couplés par mutuelle induction, est :

$$\mathcal{E}_{\text{magné}} = \frac{1}{2} L_1 i_1^2 + \frac{1}{2} L_2 i_2^2 + M i_1 i_2.$$

3 Transformateur de tension (PTSI)

3.1 Constitution

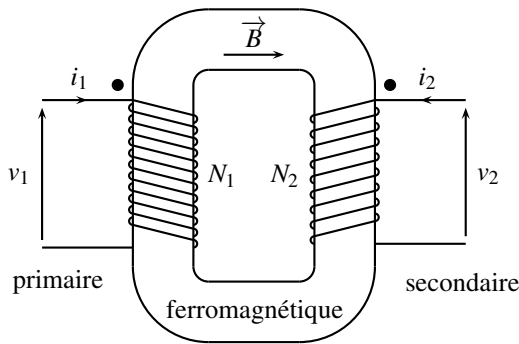


Figure 29.9 – Transformateur.

Un **transformateur** est un appareil qui modifie l'amplitude de tensions et de courants alternatifs.

Il se compose d'une carcasse **ferromagnétique**, et de deux enroulements. Le matériau ferromagnétique est choisi pour sa capacité à canaliser les lignes de champ magnétique, de même qu'un tuyau canalise un liquide en écoulement. Ainsi, il n'existe de champ magnétique qu'à l'intérieur du « circuit magnétique » formé par le ferromagnétique. Les enroulements sont constitués de fils de cuivre, bobinés autour du circuit magnétique.

L'enroulement primaire, ou plus simplement le **primaire**, ici constitué de N_1 spires, est celui qui reçoit l'énergie électrique, que restitue à la charge l'enroulement secondaire, ou **secondaire**, constitué de N_2 spires.

La signification des points ● sera abordée au paragraphe 3.3.

3.2 Principe de fonctionnement

Le primaire, soumis à la tension alternative $v_1(t)$, est parcouru par le courant alternatif d'intensité $i_1(t)$. Ce courant $i_1(t)$ crée un champ magnétique variable $\vec{B}(t)$. Le ferromagnétique canalise ce champ jusqu'au secondaire. Le champ variable $\vec{B}(t)$ crée alors un flux variable dans l'enroulement secondaire. Une f.é.m. est donc induite au secondaire.

Quel est le lien entre $v_2(t)$, tension au secondaire, et $v_1(t)$, tension au primaire ?

Soit S la section du circuit magnétique. Le flux de $\vec{B}(t)$ à travers une section droite, c'est-à-dire orthogonale au champ magnétique, est donc $B(t)S$. Le flux total à travers les N_1 spires du primaire est donc $\varphi_1(t) = N_1 B(t)S$, celui au secondaire est $\varphi_2(t) = N_2 B(t)S$. Les f.é.m. induites au primaire et au secondaire sont alors :

$$e_1(t) = -\frac{d\varphi_1}{dt} = -N_1 S \frac{dB}{dt} \quad \text{et} \quad e_2(t) = -N_2 S \frac{dB}{dt}.$$

Le schéma électrique équivalent est donc, en se souvenant bien que les f.é.m. induites sont dans le même sens que les courants :

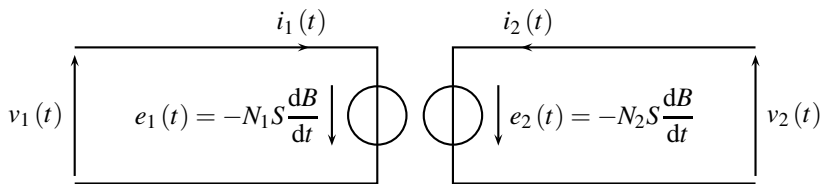


Figure 29.10 – Circuit électrique équivalent du transformateur.

Ainsi $v_1 = -e_1$ et $v_2 = -e_2$. Attendu que $B(t)$ est variable, $\frac{dB}{dt}$ est différent de 0 :

$$\begin{cases} v_1(t) = N_1 S \frac{dB}{dt} \\ v_2(t) = N_2 S \frac{dB}{dt} \end{cases} \text{ implique } \frac{v_2(t)}{v_1(t)} = \frac{N_2}{N_1} = m.$$

où $m = \frac{N_2}{N_1}$ est le rapport de transformation.

⚠ La formule précédente n'est valable que pour des tensions alternatives. En régime indépendant du temps, $v_2 = 0$: **un transformateur ne laisse pas passer le continu.**

Les tensions primaire et secondaire, dans un transformateur en régime alternatif, sont reliées par le **rapport de transformation** m :

$$\frac{v_2(t)}{v_1(t)} = \frac{N_2}{N_1} = m.$$

3.3 Normalisation et orientation des courants

Le schéma normalisé du transformateur est :

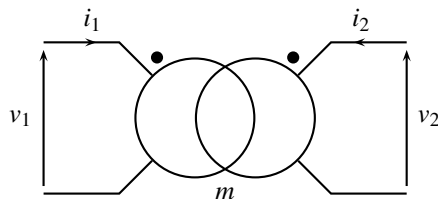


Figure 29.11 – Schéma normalisé d'un transformateur.

Il apparaît ici le besoin de préciser très correctement l'orientation des courants et la signification des points (●) au primaire et au secondaire.

Observons la figure 29.9. Au primaire, le courant d'intensité i_1 entre par la borne marquée du point. D'après la règle de la main droite, il crée un champ magnétique orienté dans le sens positif, arbitrairement choisi et indiqué sur la figure. De même au secondaire, le courant

CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE

d'intensité i_2 entre lui aussi par la borne marquée du point et crée, pareillement, un champ magnétique orienté dans le sens positif.

La signification des points est claire : tout courant qui entre par le point (●) crée, dans le matériau ferromagnétique, un champ magnétique orienté dans le sens positif choisi.

Quand on dessine les enroulements, comme sur la figure 29.9, le sens d'enroulement et le point sont redondants. Mais le sens des enroulements n'existe plus sur le schéma de la figure 29.11 ; le marquage des points y est impératif.

Que se passe-t-il quand le courant n'entre pas par le point ?

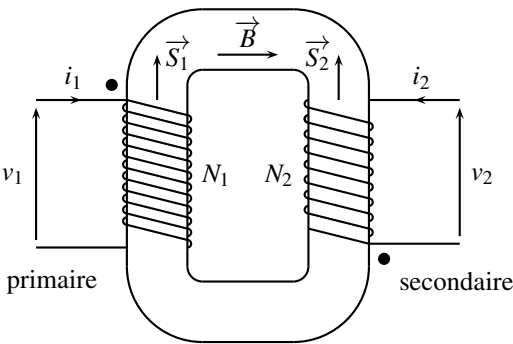


Figure 29.12 – Gestion des points dans un transformateur.

Sur la figure 29.12, le courant d'intensité i_2 n'entre pas par le point. Ce courant oriente les spires du secondaire, donc leur vecteur surface $\vec{S}_2$ d'après la règle de la main droite, de telle manière que le flux au secondaire devient négatif, $\varphi_2(t) = -N_2SB(t)$. Dès lors, $u_2(t) = -e_2(t) = \frac{d\varphi_2}{dt} = -N_2S\frac{dB}{dt}$. Le lien entre la tension primaire et secondaire devient :

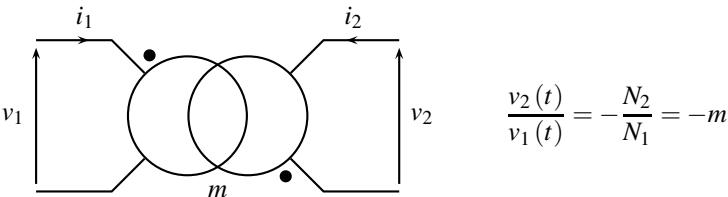


Figure 29.13 – Gestion des points avec le schéma normalisé d'un transformateur.

Il convient donc de vérifier le câblage du transformateur pour utiliser correctement le lien entre les tensions primaire et secondaire.

3.4 Courants de Foucault

Le matériau ferromagnétique est soumis à un champ magnétique variable. Des f.é.m. sont donc induites dans le volume. Attendu que le matériau est conducteur, des courants induits apparaissent. Ces courants, qui circulent dans la masse du ferromagnétique, sont nommés **courants de Foucault**⁴, ou *eddy currents* en anglais.

Les courants de Foucault sont responsables de pertes par effet Joule. Afin de les minimiser, voire même de les supprimer, on **feuillette** le matériau ferromagnétique en le constituant de tôles minces, jusqu'à 0,3 mm d'épaisseur, séparées par une couche isolante aussi mince que possible, 10^{-2} mm ou moins. La couche isolante bloque la circulation du courant et prévient les pertes par effet Joule.

Les pertes par courants de Foucault sont supprimées dans un transformateur en feuilletant le matériau ferromagnétique.

3.5 Utilisation

La principale utilité des transformateurs est d'abaisser ou d'élever des tensions alternatives. Citons, par exemple, les transformateurs EDF qui transforment la moyenne tension 20 kV en tension domestique 400 V, ou bien les transformateurs des chargeurs de portable, aisément identifiables car ils en forment la partie la plus volumineuse et la plus lourde.

Un transformateur permet aussi d'isoler électriquement deux circuits. On l'utilise en **transformateur d'isolement** pour lequel le rapport de transformation est $m = 1$. La tension secondaire est une copie conforme de la tension primaire alternative. Les circuits primaire et secondaire n'ont aucun lien électrique. Ils peuvent en particulier avoir deux masses différentes, sans risquer de court-circuit.

4. En l'honneur de leur découvreur, le français Léon Foucault, 1819 – 1868. Léon Foucault est très célèbre pour son expérience dite du pendule de Foucault, réalisée sous la coupole du Panthéon, à Paris, en 1851, grâce à laquelle il démontra l'influence de la rotation de la Terre sur le mouvement d'un pendule pesant. On lui doit aussi une mesure de la vitesse de la lumière.

CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE

SYNTHÈSE

SAVOIRS

- ordre de grandeur de l'inductance propre d'une bobine de grande longueur
- énergie magnétique mutuelle
- utilisation du transformateur

SAVOIR-FAIRE

- différencier le flux propre des flux extérieurs
- calcul de l'inductance propre d'une bobine de grande longueur
- mesurer la valeur de l'inductance propre d'une bobine
- établir le système d'équations couplées en régime sinusoïdal de deux circuits couplés par mutuelle induction
- conduire un bilan d'énergie pour l'auto-induction
- conduire un bilan d'énergie pour l'induction mutuelle
- établir la loi des tensions pour un transformateur (PTSI)

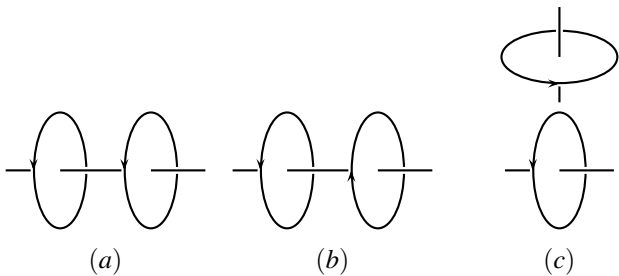
MOTS-CLÉS

- | | | |
|------------------|----------------------|-------------------------|
| • auto-induction | • auto-inductance | • inductance mutuelle |
| • flux propre | • induction mutuelle | • transformateur (PTSI) |

S'ENTRAÎNER

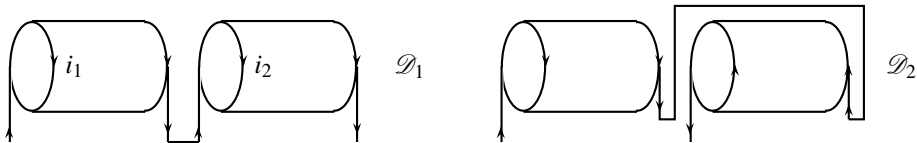
29.1 Bobines en séries (d'après CAPES) (★)

1. Deux spires orientées sont placées suivant trois dispositions (a), (b) et (c). Les spires sont coaxiales dans les cas (a) et (b) et d'axes orthogonaux dans le cas (c). Indiquer le signe de M dans chaque cas.



2. Calculer l'inductance propre d'une bobine de longueur ℓ , comportant N spires de surface S . On précisera clairement les hypothèses de calcul.

3. On dispose de deux bobines circulaires identiques, $\mathcal{B}_1$ et $\mathcal{B}_2$, d'inductance propre L . $\mathcal{B}_1$ est parcourue par le courant d'intensité i_1 et $\mathcal{B}_2$ par i_2 . Elles sont placées sur le même axe Δ et sont branchées en série afin de constituer deux nouveaux dipôles, $\mathcal{D}_1$ et $\mathcal{D}_2$. On note M leur mutuelle dans le cas du dipôle $\mathcal{D}_1$.



Les bobines étant très proches, on émet l'hypothèse $\mathcal{H}$: les champs magnétiques créés par chaque bobine sont tous portés par l'axe Δ et leur norme est constante sur toute section droite.

- a. Rappeler l'expression du flux passant dans $\mathcal{B}_1$ en fonction de i_1 et de i_2 . De même pour le flux dans $\mathcal{B}_2$.
- b. Calculer le flux total à travers $\mathcal{D}_1$ et en déduire l'inductance L_1 de $\mathcal{D}_1$ en fonction de M et de L . Même question pour $\mathcal{D}_2$ pour trouver L_2 .
- c. À quelle condition peut-on avoir $L = L_1 + L_2$?
- d. Discuter l'hypothèse $\mathcal{H}$.

29.2 Autoinductance d'un solénoïde (★)

On double le nombre de spire d'un solénoïde tout en réduisant le courant qui le traverse de moitié. Comment varie son inductance propre ?

CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE

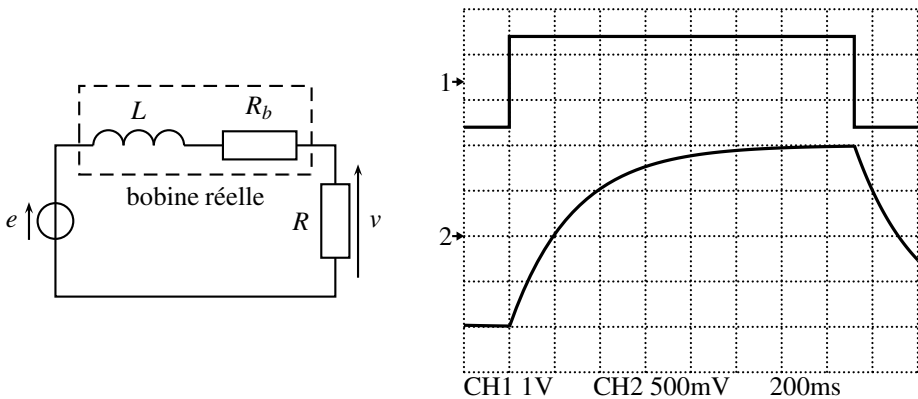
29.3 Spire autour d'un solénoïde (★)

Un solénoïde de rayon $R_1 = 2 \text{ cm}$, constitué de $n = 10$ spires par cm, est alimenté par un générateur de f.é.m. $U = 30 \text{ V}$. La résistance interne du générateur est de $1,2 \Omega$ et celle du fil du solénoïde est $6,8 \Omega$. Une spire conductrice $\mathcal{S}$, de rayon $R_2 = 4 \text{ cm}$, est placée autour du solénoïde ; elle a le même axe que celui-ci.

1. Quel est le flux magnétique à travers la spire ?
2. Par modification du circuit alimentant le solénoïde à la date $t = 0$, l'intensité du courant qui le traverse décroît au cours du temps selon la loi $i(t) = i_0 \exp\left(-\frac{t}{\tau}\right)$. Quelle est l'unité de τ ?
3. Quelle est la f.é.m. induite dans la spire pour $t > 0$?

29.4 Mesure d'une autoinductance (★)

On réalise le montage suivant, avec $R = 10^3 \Omega$ et où le générateur de f.é.m. e délivre une tension créneaux de valeur moyenne nulle. Les tensions e et v sont observées à l'oscilloscope.

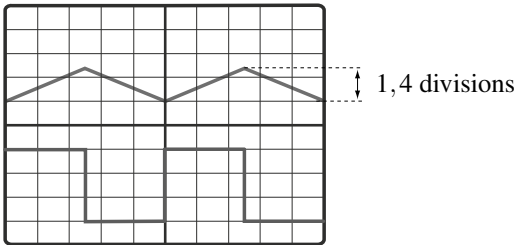


1. On mesure $R_b = 1,2 \Omega$. Proposer un montage qui permette de mesurer cette valeur. On dispose de tout l'appareillage classique d'électricité : voltmètre, ampèremètre, oscilloscope, générateur de tension constante ou alternative...
2. Dédurre des graphes la valeur de L .

29.5 Mesure d'une inductance mutuelle (★)

On considère deux bobines identiques, formées de N spires circulaires de rayon R (bobinées sur une seule épaisseur), d'inductance L , que l'on place de façon que les deux bobinages soient coaxiaux, avec le même sens d'enroulement, la distance entre leurs centres étant repérée le long de l'axe commun (Oz) par la longueur d . On se propose de mesurer le couplage entre les deux bobines en envoyant dans l'une d'elles, dite la première, une tension triangulaire et en comparant à l'oscilloscope cette tension avec la tension induite dans l'autre, celle-ci étant en circuit ouvert. On a branché en série entre le générateur de fonction et la première bobine une résistance $R = 100 \Omega$. On néglige la résistance r des bobines.

- 1. Devant quoi la résistance r des bobines est-elle négligeable ?
- 2. Dessiner le schéma du montage.
- 3. Les traces observées à l'oscilloscope ont l'allure suivante :



Les réglages de l'oscilloscope sont les suivants :

- balayage horizontal : 0,2 ms/div ;
- trace supérieure : 1 V/div ;
- trace inférieure variable (voir tableau).

En faisant varier la distance d entre les bobines, on observe pour l'amplitude crête à crête A du signal induit, mesurée en divisions de l'écran, les valeurs suivantes :

Calibre	0.01 V/div			5 mV/div			2 mV/div		1 mv/div
d (cm)	4	5	6	7	8	10	12	16	20
A	4,3	3,3	2,6	4,3	3,4	2,3	4	2,1	2,4

Écrire les équations électriques du circuit.

Établir l'expression de l'inductance mutuelle M entre les deux bobines en fonction de la période T du signal d'entrée, de son amplitude crête à crête Δe , de l'amplitude crête à crête A du signal induit et de la résistance R . Calculer alors, en mH, l'inductance mutuelle M entre les deux bobines pour chaque valeur de d .

APPROFONDIR

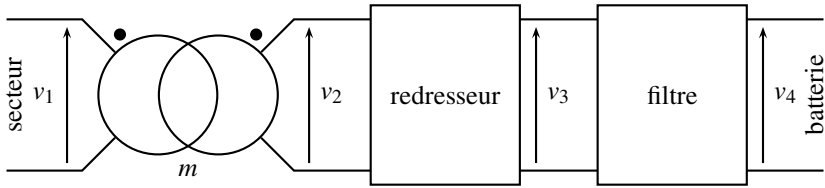
29.6 Dimensionnement d'un transformateur (PTSI) (★)

On cherche à dimensionner le transformateur utilisé pour recharger un portable. La chaîne d'énergie, logée dans un boîtier placé sur le cordon d'alimentation du portable, se compose successivement :

- de l'alimentation EDF du secteur qui délivre la tension $v_1(t) = V_0 \sin(2\pi f_0 t)$, où $f_0 = 50$ Hz et $V_0 = 240$ V,
- d'un transformateur, dont la sortie est $v_2(t) = V_{0,2} \sin(2\pi f_0 t)$, et dont le rapport de transformation est noté m ,
- d'un redresseur, montage qui délivre la valeur absolue v_3 de la tension d'entrée v_2 ,
- d'un filtre moyennneur, dont la sortie v_4 est la valeur moyenne de la tension d'entrée v_3 . La batterie du portable est branchée à la sortie, elle requiert une tension de charge constante $v_4 = 12$ V.

- 1. Que vaut $V_{0,2}$ en fonction de V_0 ?

CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE



2. Tracer le graphe de la tension $v_3(t)$.
3. Quelle est la nature du filtre utilisé entre v_3 et v_4 (passe-bas, haut, bande...)? Proposer une valeur pour sa fréquence de coupure.
4. Établir l'expression de la tension v_4 en fonction de V_0 .
5. En déduire la valeur de m .

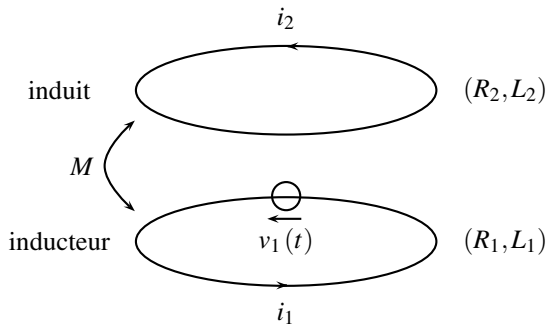
29.7 Table à induction (d'après CCP) (★)

Le chauffage du fond métallique des récipients de cuisson peut être directement réalisé au moyen de courants de Foucault induits par un champ magnétique variable.

Logé dans une table en céramique, un bobinage, nommé l'inducteur, alimenté en courant sinusoïdal génère ce champ. Le transfert d'énergie électrique s'effectue par induction mutuelle entre ce bobinage et la plaque circulaire assimilable à une spire unique fermée sur elle-même, situé au fond d'une casserole.

L'inducteur, de 5 cm de rayon, comporte 20 spires de cuivre de résistance électrique $R_1 = 1,8 \cdot 10^{-2} \Omega$ et d'autoinductance $L_1 = 30 \mu\text{H}$.

La plaque de résistance $R_2 = 8,3 \text{ m}\Omega$ et d'autoinductance $L_2 = 0,24 \mu\text{H}$, nommée l'induit, est assimilable à une spire unique refermée sur elle-même. L'inducteur est alimenté par une tension $v_1(t)$. L'ensemble plaque (induit) – inducteur se comporte comme deux circuits couplés par une mutuelle M .



1. Écrire les équations électriques relatives aux deux circuits (équations de couplage entre i_1 et i_2).
2. En déduire l'expression littérale du rapport des amplitudes complexes $\frac{I_2}{I_1}$.
3. En déduire l'expression littérale de l'impédance d'entrée complexe du système : $Z_e = \frac{V_1}{I_1}$.

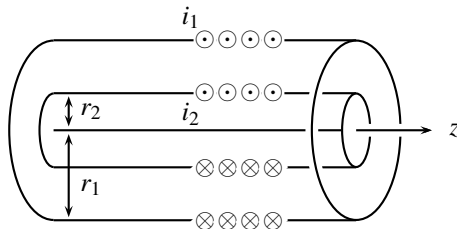
4. On choisit ω telle que $R_1 \ll L_1 \omega$ et $R_2 \ll L_2 \omega$. Simplifier les deux expressions littérales précédentes, puis effectuer le calcul numérique de leur module, sachant que l'inductance mutuelle est estimée à $M = 2 \mu\text{H}$.
5. On soulève la plaque à chauffer ; on demande un raisonnement purement qualitatif. L'amplitude du courant i_1 appelé par l'inducteur augmente-t-il ou décroît-il ?

29.8 Écrantage d'un champ magnétique (d'après CCP) (★★)

On utilisera les coordonnées cylindriques (r, θ, z) et la base locale $(\vec{u}_r, \vec{u}_\theta, \vec{u}_z)$.

On se place dans tout le problème dans l'approximation des régimes quasi permanents.

On considère deux solénoïdes Σ_1 et Σ_2 coaxiaux, d'axe (Oz) , de même longueur $L = 20 \text{ cm}$, de rayons $r_1 = 10 \text{ cm}$ et $r_2 = 5 \text{ cm}$ et comportant respectivement $N_1 = 700$ et $N_2 = 500$ spires jointives, enroulées dans le même sens :



Dans toute la suite on négligera les effets de bord ; on considèrera donc les solénoïdes comme très longs. Ces deux bobines ont pour résistance respectivement R_1 et $R_2 = 50 \Omega$. On pourra introduire les nombres de spires par unité de longueur $n_i = N_i/L$.

1. Le solénoïde Σ_1 est parcouru par un courant d'intensité i , Σ_2 étant en circuit ouvert.
- a. Exprimer le champ magnétique $\vec{B}_1$ créé dans tout l'espace.
- b. En déduire que le coefficient d'inductance L_1 de Σ_1 vaut $\mu_0 \frac{N_1^2}{L} \pi r_1^2$; donner l'expression de L_2 , inductance de Σ_2 et calculer sa valeur numérique.
2. Le solénoïde Σ_1 est alimenté par un générateur idéal de courant électromoteur $i_0(t) = I_0 \cos(\omega t)$ avec $I_0 = 1 \text{ A}$; les deux extrémités du solénoïde Σ_2 sont reliées par un fil sans résistance. On admet que le coefficient de mutuelle inductance M entre les deux solénoïdes est donné par $M = \mu_0 \frac{N_1 N_2}{L} \pi r_2^2$.
- a. Déterminer l'amplitude complexe du courant $i_2(t)$ circulant dans Σ_2 en fonction de M , L_2 et R_2 . La mettre sous la forme : $i_2 = K \frac{j \frac{\omega}{\omega_c}}{1 + j \frac{\omega}{\omega_c}} i_0$. On donnera l'expression de K en fonction de N_1 et N_2 et celle de ω_c en fonction de R_2 et L_2 .
- b. En déduire l'expression de l'amplitude complexe $\underline{B}_2$ du champ magnétique total à l'intérieur du solénoïde Σ_2 . Montrer que ce champ tend vers 0 à haute fréquence. Commenter.

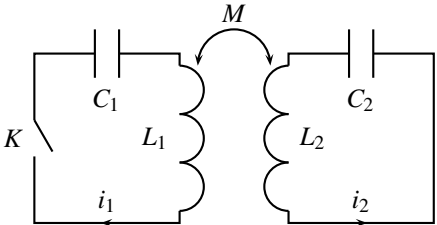
CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE

c. Application numérique : calculer ω_c ainsi que les amplitudes de i_2 et de B_2 pour une fréquence de 11 kHz. Calculer le rapport des amplitudes B_2/B_1 .

29.9 Circuits couplés par mutuelle (★★)

Un circuit LC série oscille naturellement à la pulsation $\omega_0 = \frac{1}{\sqrt{LC}}$. Cette pulsation est modifiée lorsqu'on approche un autre circuit LC , identique au premier, mais dans une configuration telle que les deux circuits deviennent couplés par mutuelle induction.

Dans le circuit suivant, le condensateur de capacité C_1 est chargé sous la tension u_0 à la date $t = 0$ où l'on ferme l'interrupteur K . On prendra dans toute la suite $C_1 = C_2 = C$ et $L_1 = L_2 = L$.



1. Que signifie, pour les lignes de champ magnétique, que les circuits soient couplés par mutuelle induction ?
2. Établir deux équations différentielles couplées sur les tensions u_{C1} et u_{C2} aux bornes des condensateurs.
3. Découpler ces équations en formant les équations sur les fonctions somme $\sigma = u_{C1} + u_{C2}$ et différence $\delta = u_{C1} - u_{C2}$. Les intégrer et en déduire les expressions des tensions aux bornes des condensateurs.

On pourra poser $\omega_1 = \frac{1}{\sqrt{C(L+M)}}$ et $\omega_2 = \frac{1}{\sqrt{C(L-M)}}$.

4. Si $M \ll L$, comparer ω_1 et ω_2 Quelle est alors l'allure du graphe de u_{C1} ? Comment s'appelle le phénomène observé ? Cette question ne requiert aucun calcul.
5. Dans le cas où $M \ll L$, montrer que ω_1 et ω_2 s'écrivent :

$$\omega_1 = \omega_0 \left(1 - \frac{M}{nL} \right) \quad \text{et} \quad \omega_2 = \omega_0 \left(1 + \frac{M}{nL} \right),$$

où n est un entier à préciser. En déduire l'expression de $u_{C1}(t)$ sous la forme d'un produit de cosinus, puis une méthode qui permette de mesurer expérimentalement le rapport M/L à l'oscilloscope, avec les mesures des périodes des phénomènes.

CORRIGÉS

29.1 Bobines en séries

1. Dans le cas (a), les champs magnétiques créés par les deux spires sont orientés de la gauche vers la droite, d’après la règle de la main droite. Les vecteurs surfaces sont orientés de même, de la gauche vers la droite. Les flux magnétiques sont donc positifs. Attendu que $\phi_{1 \rightarrow 2} = M i_1$ et que i_1 est positif, alors M est positif.

Dans le cas (b) : $M < 0$ car les champs magnétiques ont un produit scalaire négatif avec les vecteurs surfaces.

Dans le cas (c), les champs magnétiques sont orthogonaux aux vecteurs surface. Les flux sont donc nuls et $M = 0$.

2. On suppose que la bobine est assez longue pour être assimilée à un solénoïde infini (on néglige les effets de bord).

Le flux de $\vec{B}$ à travers une spire vaut : $\phi_1 = BS = \mu_0 n i S$.

À travers toutes les spires : $\phi = N \phi_1 = N \mu_0 \frac{N}{\ell} i S = L i$. D’où : $L = \frac{\mu_0 N^2 S}{\ell}$.

3. a. $\phi_1 = L i_1 + M i_2$ et de même $\phi_2 = M i_1 + L i_2$.

b. i_1 et i_2 sont pris positifs dans le sens indiqué que la figure associée au dipôle $\mathcal{D}_1$. Le flux Φ_1 à travers $\mathcal{D}_1$ est la somme des deux flux précédents : $\Phi_1 = \phi_1 + \phi_2 = (L + M)(i_1 + i_2)$. Les deux bobines sont en série et $i_1 = i_2 = i$, $\Phi_1 = 2(L + M)i$ et donc $L_1 = 2(L + M)$.

Pour Φ_2 , on remarque que $i_2 = -i_1 = -i$, $\Phi_2 = 0$ et donc $L_2 = 0$.

c. Ces valeurs inhabituelles des coefficients d’autoinductance viennent du couplage entre les deux bobines, c’est à dire de leur mutuelle. Si on éloigne les deux bobines de telle façon que $M = 0$, a lors on retrouve la loi classique d’additivité des inductances.

d. Les lignes de $\vec{B}$ divergent en sortie des bobines. Les deux bobines doivent être très proches pour que tout le champ $\vec{B}$ créé par l’une passe à travers toutes les spires de l’autre.

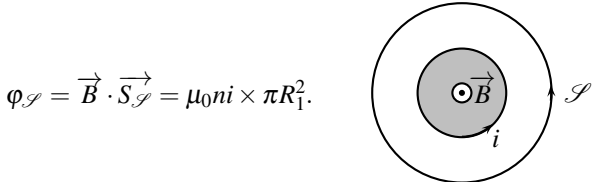
29.2 Autoinductance d’un solénoïde

L’inductance ne dépend pas de l’intensité du courant qui traverse le solénoïde. Elle ne dépend que du nombre de spires. Quand on double le nombre de spires, on double à la fois le champ magnétique propre et la surface donc le flux. L’inductance propre est donc multipliée par quatre.

29.3 Spire autour d’un solénoïde

1. L’énoncé ne précise pas l’orientation de la spire $\mathcal{S}$. On la choisit arbitrairement comme le montre le schéma. Le champ magnétique créé par le solénoïde vaut, à l’intérieur $\vec{B} = \mu_0 n i \vec{u}_z$ (zone grisée), où (Oz) est l’axe commun au solénoïde et à la spire, et est nul à l’extérieur (zone laissée en blanc). La spire $\mathcal{S}$ ne « voit » de champ magnétique que dans la zone grisée, c’est à dire sur une surface πR_1^2 . D’où un flux magnétique à travers $\mathcal{S}$:

CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE



L'intensité du courant qui traverse la solénoïde est $i = \frac{30}{1,2 + 6,7} = 3,7(5) \text{ A}$ et donc $\varphi_{\mathcal{S}} = 5,9 \cdot 10^{-6} \text{ Wb}$.
Attention aux unités légales : $R_1 = 2 \text{ cm} = 2 \cdot 10^{-2} \text{ m}$ et :

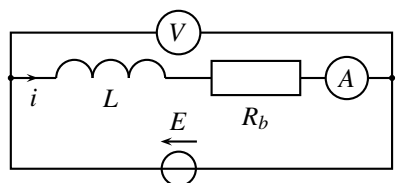
$$n = 10 \text{ spires par cm} = \frac{10 \text{ spires}}{1 \text{ cm}} = \frac{10 \text{ spires}}{10^{-2} \text{ m}} = 10^3 \text{ spires.m}^{-1}.$$

2. L'exponentielle travaille sur un nombre sans dimension et délivre un nombre sans dimension. Ainsi t et τ ont la même unité, la seconde.

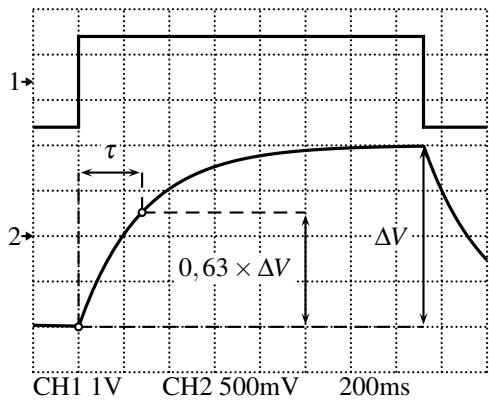
3. $\varphi_{\mathcal{S}} = \vec{B} \cdot \vec{S}_{\mathcal{S}} = \mu_0 n i(t) \pi R_1^2$ et : $e_{\mathcal{S}} = -\frac{d\varphi_{\mathcal{S}}}{dt} = -\mu_0 n \pi R_1^2 \frac{di}{dt} = \frac{\mu_0 n i_0 \pi R_1^2}{\tau} \exp\left(-\frac{t}{\tau}\right).$

29.4 Mesure d'une autoinductance

1. On câble le montage suivant, où E est un générateur de tension constante. L'équation différentielle à laquelle obéit l'intensité $i(t)$ du courant est $E = L \frac{di}{dt} + R_b i$ (on rappelle que le voltmètre a une impédance infinie et ne se laisse traverser par aucun courant). Au bout de quelques L/R_b , on arrive en régime permanent où l'intensité est décrite par la solution particulière $i = E/R$. Dans ce cas, $i =$ constante, la tension aux bornes de la bobine réelle, que mesure le voltmètre, vaut $E = R_b i$. R_b est alors le rapport de la tension mesurée au voltmètre par l'intensité du courant que traverse l'ampèremètre.



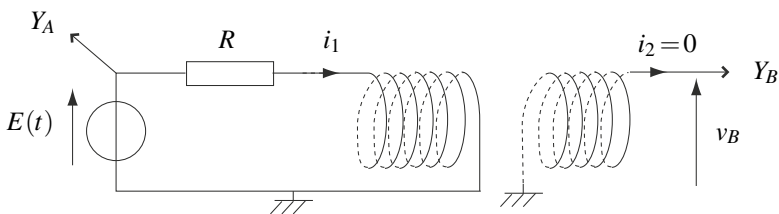
2. On néglige $R_b = 1,2 \Omega$ devant $R = 10^3 \Omega$. La fonction de transfert, entre les tensions $\underline{v}$ et $\underline{e}$, s'obtient avec la formule du diviseur de tensions : $\frac{\underline{v}}{\underline{e}} = \frac{R}{R + jL\omega} = \frac{1}{1 + j\tau\omega}$ où $\tau = \frac{L}{R}$.
On reconnaît un passe-bas du premier ordre ; on utilise donc la méthode du temps de montée à 63%. On mesure sur le graphe $\Delta V = 2 \text{ V}$ d'où $0,63 \times \Delta V = 1,26 \text{ V}$ et : $\tau = 1,4 \text{ carreau} \times 200 \text{ ms par carreau} = 280 \text{ ms} = \frac{L}{R}$. Finalement, $L = \tau R = 336 \text{ mH}$. Attendu l'imprécision graphique, on gardera comme résultat $L = 3,4 \cdot 10^2 \text{ mH}$.
Petit rappel de la méthode du temps de montée à 63% :



L'équation différentielle qui modélise l'évolution du circuit est : $(j\tau\omega + 1) \underline{v} = \underline{e}$ donc $\tau \frac{dv}{dt} + v(t) = e(t)$. En notant E_0 la valeur haute de $e(t)$, qui vaut 1 V par lecture graphique, $v(t)$ s'écrit : $v(t) = E_0 + \mu \exp\left(-\frac{t}{\tau}\right)$. Avec v_0 la condition initiale à la date du front montant de $e(t)$ ($v_0 = -1$ V par lecture graphique), $\mu = v_0 - E_0$. Dès lors : $v(t) = E_0 + (v_0 - E_0) \exp\left(-\frac{t}{\tau}\right)$. Ainsi : $v(\tau) = E_0 + (v_0 - E_0) \times 0,37 = v_0 + (E_0 - v_0) \times 0,63$. Au bout d'une durée τ , on ajoute à la valeur initiale v_0 du signal, une tension $0,63\Delta V$.

29.5 Mesure d'une inductance mutuelle

- 1. $r \ll R$.
- 2. Le schéma du montage est le suivant :



Dans la première bobine, la f.e.m. induite est $e_1(t) = -L \frac{di_1}{dt} - M \frac{di_2}{dt} = -L \frac{di_1}{dt}$, puisque $i_2(t) = 0$. De même, dans la seconde bobine, $e_2(t) = -L \frac{di_2}{dt} - M \frac{di_1}{dt} = -M \frac{di_1}{dt}$. Les équations électriques des circuits sont donc :

$$\begin{cases} E(t) = Ri_1(t) + L \frac{di_1}{dt} \\ v_B(t) = M \frac{di_1}{dt} \end{cases}$$

Sur chaque demi-période, la tension $v_B(t)$ est constante donc $\frac{di_1}{dt} = \text{cste} = v_B$. Dérivons la

CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE

première équation : $\frac{dE}{dt} = R \frac{v_B}{M}$ et écrivons cette équation dans chaque demi-période. Nous obtenons sur la première demi-période : $\frac{2\Delta e}{T} = \frac{Rv_{B1}}{M}$, et sur la deuxième demi-période : $-\frac{2\Delta e}{T} = \frac{Rv_{B2}}{M}$.

Finalement, avec $A = v_{B1} - v_{B2}$ et $\Delta e = 1,4 \text{ V}$: $M = \frac{RAT}{4\Delta e}$.

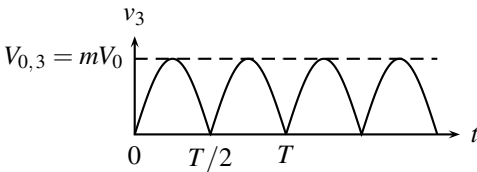
On lit la période sur l'oscillogramme : $T = 5 \times 0,2 \text{ ms} = 1 \text{ ms}$. L'application numérique donne $M = 17,9 \text{ A}$ avec M en mH et A en V :

$d \text{ (cm)}$	4	5	6	7	8	10	12	16	20
$M \text{ (mH)}$	0,77	0,59	0,46	0,38	0,30	0,21	0,14	0,075	0,043

Plus les bobines sont éloignées, plus leur mutuelle M diminue.

29.6 Dimensionnement d'un transformateur (PTSI)

- $V_{0,2} = mV_0$.
- $v_3(t)$ est la valeur absolue de $v_2(t) = mV_0 \sin(2\pi ft)$:



3. Le signal de tension v_3 , de période $T/2$, est composé d'une valeur moyenne, à conserver, d'un fondamental à la pulsation $\omega_1 = \frac{2\pi}{\text{période}} = \frac{4\pi}{T} = 4\pi f_0$ et d'harmoniques de pulsations multiples entières du fondamental, $\omega_n = n\omega_1$. Pour ne garder que la valeur moyenne, il faut filtrer les hautes fréquences. On utilise donc un passe-bas de fréquence de coupure très inférieure à $f_1 = \frac{\omega_1}{2\pi} = 2f_0 = 100 \text{ Hz}$, par exemple inférieure à 10 Hz.

4. Calculons sur une période la valeur moyenne du signal v_3 , de période $T/2$ (d'où le $1/\text{période} = 2/T$ en début de formule) :

$$\langle v_3 \rangle = \frac{2}{T} \int_0^{T/2} mV_0 \sin(2\pi f_0 t) dt = \frac{2mV_0}{T} \left[-\frac{\cos(2\pi f_0 t)}{2\pi f_0} \right]_0^{T/2} = \frac{2mV_0}{\pi} = v_4.$$

5. $m = \frac{\pi v_4}{2V_0} = 7,8 \cdot 10^{-2} < 1$, il y a plus de spires au primaire qu'au secondaire.

29.7 Table à induction (d'après CCP)

$$1. \begin{cases} v_1 = R_1 i_1 + L_1 \frac{di_1}{dt} + M \frac{di_2}{dt} \\ 0 = R_2 i_2 + L_2 \frac{di_2}{dt} + M \frac{di_1}{dt} \end{cases}$$

2.
$$\begin{cases} \underline{v}_1 = (R_1 + iL_1\omega)\underline{i}_1 + iM\omega\underline{i}_2 \\ 0 = (R_2 + iL_2\omega)\underline{i}_2 + iM\omega\underline{i}_1 \end{cases} \quad \text{donc} \quad \frac{L_2}{L_1} = -\frac{iM\omega}{R_2 + iL_2\omega}.$$

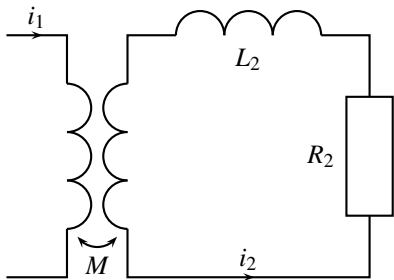
3. $\underline{v}_1 = (R_1 + iL_1\omega)\underline{i}_1 + iM\omega\left(-\frac{iM\omega}{R_2 + iL_2\omega}\underline{i}_1\right)$, donc $\underline{Z}_e = R_1 + iL_1\omega + \frac{(M\omega)^2}{R_2 + iL_2\omega}.$

4.
$$\begin{cases} R_1 \ll L_1\omega \\ R_2 \ll L_2\omega \end{cases} \quad \text{donc} \quad \begin{cases} \frac{L_2}{L_1} = -\frac{M}{L_2} \\ \underline{Z}_e = iL_1\omega\left(1 - \frac{M^2}{L_1L_2}\right) \end{cases} \quad \text{donc} \quad \begin{cases} \frac{L_2}{L_1} = -8,3 \\ |\underline{Z}_e| = 2,1\,\Omega \end{cases}$$

5. Le champ magnétique, créé par l'inducteur et vu par la plaque, diminue lorsqu'on éloigne la plaque. Le flux de $\vec{B}$ à travers la plaque diminue et donc M diminue. Alors Z_e augmente et, pour une même tension d'alimentation, le courant décroît.

29.8 Écrantage d'un champ magnétique (d'après CCP)

1. a.
$$\vec{B}_1 = \begin{cases} \mu_0 n_1 i \vec{u}_z^> & (r < r_1) \\ \vec{0} & (r > r_1) \end{cases}$$
- b. Le flux de $\vec{B}_1$ à travers les N_1 spires de Σ_1 est : $\phi_1 = N_1 \mu_0 n_1 i \pi r_1^2 = L_1 i$ donc $L_1 = \mu_0 \frac{N_1^2}{L} \pi r_1^2$. De même : $L_2 = \mu_0 \frac{N_2^2}{L} \pi r_2^2 = 1,2 \cdot 10^{-2}$ H.
- c. Le flux de $\vec{B}_1$ à travers les N_2 spires de Σ_2 est : $\phi_{1 \rightarrow 2} = N_2 \mu_0 n_1 i \pi r_2^2 = M i$ donc $M = \mu_0 \frac{N_1 N_2}{L} \pi r_2^2$.
2. a. Le système est équivalent à :



D'où : $L_2 \frac{di_2}{dt} + M \frac{di_1}{dt} + Ri_2 = 0$. En régime harmonique à la pulsation ω : $(R_2 + jL_2\omega)\underline{i}_2 + jM\omega\underline{i}_1 = 0$ donc $\underline{i}_2 = -\frac{\frac{M}{L_2} j \frac{L_2}{R_2} \omega}{1 + j \frac{L_2}{R_2} \omega} \underline{i}_1$. Ainsi $\omega_c = \frac{R_2}{L_2}$ et $K = -\frac{M}{L_2} = -\frac{N_1}{N_2}$.

Remarque physique : $i_1(t)$ crée $\vec{B}_1(t)$. Donc Σ_2 devient un circuit fixe dans $\vec{B}_1$ variable, il reçoit un flux de Σ_1 ($\phi_{1 \rightarrow 2} = M i_1$) ; une f.é.m. $e_2 = -M di_1/dt$ est induite. Σ_2 est fermé et conducteur donc un courant induit i_2 circule. i_2 crée un champ magnétique propre qui a un flux à travers Σ_2 ($\phi_{2 \rightarrow 2} = L_2 i_2$). Ce flux propre crée une seconde f.é.m. induite $e'_2 = -L_2 di_2/dt$.

CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE

b. Le champ magnétique total $\vec{B}_2$ est créé par i_1 et i_2 : $\vec{B}_2 = (\mu_0 n_1 i_1 + \mu_0 n_2 i_2) \vec{u}_z$. En régime harmonique à la pulsation ω et en tenant compte de l'expression de i_2 :

$$\underline{B}_2 = \frac{\mu_0}{L} \left(N_1 \underline{i}_1 - N_2 \frac{N_1}{N_2} \frac{j \frac{\omega}{\omega_c}}{1 + j \frac{\omega}{\omega_c}} \underline{i}_1 \right) = \frac{\mu_0}{L} N_1 \left(1 - \frac{j \frac{\omega}{\omega_c}}{1 + j \frac{\omega}{\omega_c}} \right) \underline{i}_1 \xrightarrow{\omega \gg \omega_c} 0.$$

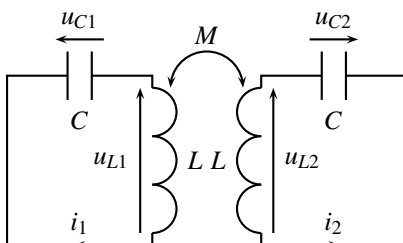
Σ_2 réagit par induction en créant un champ magnétique qui s'oppose à celui de Σ_1 (loi de Lenz).

c. $\omega_c = 4,0 \cdot 10^3 \text{ rad.s}^{-1}$, $i_2 = |\underline{i}_2| = 1,40 \text{ A}$, $B_2 = |\underline{B}_2| = 2,57 \cdot 10^{-4} \text{ T}$ et $\frac{B_1}{B_2} = 17,1$.

29.9 Circuits couplés par mutuelle

1. Si les circuits sont couplés, alors des lignes du champ magnétique $\vec{B}_1$, créé par i_1 , traversent le circuit 2 ; de même, celle de $\vec{B}_2$, créé par i_2 , traversent le circuit 1.

2. Les tensions doivent être dessinées en convention récepteur.



La loi des mailles dans le circuit 1 impose $u_{L1} + u_{C1} = 0$. De plus, $u_{L1} = L \frac{di_1}{dt} + M \frac{di_2}{dt}$, $i_1 = C \frac{du_{C1}}{dt}$ et $i_2 = C \frac{du_{C2}}{dt}$. Ainsi : $LC \frac{d^2 u_{C1}}{dt^2} + MC \frac{d^2 u_{C2}}{dt^2} + u_{C1} = 0$. De même avec la maille 2 : $LC \frac{d^2 u_{C2}}{dt^2} + MC \frac{d^2 u_{C1}}{dt^2} + u_{C2} = 0$.

3. Sommons ces deux équations : $C(L+M) \frac{d^2}{dt^2} (u_{C1} + u_{C2}) + (u_{C1} + u_{C2}) = 0$ soit $\frac{1}{\omega_1^2} \frac{d^2 \sigma}{dt^2} + \sigma = 0$, qui s'intègre en : $\sigma(t) = \sigma_0 \cos(\omega_1 t) + \sigma'_0 \sin(\omega_1 t)$.

Et pour la différence : $C(L-M) \frac{d^2}{dt^2} (u_{C1} - u_{C2}) + (u_{C1} - u_{C2}) = 0$ soit $\frac{1}{\omega_2^2} \frac{d^2 \delta}{dt^2} + \delta = 0$, qui s'intègre en : $\delta(t) = \delta_0 \cos(\omega_2 t) + \delta'_0 \sin(\omega_2 t)$.

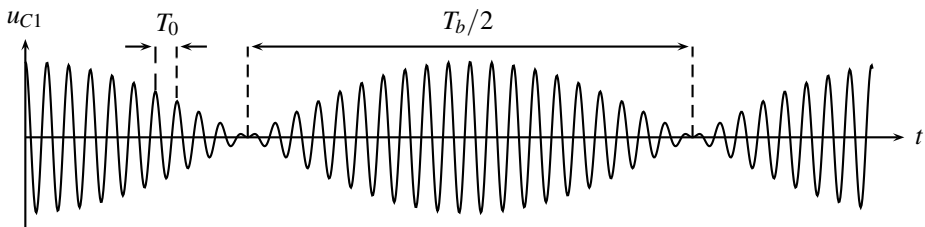
Les tensions aux bornes des condensateurs sont continues donc les conditions initiales sont $u_{C1}(0) = u_0$ et $u_{C2}(0) = 0$. Ainsi : $\sigma_0 = \delta_0 = u_0$.

Les intensités des courants qui traversent les bobines sont continues donc les conditions initiales sont $\left. \frac{du_{C1}}{dt} \right|_0 = \frac{i_1(0)}{C} = 0$ et $\left. \frac{du_{C2}}{dt} \right|_0 = \frac{i_2(0)}{C} = 0$. Ainsi : $\sigma'_0 = \delta'_0 = 0$. Finalement :

$$u_{C1}(t) = \frac{\sigma(t) + \delta(t)}{2} = \frac{u_0}{2} (\cos(\omega_1 t) + \cos(\omega_2 t)), \text{ et}$$

$$u_{C2}(t) = \frac{\sigma(t) - \delta(t)}{2} = \frac{u_0}{2} (\cos(\omega_1 t) - \cos(\omega_2 t)).$$

4. Si $M \ll L$ alors $\omega_1 \approx \omega_2$. Le signal u_{C1} est la somme de deux signaux sinusoïdaux de pulsations proches, on observe des battements :



$$5. \omega_0 = \frac{1}{\sqrt{LC}} \text{ et : } \begin{cases} \omega_1 = \frac{1}{\sqrt{C(L+M)}} = \frac{1}{\sqrt{LC(1+\frac{M}{L})}} \\ \omega_2 = \omega_0 \left(1 + \frac{M}{L}\right)^{-1/2} = \omega_0 \left(1 - \frac{M}{2L} + o\left(\frac{M}{L}\right)\right). \end{cases}$$

Au premier ordre en $\frac{M}{L}$, $\omega_1 = \omega_0 \left(1 - \frac{M}{2L}\right)$. De même, au premier ordre en $\frac{M}{L}$, $\omega_2 = \omega_0 \left(1 + \frac{M}{2L}\right)$. Et la tension u_{C1} devient :

$$\begin{aligned} u_{C1}(t) &= \frac{u_0}{2} (\cos(\omega_1 t) + \cos(\omega_2 t)) = u_0 \cos\left(\frac{\omega_1 + \omega_2}{2} t\right) \cos\left(\frac{\omega_1 - \omega_2}{2} t\right) \\ &= u_0 \cos(\omega_0 t) \cos\left(-\omega_0 \frac{M}{2L} t\right). \end{aligned}$$

Le second cosinus a une pulsation $\omega_b = \omega_0 \frac{M}{2L}$ beaucoup plus faible que ω_0 . On lit sur le graphe les valeurs de $T_0 = \frac{2\pi}{\omega_0}$ et $T_b = \frac{2\pi}{\omega_b}$. On en déduit $\frac{M}{L} = \frac{2T_0}{T_b}$.

CHAPITRE 29 – CIRCUIT FIXE DANS UN CHAMP MAGNÉTIQUE VARIABLE

Circuit mobile dans un champ magnétique stationnaire

30

Dans ce chapitre on étudie des phénomènes d'induction dans des circuits mobiles dans un champ magnétique variable. Les dispositifs modèles présentés fonctionnent soit en générateur, quand ils transforment une puissance mécanique en puissance électrique, soit en moteur quand ils transforment une puissance électrique en puissance mécanique.

1 Conversion de puissance mécanique en puissance électrique

1.1 Rails de Laplace générateurs

a) Présentation

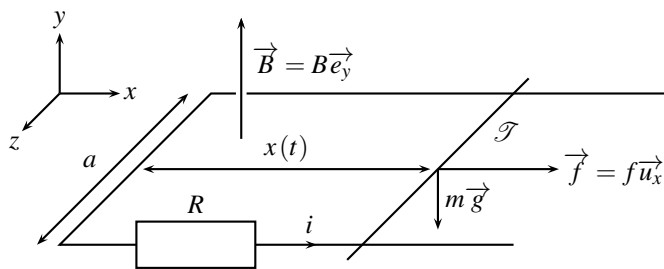


Figure 30.1 – Rails de Laplace générateurs.

Une tige $\mathcal{T}$, de masse m , conductrice, glisse sans frottement sur deux rails conducteurs, à la vitesse $\vec{v} = v(t)\vec{u}_x$. Elle est tirée par une force $\vec{f} = f\vec{u}_x$ constante. $\mathcal{T}$ reste toujours parallèle à $\vec{u}_z$ lors de son mouvement, comme représenté sur le schéma. Les rails sont dans le plan horizontal (Oxz) . L'ensemble est plongé dans un champ magnétique uniforme et stationnaire $\vec{B} = B\vec{u}_y$, orthogonal au plan des rails.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

b) Analyse physique

La mise en équation d'un problème d'électromécanique doit toujours être précédée d'une analyse qualitative :

- la tige $\mathcal{T}$ est mise en mouvement par la force $\vec{f}$,
- $\mathcal{T}$ est alors mobile dans le champ $\vec{B}$ stationnaire : il apparaît dans ce conducteur une f.é.m. induite e ,
- le circuit, constitué des rails et de $\mathcal{T}$, est fermé et conducteur : un courant induit i circule,
- i et $\vec{B}$ créent une force de Laplace $\vec{f}_L$.

On verra que $\vec{f}_L$ s'oppose au mouvement de $\mathcal{T}$. C'est une manifestation de la loi de Lenz : le phénomène d'induction provoque l'apparition d'une force qui s'oppose au mouvement de la barre qui est le point de départ du raisonnement.

c) Choix des orientations

Il faut orienter le circuit électrique constitué des rails et de la tige $\mathcal{T}$, c'est-à-dire le sens positif du courant. Ce choix est arbitraire, mais il faut s'y tenir jusqu'à la fin, en restant cohérent. Si une orientation différente est choisie, tous les résultats intermédiaires changent de signe, mais pas le résultat final.

On choisit arbitrairement le sens proposé par le schéma, afin que le circuit présente un vecteur surface $\vec{S}$ parallèle à $\vec{B}$ et de même sens.

L'analyse physique a fait apparaître deux grandeurs qu'il convient d'établir : la f.é.m. induite e et la force de Laplace $\vec{f}_L$.

d) F.é.m. induite et équation électrique

Le vecteur surface $\vec{S}$ du circuit est orienté par i et la règle de la main droite ; il est donc porté par $\vec{u}_y$: $\vec{S} = S\vec{u}_y$ où $S = ax(t)$ est la surface du circuit. Le flux de $\vec{B}$ à travers le circuit est :

$$\varphi = \vec{B} \cdot \vec{S} = (B\vec{u}_y) \cdot (ax(t)\vec{u}_y) = Bax(t).$$

Attendu que le circuit coupe des lignes de $\vec{B}$ dans son déplacement, la f.é.m. induite est donnée par la loi de Faraday :

$$e = -\frac{d\varphi}{dt} = -Ba\frac{dx}{dt} = -Bav(t).$$



Comment visualiser que la tige $\mathcal{T}$ coupe des lignes de champ magnétique ? Un champ magnétique $\vec{B}$ existe en tout point du circuit, quelques uns sont représentés sur le schéma de la figure 30.2. L'ensemble forme un champ de vecteur, analogue à des épis de blé dans un champ agricole. Si l'on imagine que la tige $\mathcal{T}$ est « trançante », les vecteurs $\vec{B}$ sont « coupés » par $\mathcal{T}$ dans son mouvement, de façon analogue au blé moissonné.

Soit R la résistance des rails et de la tige $\mathcal{T}$, supposée constante quelle que soit la position de $\mathcal{T}$. En orientant bien e dans le même sens que i , on peut dessiner le circuit électrique de la

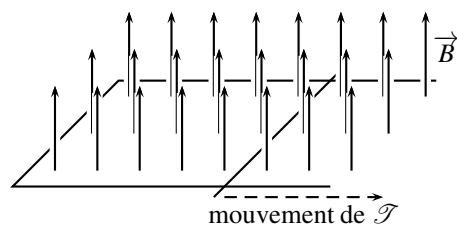


Figure 30.2 – Justification de la loi de Faraday.

figure 30.3. La loi des mailles est l'équation électrique (EE) du circuit :

$$e(t) - Ri(t) = 0 \text{ (EE)}$$

$$\text{d'où } i(t) = \frac{e(t)}{R} = -\frac{Bav(t)}{R}.$$

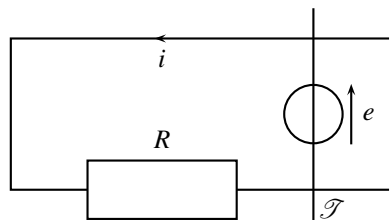


Figure 30.3 – Schéma électrique équivalent.

Remarque

Le flux φ n'est pas le seul flux magnétique à travers le circuit. En effet, le courant i crée un champ magnétique propre $\vec{B}_p$, qui crée un flux φ_p à travers le circuit. Le coefficient d'auto-inductance L permet d'écrire $\varphi_p = Li$. Le flux total, qui mène à la f.é.m. induite, est donc la somme de ces deux flux : $\varphi_{tot} = \varphi + \varphi_p$. On considère donc, dans cette première approche, que la f.é.m. d'auto-induction est négligeable devant la f.é.m. due au champ magnétique extérieur $\vec{B}$. Dans le cas du dispositif des rails de Laplace, cette approximation est tout à fait justifiée.

e) Force de Laplace et équation mécanique

La force de Laplace $\vec{f}_L$ qui s'exerce sur la tige $\mathcal{T}$ est due à la présence simultanée du courant et du champ magnétique. Elle a été calculée au chapitre 27 et son expression est :

$$\vec{f}_L = i(t)Ba\vec{u}_x.$$



Une erreur classique consiste à oublier que l'intensité i dans le circuit dépend du temps.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

La tige $\mathcal{T}$ est soumise à :

- la force de Laplace $\vec{f}_L = i(t) a B \vec{u}_x$,
- la force de traction $\vec{f} = f \vec{u}_x$,
- le poids $m \vec{g} = -mg \vec{u}_y$,
- la réaction des rails, normale au déplacement car $\mathcal{T}$ glisse sans frottement : $\vec{R} = R \vec{u}_y$.

Le principe fondamental de la dynamique, appliqué à la tige $\mathcal{T}$, en projection sur $\vec{u}_x$, mène à l'équation mécanique (EM) :

$$m \frac{dv}{dt} = i(t) a B + f \quad (EM).$$

f) Établissement de la vitesse

L'équation électrique (EE) et l'équation mécanique (EM) sont couplées : la vitesse $v(t)$ intervient dans (EE), l'intensité du courant $i(t)$ dans (EM). Il faut les découpler pour établir les expressions de $v(t)$ et de $i(t)$. Pour $v(t)$:

$$\begin{cases} m \frac{dv}{dt} = i(t) a B + f \\ e(t) = -B a v(t) = R i(t) \end{cases} \quad \text{implique} \quad m \frac{dv}{dt} = -\frac{(Ba)^2}{R} v(t) + f.$$

Le premier enseignement de cette équation, est que la force de Laplace s'écrit *in fine* sous la forme d'une **force de frottement fluide** en $-\text{constante} \times v$. Comme indiqué dans l'analyse physique préalable, la force de Laplace s'oppose ainsi au déplacement de la tige, conformément à la **loi de Lenz**. On reviendra sur cet aspect dans le paragraphe 1.2

De plus, le système est décrit par une équation différentielle du premier ordre qu'on écrit sous forme canonique :

$$\frac{mR}{(Ba)^2} \frac{dv}{dt} + v(t) = \frac{Rf}{(Ba)^2} \quad \text{soit} \quad \tau \frac{dv}{dt} + v(t) = \frac{Rf}{(Ba)^2},$$

où $\tau = \frac{mR}{(Ba)^2}$ est la constante de temps du système.

Si par exemple, la tige est initialement au repos, $v(0) = 0$, la solution du couple {équation différentielle, condition initiale} est :

$$v(t) = \frac{Rf}{(Ba)^2} \left(1 - \exp\left(-\frac{t}{\tau}\right) \right).$$

Quant à l'intensité du courant $i(t)$, l'équation électrique (EE) précise sa valeur :

$$i(t) = -\frac{B a v(t)}{R} = -\frac{f}{Ba} \left(1 - \exp\left(-\frac{t}{\tau}\right) \right).$$

g) Bilan de puissance

La méthode pour établir un bilan de puissance d'un système électromécanique est systématique. L'équation électrique doit être multipliée par i et l'équation mécanique par v :

$$\begin{cases} (EE) \times i \\ (EM) \times v \end{cases} \Rightarrow \begin{cases} ei = Ri^2 \\ m \frac{dv}{dt} v = f v + f_L v \end{cases} \text{ soit } \begin{cases} -Bavi = Ri^2 \\ \frac{d}{dt} \left(\frac{1}{2} m v^2 \right) = f v + iaBv \end{cases}$$

Le même terme de couplage $Bavi$ apparaît dans les deux équations. En l'éliminant on obtient :

$$f v = \frac{d}{dt} \left(\frac{1}{2} m v^2 \right) + Ri^2.$$

Cette dernière équation explicite les transferts de puissance. La force $\vec{f}$ fournit une puissance $\vec{f} \cdot \vec{v} = f v$ à la tige $\mathcal{T}$, qui se répartit en :

- puissance cinétique $\frac{d}{dt} \left(\frac{1}{2} m v^2 \right)$, dérivée de l'énergie cinétique de $\mathcal{T}$,
- puissance de l'effet Joule Ri^2 .

La création du courant d'intensité $i(t)$ justifie l'adjectif de générateur donné aux rails de Laplace : une puissance mécanique se transforme en une puissance électrique, ici modélisée par de l'effet Joule.

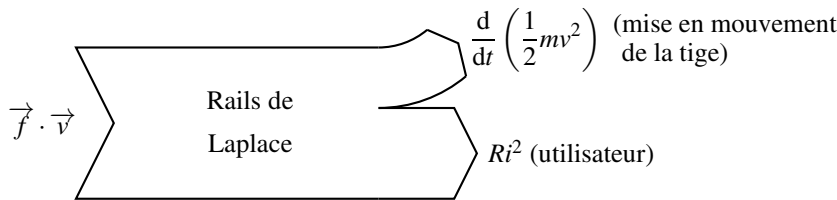


Figure 30.4 – Transferts de puissance pour des rails de Laplace générateurs.

Remarque

En regime permanent, lorsque v et i sont constants, le bilan de puissance devient $f v = Ri^2$. Toute la puissance mécanique en entrée est convertie en puissance électrique disponible pour l'utilisateur.

h) Relation des puissances

Que représente le terme de couplage $Bavi$? C'est à la fois :

- la puissance de la f.é.m. : $\mathcal{P}_{fém} = ei = -Bavi$,
- l'opposé de la puissance de la force de Laplace : $\mathcal{P}_L = \vec{f}_L \cdot \vec{v} = iab\vec{u}_y \cdot v\vec{u}_y = Bavi$.

Le fait que ces deux puissances soient égales et opposées est un fait général que l'on admettra.

Pour un circuit mobile dans un champ magnétique stationnaire, la somme de la puissance de Laplace et de la f.é.m. induite est nulle : $\mathcal{P}_{fém} + \mathcal{P}_L = 0$.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

1.2 Freinage par induction

a) Application de la loi de Lenz

Dans l'exemple précédent, la force de Laplace se met sous la forme :

$$\vec{f}_L = -\frac{(Ba)^2}{R} \vec{v}.$$

Le coefficient devant la vitesse est négatif car $(Ba)^2$ est positif et une résistance électrique R est toujours positive. La force de Laplace est de direction opposée à celle de la vitesse, c'est une **force de freinage**.

Comme on l'a expliqué, il s'agit d'une manifestation de la loi de Lenz : la force magnétique qui apparaît à cause du phénomène d'induction s'oppose au mouvement du conducteur qui a été la cause du phénomène d'induction. Ainsi :

Dans tous les dispositifs où il y a conversion de puissance mécanique en puissance électrique, l'action mécanique de Laplace est une action de freinage.



Cette action prend dans presque tous les cas une forme mathématique analogue à une action de frottements fluides.

b) Application : freinage électromagnétique

On en déduit que les phénomènes d'induction ont, par exemple dans le domaine de la traction automobile, deux utilisations importantes :

- le freinage par induction,
- la récupération d'énergie lors de ce freinage, afin de convertir l'énergie cinétique du véhicule en énergie électrique.

Le freinage électromagnétique est utilisé sur le TGV et sur des camions.

c) Courants de Foucault

Lorsque le conducteur en mouvement n'est plus filiforme, mais est formé d'un bloc métallique, si la modélisation des phénomènes d'induction, développée dans ce cours, n'est plus valable dans le détail, les phénomènes physiques restent les mêmes. Par exemple, lorsqu'on applique un champ magnétique $\vec{B}$ à une roue d'automobile, constituée d'un bloc métallique massif :

- le roue devient un circuit mobile dans un champ magnétique stationnaire, des f.é.m. sont induites,
- la roue étant conductrice, des courants sont induits. Ces courants, répartis dans tout le volume du conducteur, sont nommés **courants de Foucault**, *eddy currents* en anglais.
- Les courants de Foucault créent, avec le champ magnétique, des efforts de Laplace qui s'opposent au mouvement de la roue (loi de Lenz) : la roue est freinée.

La modélisation des phénomènes d'induction dans le volume du conducteur relève du cours de seconde année.

1.3 Alternateur

a) Présentation

Un **alternateur** sert à transformer une puissance mécanique en une puissance électrique. Ce dispositif est par exemple utilisé sur les vélos, dont une roue entraîne en rotation l'alternateur qui alimente des ampoules ou une batterie. On peut aussi citer les centrales électriques, où l'alternateur est entraîné par de la vapeur (centrales thermiques), de l'eau sous pression (centrales nucléaires françaises), de l'eau en écoulement (centrales hydrauliques). Plusieurs technologies existent, seul le principe général de l'alternateur est présenté ici.

L'alternateur est modélisé par une spire rectangulaire $\mathcal{S}$, de surface $a \times b$, conductrice de résistance électrique R .

Cette spire qui constitue un rotor, est en liaison pivot d'axe (Oy) par rapport à un stator. Elle est en rotation autour de l'axe (Oy) à la vitesse angulaire ω constante. Son moment d'inertie par rapport à l'axe (Oy) est noté J .

$\mathcal{S}$ est plongée dans un champ magnétique uniforme et stationnaire $\vec{B}$ perpendiculaire à (Oy) , créé par un environnement extérieur.

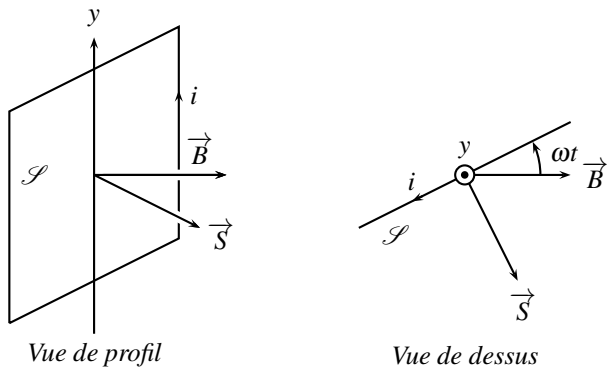


Figure 30.5 – Schéma de principe d'un alternateur.

b) Analyse physique

- La spire $\mathcal{S}$ est entraînée en rotation,
- $\mathcal{S}$ est alors un circuit mobile dans le champ magnétique stationnaire $\vec{B}$: il apparaît dans ce conducteur une f.é.m. induite e ,
- $\mathcal{S}$ est un circuit électrique fermé et conducteur : un courant induit i circule,
- i et $\vec{B}$ créent un moment des forces de Laplace qui s'oppose à la rotation de la spire d'après la loi de Lenz.



La loi de Lenz marque la fin de l'analyse physique. Un effet mécanique (le moment des forces de Laplace) s'oppose à un autre effet mécanique (la rotation de la spire).

c) Choix des orientations

Le sens du courant est arbitrairement choisi comme sur la figure ci-dessus. Ce sens impose le sens du vecteur surface $\vec{S}$ d'après la règle de la main droite.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

d) F.é.m. induite et équation électrique

Attendu que la spire $\mathcal{S}$ tourne, son orientation par rapport à $\vec{B}$ varie. L'angle entre le vecteur $\vec{S}$ et le vecteur $\vec{B}$ est égal à $\frac{\pi}{2} - \omega t$ (voir figure 30.5). Le flux magnétique à travers $\mathcal{S}$ est :

$$\varphi(t) = \vec{B} \cdot \vec{S} = BS \cos\left(\frac{\pi}{2} - \omega t\right) = BS \sin(\omega t).$$

Ce flux varie dans le temps. De plus la spire coupe des lignes de champ magnétique lors de sa rotation, donc la f.é.m. induite dans la spire vaut :

$$e(t) = -\frac{d\varphi}{dt} = -BS\omega \cos(\omega t).$$

Cette f.é.m. fait circuler un courant d'intensité i , qui crée son propre champ magnétique. Le flux de ce champ, flux propre φ_p , s'exprime à l'aide du coefficient d'autoinductance L : $\varphi_p = Li$. La constitution des alternateurs est telle que la variation de ce flux propre n'est, en général, pas négligeable devant la variation du flux du champ magnétique $\vec{B}$ extérieur. Ce phénomène d'auto-induction est pris en compte dans le circuit électrique équivalent par l'inductance propre L :

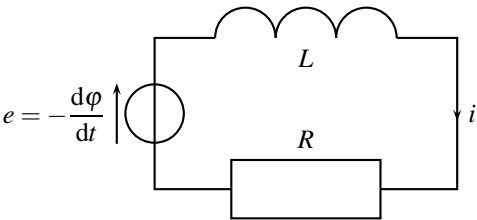


Figure 30.6 – Schéma électrique équivalent de l'alternateur.

La loi des mailles constitue l'équation électrique (EE) :

$$-\frac{d\varphi}{dt} - L\frac{di}{dt} - Ri(t) = 0 \quad (EE).$$

e) Moment des forces de Laplace et équation mécanique

Le moment résultant des forces de Laplace qui s'exerce sur la spire $\mathcal{S}$ est :

$$\vec{\Gamma} = \vec{\mathcal{M}} \wedge \vec{B} = i\vec{S} \wedge \vec{B},$$

où $\vec{\mathcal{M}} = i\vec{S}$ est le moment magnétique de la spire. Ainsi :

$$\vec{\Gamma}(t) = iSB \sin\left(\frac{\pi}{2} - \omega t\right) \vec{u}_y = iSB \cos(\omega t) \vec{u}_y.$$

La spire $\mathcal{S}$ est entraînée en rotation sous l'effet d'un couple extérieur $\vec{\Gamma}_{\text{ext}} = \Gamma_{\text{ext}} \vec{u}_y$, dû, par exemple, à la roue de vélo qui entraîne l'alternateur dans l'exemple introductif. Ce couple

est nécessaire pour avoir une vitesse de rotation constante malgré le couple magnétique et les frottements de la liaison pivot de la spire, frottements que l'on négligera en supposant la liaison pivot parfaite. La loi scalaire du moment cinétique, appliquée à la spire $\mathcal{S}$, autour de l'axe (Oy) , mène, puisque $\omega = \text{constante}$, à l'équation mécanique (EM) :

$$J \frac{d\omega}{dt} = 0 = iSB \cos(\omega t) + \Gamma_{\text{ext}} \quad (EM).$$

f) Calcul du régime permanent sinusoïdal

En combinant les équations électrique (EE) et mécanique (EM) on montre que l'intensité obéit à l'équation électrique, écrite sous forme canonique :

$$\frac{L}{R} \frac{di}{dt} + i(t) = -\frac{BS\omega}{R} \cos(\omega t).$$

La constante de temps du circuit est $\tau = \frac{L}{R}$. Au bout de 5τ , durée du régime transitoire, un régime permanent sinusoïdal à la pulsation ω s'établit. Le calcul de la solution particulière sinusoïdale de l'équation est fait dans l'annexe mathématique page 1068 et conduit au résultat :

$$i(t) = -\frac{\omega}{1 + (\tau\omega)^2} \frac{BS}{R} (\cos(\omega t) + \omega\tau \sin(\omega t)).$$

L'équation mécanique donne l'expression du couple extérieur :

$$\Gamma_{\text{ext}} = -i(t)SB \cos(\omega t) = \frac{\omega}{1 + (\tau\omega)^2} \frac{(BS)^2}{R} (\cos^2(\omega t) + \omega\tau \sin(\omega t) \cos(\omega t)).$$

La moyenne dans le temps de $\cos^2(\omega t)$ est égale à $\frac{1}{2}$ et la moyenne de $\cos(\omega t) \sin(\omega t)$ est nulle. Ainsi la moyenne dans le temps du couple Γ_{ext} est :

$$\langle \Gamma_{\text{ext}} \rangle = \frac{\omega}{1 + (\tau\omega)^2} \frac{(BS)^2}{2R}.$$

Quel est le bilan de puissance ? L'alternateur délivre en sortie une puissance électrique moyenne :

$$\langle Ri^2 \rangle = \frac{\omega^2}{1 + (\tau\omega)^2} \frac{(BS)^2}{2R}.$$

Il faut donc lui fournir en entrée une puissance mécanique identique. C'est la puissance mécanique du couple $\vec{\Gamma}_{\text{ext}}$:

$$\langle \Gamma_{\text{ext}} \omega \rangle = \frac{\omega^2}{1 + (\tau\omega)^2} \frac{(BS)^2}{2R}.$$

En régime établi, l'intégralité de la puissance mécanique injectée dans l'alternateur est convertie en puissance électrique.

2 Conversion de puissance électrique en puissance mécanique

2.1 Rails de Laplace moteurs

a) Présentation

Le dispositif des rails de Laplace peut aussi être utilisé pour convertir une puissance électrique en puissance mécanique. Les conditions d'utilisation sont cependant différentes : un générateur impose un échelon de tension, c'est-à-dire une tension E positive pour $t > 0$ et une tension nulle pour $t < 0$.

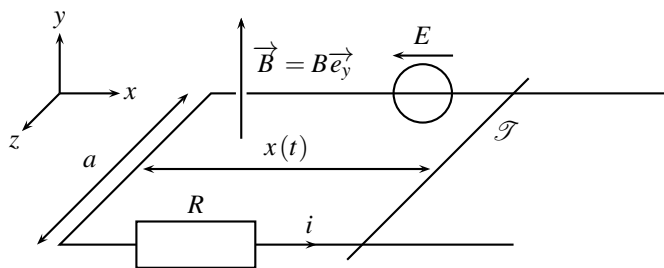


Figure 30.7 – Rails de Laplace moteurs.

b) Analyse physique

- Le générateur impose la tension E et fait circuler un courant d'intensité i dans le circuit formé par les rails la tige $\mathcal{T}$,
- i et $\vec{B}$ créent une force de Laplace $\vec{f}_L$,
- la tige $\mathcal{T}$ est mise en mouvement par la force de Laplace $\vec{f}_L$,
- $\mathcal{T}$ devient un circuit mobile dans $\vec{B}$ stationnaire : il apparaît dans ce conducteur une f.é.m. induite e , qui s'oppose à la tension E du générateur (loi de Lenz).



La loi de Lenz marque encore la fin de l'analyse physique. Dans le présent cas, une grandeur électrique, la f.é.m. induite, s'oppose bien à un phénomène électrique, la tension E du générateur.

c) F.é.m. induite et équation électrique

Le sens positif choisi pour le courant est indiqué sur la figure ci-dessus.

Le calcul de la f.é.m. et la justification de la validité de la loi de Faraday ne sont pas modifiés par rapport à l'étude du début du chapitre et : $e(t) = -Bav(t)$.

En orientant la f.é.m. induite e dans le même sens que i , on peut représenter le schéma électrique équivalent au circuit, sur la figure 30.8 suivante.

La loi des mailles mène à l'équation électrique (EE) du circuit :

$$e(t) + E - Ri(t) = 0 \quad (EE) \quad \text{d'où} \quad i(t) = \frac{e(t) + E}{R} = \frac{E - Bav(t)}{R}.$$

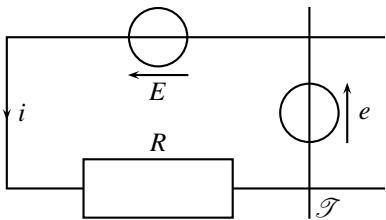


Figure 30.8 – Schéma électrique équivalent.

d) Force de Laplace et équation mécanique

La force de Laplace $\vec{f}_L$ qui s'exerce sur la tige $\mathcal{T}$ est la même que dans l'étude du début du chapitre et : $\vec{f}_L = iaB\vec{u}_x$.

Le principe fondamental de la dynamique, appliqué à la tige $\mathcal{T}$, en projection sur $\vec{u}_x$, mène à l'équation mécanique (EM) :

$$m \frac{dv}{dt} = i(t) a B \quad (EM).$$

e) Établissement de la vitesse

En injectant $i(t) = \frac{E - Bav(t)}{R}$ dans l'équation (EM) on trouve une équation différentielle dans laquelle la seule fonction inconnue est $v(t)$:

$$m \frac{dv}{dt} = \frac{BaE}{R} - \frac{(Ba)^2}{R} v(t).$$

La force de Laplace s'écrit *in fine*, en partie sous la forme d'une force d'entraînement constante BaE/R , et en partie sous la forme d'une **force de frottement fluide** en $-\text{constante} \times v$, caractéristique des phénomènes d'électromécanique.

De plus, le système est décrit par une équation différentielle du premier ordre qu'on écrit sous forme canonique :

$$\tau \frac{dv}{dt} + v(t) = \frac{E}{Ba},$$

où $\tau = \frac{mR}{(Ba)^2}$ est la constante de temps du système.

Si par exemple, la tige est initialement au repos, $v(0) = 0$, la solution du couple {équation différentielle, condition initiale} est :

$$v(t) = \frac{E}{Ba} \left(1 - \exp\left(-\frac{t}{\tau}\right) \right).$$

Quant à l'intensité du courant $i(t)$, l'équation électrique (EE) précise sa valeur :

$$i(t) = \frac{E - Bav(t)}{R} = \frac{E}{R} \exp\left(-\frac{t}{\tau}\right).$$

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

f) Bilan de puissance

L'équation électrique est multipliée par i et l'équation mécanique par v :

$$\begin{cases} (EE) \times i \\ (EM) \times v \end{cases} \Rightarrow \begin{cases} ei + Ei = Ri^2 \\ m \frac{dv}{dt} v = f_L v \end{cases} \text{ soit } \begin{cases} Ei - Bavi = Ri^2 \\ \frac{d}{dt} \left(\frac{1}{2} mv^2 \right) = iaBv \end{cases}$$

Le même terme de couplage $Bavi$ apparaît dans les deux équations. En l'éliminant on arrive à :

$$Ei = \frac{d}{dt} \left(\frac{1}{2} mv^2 \right) + Ri^2.$$

Cette dernière équation explicite les transferts de puissance. Le générateur fournit une puissance électrique Ei , qui se répartie en :

- puissance cinétique $\frac{d}{dt} \left(\frac{1}{2} mv^2 \right)$, dérivée de l'énergie cinétique de la tige $\mathcal{T}$,
- puissance de l'effet Joule Ri^2 qui représente ici une perte puisqu'on souhaite obtenir de la puissance mécanique.

La mise en mouvement de la tige $\mathcal{T}$ sous l'effet d'un générateur électrique justifie l'adjectif de moteur donné aux rails de Laplace : une puissance électrique se transforme en une puissance mécanique.

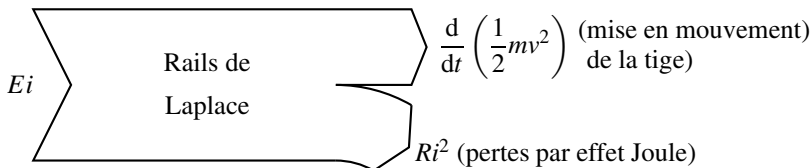


Figure 30.9 – Transferts de puissance pour des rails de Laplace moteurs.

g) Bilan d'énergie

Le bilan d'énergie est l'intégrale du bilan de puissance entre $t = 0$ et $t \rightarrow \infty$.

L'énergie électrique délivrée par le générateur entre l'instant initial $t = 0$ et l'instant final $t \rightarrow \infty$ est :

$$\begin{aligned} \mathcal{E}_{\text{gén}} &= \int_0^\infty Eidt = \frac{E^2}{R} \int_0^\infty \exp\left(-\frac{t}{\tau}\right) dt = \frac{E^2}{R} \left[-\tau \exp\left(-\frac{t}{\tau}\right) \right]_0^\infty \\ &= \frac{E^2}{R} \tau = \frac{mE^2}{(Ba)^2}. \end{aligned}$$

L'énergie cinétique reçue par la tige $\mathcal{T}$ est :

$$E_c = \int_0^\infty \frac{d}{dt} \left(\frac{1}{2} mv^2 \right) dt = \left[\frac{1}{2} mv^2 \right]_0^\infty = \frac{1}{2} m \left(v(\infty)^2 - v(0)^2 \right) = \frac{mE^2}{2(Ba)^2}.$$

L'énergie électrique dissipée en effet Joule est :

$$\begin{aligned}\mathcal{E}_J &= \int_0^\infty Ri^2 dt = \frac{E^2}{R} \int_0^\infty \exp\left(-2\frac{t}{\tau}\right) dt = \frac{E^2}{R} \left[-\frac{\tau}{2} \exp\left(-\frac{t}{\tau}\right)\right]_0^\infty \\ &= \frac{E^2}{R} \frac{\tau}{2} = \frac{mE^2}{2(Ba)^2}.\end{aligned}$$

L'énergie du générateur est répartie à parts égales entre la tige $\mathcal{T}$ et les pertes par effet Joule : $\mathcal{E}_{\text{gén}} = E_c + \mathcal{E}_J$. Ce moteur présente donc un rendement de 0,5 (ce qui est très mauvais pour un moteur électrique).

2.2 Haut-parleur électrodynamique

a) Présentation

Un **haut-parleur** est un appareil électromécanique qui transforme un signal électrique en signal sonore.

On va en étudier le principe physique de fonctionnement dans la géométrie simplifiée des rails de Laplace. Comme le montre le schéma suivant, une membrane est solidaire de la tige $\mathcal{T}$. Elle sert à émettre une onde sonore.

L'équipage mobile, constitué de $\mathcal{T}$ et de la membrane, est de masse totale m_0 . La tige $\mathcal{T}$ est reliée au bâti par l'intermédiaire d'un ressort de rappel, de constante de raideur k . Le générateur de tension $E(t)$ délivre le signal électrique à transformer en signal sonore.

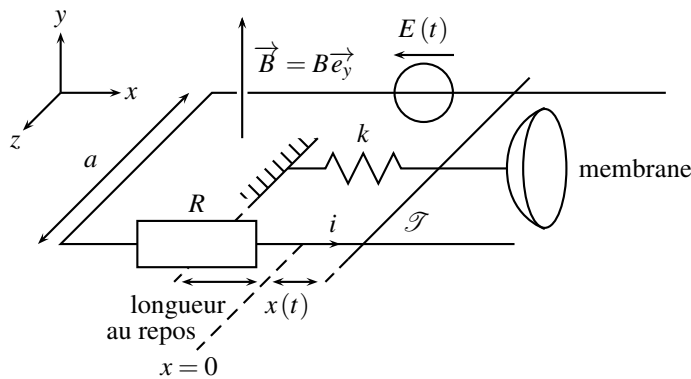


Figure 30.10 – Schéma de principe d'un haut-parleur.

b) Analyse physique

- Le générateur impose la tension $E(t)$ et fait circuler un courant d'intensité i ,
- i et $\vec{B}$ créent une force de Laplace $\vec{f}_L$,
- la tige $\mathcal{T}$ et la membrane vibrent sous l'effet de la force $\vec{f}_L$,
- la vibration de la membrane émet une onde sonore, image du signal électrique $E(t)$,
- $\mathcal{T}$ devient un circuit mobile dans $\vec{B}$ stationnaire : il apparaît dans ce conducteur une f.é.m. induite e , qui s'oppose à la tension $E(t)$ du générateur (loi de Lenz).

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

Une partie de la puissance électrique délivrée par le générateur est transformée en puissance mécanique.

Remarque

Le même système peut aussi transformer une puissance mécanique en une puissance électrique. C'est alors un microphone.

c) F.é.m. induite et équation électrique

Le sens positif du courant est indiqué sur la figure ci-dessus.

Le calcul de la f.é.m., la justification de la validité de la loi de Faraday, le schéma électrique équivalent et l'équation électrique sont identiques à ceux du paragraphe 2.1. La f.é.m. induite est $e(t) = -Bav(t)$ et l'équation électrique est :

$$e(t) + E(t) - Ri(t) = 0 \quad \text{soit} \quad -Bv(t)a + E(t) = Ri(t) \quad (EE).$$

d) Force de Laplace et équation mécanique

La force de Laplace $\vec{f}_L$ qui s'exerce sur la tige $\mathcal{T}$ est la même que celle du paragraphe 2.1 : $\vec{f}_L = i(t)aB\vec{u}_x$.

Il faut de plus tenir compte de la force $\vec{f}_R$ de rappel du ressort, ainsi que de la force $\vec{f}_S$ modélisant l'effet de l'onde sonore sur la membrane.

En repérant la position de la tige $\mathcal{T}$ par rapport à sa position au repos :

$$\vec{f}_R = -kx\vec{u}_x.$$

De plus, pour tenir compte de la perte d'énergie de la membrane liée à l'émission de l'onde sonore on ajoute une force de frottement fluide :

$$\vec{f}_S = -\alpha \vec{v}.$$

Le principe fondamental de la dynamique, appliqué à l'équipage mobile constitué de la tige $\mathcal{T}$ et de la membrane, en projection sur $\vec{u}_x$, mène à l'équation mécanique :

$$m_0 \frac{dv}{dt} = i(t)aB - kx(t) - \alpha v(t) \quad (EM).$$

e) Bilan de puissance

L'équation électrique (EE) est multipliée par i et l'équation mécanique (EM) par v :

$$\begin{cases} (EE) \times i \\ (EM) \times v \end{cases} \Rightarrow \begin{cases} -Bvai + Ei = Ri^2 \\ m \frac{dv}{dt} v = iaBv - kxv - \alpha v^2 \end{cases}$$

Avec $v(t) = \frac{dx}{dt}$ et donc $kxv = kx \frac{dx}{dt} = \frac{d}{dt} \left(\frac{1}{2} kx^2 \right)$, en éliminant le terme $Bavi$ on trouve :

$$Ei = \frac{d}{dt} \left(\frac{1}{2} mv^2 + \frac{1}{2} kx^2 \right) + \alpha v^2 + Ri^2.$$

Le générateur délivre une puissance qui sert à mettre en mouvement la membrane, à produire une onde sonore (αv^2) et à compenser les pertes par effet Joule (Ri^2).

f) Régime sinusoïdal permanent, ou établi

Les grandeurs électriques et mécaniques sont couplées *via* les équations électrique et mécanique. Les grandeurs mécaniques ont ainsi une influence électrique. Ce couplage apparaît clairement quand on établit l'équivalent électrique complet du haut-parleur.

Les équations électrique et mécaniques deviennent, en régime sinusoïdal permanent à la pulsation ω de l'onde sonore, émise ou reçue :

$$\begin{cases} \underline{E} - Ba\underline{v} = R\underline{i} \\ m_0j\omega\underline{v} = Ba\underline{i} - k\frac{\underline{v}}{j\omega} - \alpha\underline{v} \end{cases}$$

On élimine $\underline{v}$:

$$\underline{E} - Ba \frac{Ba\underline{i}}{m_0j\omega + \frac{k}{j\omega} + \alpha} = R\underline{i} \quad \text{soit} \quad \underline{E} = R\underline{i} + \boxed{\frac{(Ba)^2}{m_0j\omega + \frac{k}{j\omega} + \alpha}} \underline{i}.$$

$= Z_{cin}$

Apparaît une **impédance cinétique** Z_{cin} , qui marque le couplage entre les grandeurs mécaniques et les grandeurs électriques. Elle est constituée de trois dipôles équivalents en parallèle. En effet, l'inverse de l'impédance est :

$$\frac{1}{Y_{cin}} = \boxed{\frac{m_0}{(Ba)^2}} j\omega + \boxed{\frac{k}{(Ba)^2}} \frac{1}{j\omega} + \boxed{\frac{\alpha}{(Ba)^2}}.$$

$= C_{cin} \qquad \qquad \qquad = 1/L_{cin} \qquad \qquad \qquad = 1/R_{cin}$

Le haut-parleur est donc électriquement équivalent à :

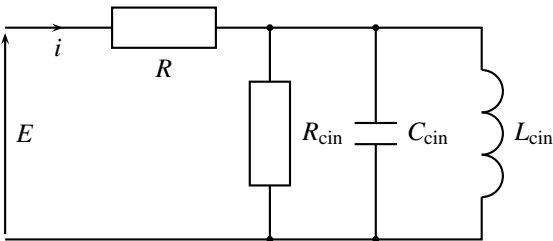


Figure 30.11 – Modèle électrique du haut-parleur.

Le comportement du haut-parleur dépend donc de la pulsation ω . Il n'a donc pas les mêmes performances pour toutes les fréquences. Dans la pratique, pour assurer une bonne restitution du son, il faut associer plusieurs haut-parleurs.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

2.3 Machine à courant continu à entrefer plan (PTSI)

a) Présentation

On considère une roue d'axe (Oz) constituée d'un cercle externe et de rayons, tous conducteurs. C'est l'analogue électrique d'une roue de voiture à cheval en bois. Un dispositif, non représenté, permet à un courant électrique de circuler radialement entre le centre et la périphérie, le long des rayons. L'ensemble est plongé dans un champ magnétique uniforme et stationnaire $\vec{B} = -B\vec{u}_z$.

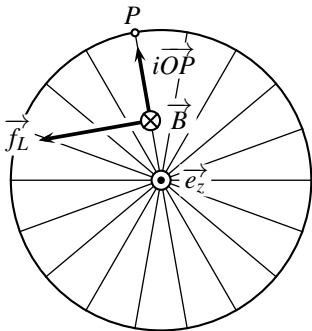


Figure 30.12 – Roue conductrice étudiée.

b) Analyse physique

Ce dispositif est réversible et fonctionne en générateur ou en moteur. Dans le cas **générateur** :

- la roue est entraînée en rotation,
- elle devient un circuit mobile dans $\vec{B}$ stationnaire : une f.é.m. e est induite dans chaque rayon,
- l'ensemble des rayons forme un circuit électrique fermé et conducteur : un courant induit i circule sur chaque rayon,
- i et $\vec{B}$ créent un moment des forces de Laplace $\vec{M}_L$ qui s'oppose à la rotation de la roue (loi de Lenz).

Dans le cas **moteur** :

- un générateur extérieur impose la circulation d'un courant d'intensité i le long des rayons,
- i et $\vec{B}$ créent un moment des forces de Laplace $\vec{M}_L$ qui met en rotation de la roue,
- la roue tourne, elle devient un circuit mobile dans $\vec{B}$ stationnaire : une f.é.m. e est induite dans chaque rayon, qui s'oppose à la tension du générateur extérieur (loi de Lenz).

c) Complément : mise en équation

Choix des orientations On se place dans une base cylindrique $(O; \vec{u}_r, \vec{u}_\theta, \vec{u}_z)$. Le sens positif de l'intensité i du courant est arbitrairement choisi du centre vers la périphérie, comme sur le schéma.

Les mêmes étapes, f.é.m. induite et moment des forces de Laplace, apparaissent dans les deux fonctionnements de l'analyse physique. La mise en équation est donc unique, quelque soit le fonctionnement envisagé.

Moment des forces de Laplace On considère un unique rayon. La force de Laplace qui s'exerce sur un rayon OP est :

$$\vec{f}_L = i\vec{OP} \wedge \vec{B} = iaB\vec{u}_\theta.$$

On admet que cette force s'exerce au milieu H du segment $[OP]$. Son moment par rapport au point O est alors :

$$\vec{M}_L = \vec{OH} \wedge \vec{f}_L = \frac{a}{2}\vec{u}_r \wedge iaB\vec{u}_\theta = \frac{iBa^2}{2}\vec{u}_z = \varphi_0 i\vec{u}_z,$$

avec $\varphi_0 = \frac{Ba^2}{2}$.

Dans une MCC, le couple de Laplace est proportionnel à l'intensité du courant.

Ce moment entraîne la roue en rotation autour de $\vec{u}_z$ (fonctionnement moteur), ou s'oppose à sa rotation (fonctionnement générateur).

F.é.m. induite La loi de Faraday est ici inopérante. En effet, elle n'est valable, dans le cas d'un circuit mobile dans un champ magnétique stationnaire, que si le circuit coupe des lignes de champ et qu'on peut définir un flux magnétique $\varphi(t)$ variable, ce qui n'est pas le cas ici, car le flux à travers la roue reste constant.

Le calcul direct sort du cadre de ce cours. On a alors recourt à la nullité de la somme des puissances de la f.é.m. et du moment de Laplace. Pour un unique rayon :

$$\mathcal{P}_{fém} + \mathcal{P}_L = ei + \frac{iBa^2}{2}\vec{u}_z \cdot \omega\vec{u}_z = 0, \quad \text{soit} \quad e = -\frac{\omega Ba^2}{2} = -\varphi_0\omega.$$

Dans une MCC, la f.é.m. est proportionnelle à la vitesse angulaire.

Cette f.é.m. est recueillie par l'utilisateur (fonctionnement générateur), ou s'oppose au générateur extérieur (fonctionnement moteur).

d) Réalisation

La machine à courant continu, nommée suivant son acronyme MCC, ici à entrefer plan, est composée d'une partie fixe, nommé **stator**, et d'une partie mobile, nommée **rotor**.

Le stator est ici constitué d'une carcasse et de deux supports fixes circulaires, sur lesquels sont montés des aimants permanents (2 jeux de 8 aimants visibles sur la figure 30.13) qui créent un champ magnétique stationnaire.

Le rotor est constitué d'un axe de rotation et d'un disque, solidaire de l'axe, sur lequel sont imprimés de nombreux circuits électriques radiaux, séparés par un isolant, non représentés afin de ne pas surcharger le schéma. Ces circuits électriques imposent une conduction radiale du courant.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

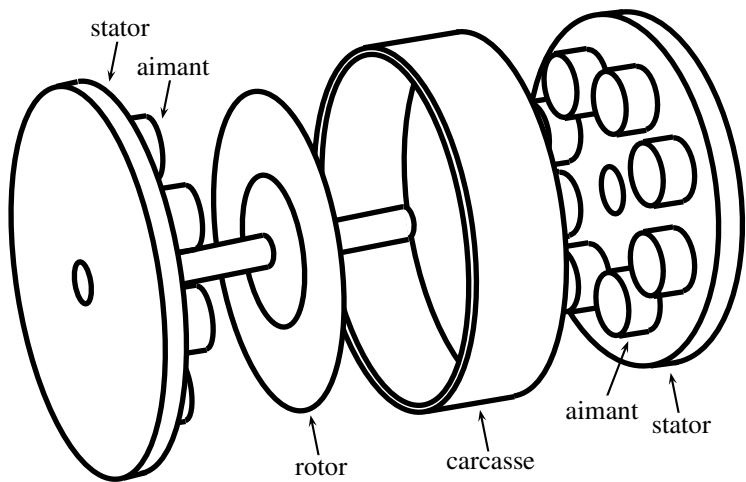


Figure 30.13 – Vue éclatée d’une MCC à entrefer plan.

Des balais, contacts métalliques frottants, assurent le passage du courant entre le circuit extérieur fixe et les circuits électriques mobiles rotoriques.

La MCC est totalement réversible, elle peut servir soit en générateur, soit en moteur. Son nom insiste sur le caractère continu des courants d’alimentation, par opposition à sinusoïdal. En effet, une vitesse de rotation ω constante est obtenue en alimentant la MCC avec une tension constante e , un couple constant C par un courant i constant.

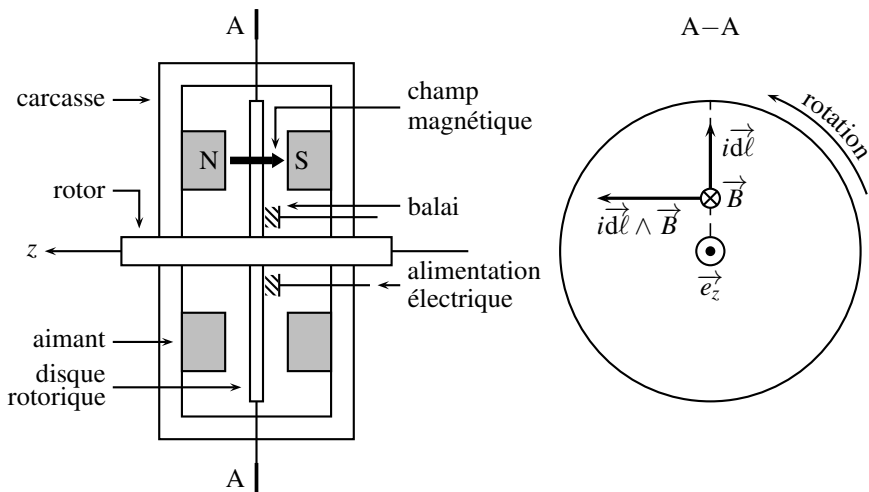


Figure 30.14 – Schéma d’une MCC à entrefer plan.

L’**entrefer** est la zone située entre deux aimants. Il s’agit ici du plan du disque rotorique, d’où le substantif d’entrefer plan.

Une MCC à entrefer plan est nommé en anglais *brushed disc DC motor* ou *brushed disc DC servo*. *Brushed* signifie que le rotor est alimenté par des balais, *disc* insiste sur la forme géométrique du rotor, *DC*, pour *direct current*, s'oppose à *AC*, pour *alternative current* et *servo* indique l'utilisation possible en servomécanisme (asservissement dans lequel la grandeur de sortie est mécanique).

Un bobinage très ingénieux des circuits électriques radiaux du rotor, associé à une alternance des pôles des aimants, permet une alimentation uniquement à partir des balais, en particulier pour la gestion du retour du courant arrivé en bout de rotor. Toutefois, cette réalisation technologique sort du cadre de ce cours. Le lecteur intéressée pourra, par exemple, se reporter à François BERNOT, *Machines à courant continu, construction*, Techniques de l'Ingénieur, traité Génie électrique, D 3 556, 2.3 Induits discoïdaux.

e) Utilisations

Les avantages des MCC à entrefer plan sont une vitesse très contrôlée et stable, de 1 à $4.10^3 \text{ tour.min}^{-1}$, une très grande accélération angulaire, jusqu'à environ $150.10^3 \text{ rad.s}^{-2}$, un couple indépendant de la vitesse de rotation. De plus, la très faible inductance du rotor, ainsi que son faible moment d'inertie, produisent des constantes de temps mécanique τ_m et électrique τ_e extrêmement faibles, jusqu'à $\tau_m = 4 \text{ ms}$ et $\tau_e < 5.10^{-2} \text{ ms}$, ainsi qu'un très faible encombrement. Toutefois, la puissance disponible reste limitée à environ 1 kW.

Dans quels domaines utilise-t-on alors une MCC à entrefer plan ? Partout où il faut créer un mouvement de rotation soit de précision, soit dans un volume limité.

La géométrie très plate sera appréciée en motorisation des bicyclettes, des chaises roulantes ; la précision de la rotation en robotique industrielle, médicale (pompes à sang, à dialyse, respirateurs), informatique (rotation de disques de stockage), militaire (chargeurs automatiques de munitions, tourelles).

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

SYNTHÈSE

SAVOIRS

- origine des courants de Foucault
- utilisation des courants de Foucault
- principe du freinage par induction
- principe de fonctionnement d'un haut-parleur
- description qualitative du principe de la MCC à entrefer plan (PTSI)
- exemples d'utilisation d'une MCC (PTSI)

SAVOIR-FAIRE

- mener une analyse physique
- préciser les orientations
- écrire les équations électrique et mécanique
- effectuer un bilan de puissance

MOTS-CLÉS

- électromécanique (PTSI)
- courants de Foucault
- MCC à entrefer plan
- haut-parleur
- alternateur
- rails de laplace

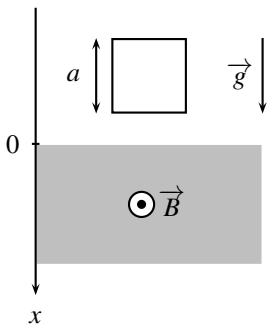
S'ENTRAÎNER

30.1 Cadre qui chute dans un champ localisé (★)

Un cadre conducteur, constitué de 4 segments de longueur a , tombe dans le plan du schéma sous l'effet de la gravité. Sa résistance électrique est notée R , son autoinductance L .

L'espace est divisé en deux régions :

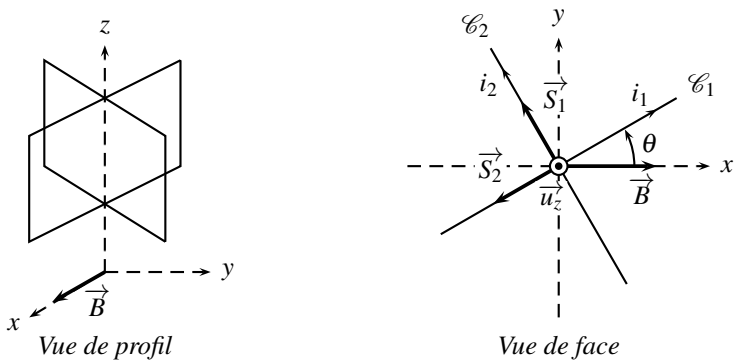
- pour $x < 0$, il n'y a pas de champ magnétique,
- pour $x > 0$, un champ magnétique est présent. Il est uniforme, stationnaire et orthogonal au plan du schéma.



Établir les équations différentielles régissant la vitesse $v(t)$ du cadre dans les 3 régions :

1. le cadre est entièrement dans la région où $\vec{B} = \vec{0}$,
2. le cadre est à cheval sur les régions où $\vec{B} = \vec{0}$ et $\vec{B} \neq \vec{0}$,
3. le cadre est entièrement dans la région où $\vec{B} \neq \vec{0}$.

30.2 Deux cadres à angle droit (★)



Deux cadres rectangulaires, identiques et solidaires, dont les plans forment un angle droit, chacun de résistance électrique R , d'auto-inductance négligeable, de surface S , sont mobiles

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

en rotation autour de l'axe vertical (Oz), sans aucun frottement. Leur moment d'inertie par rapport à cet axe est J . Leur vitesse angulaire initiale est ω_0 .

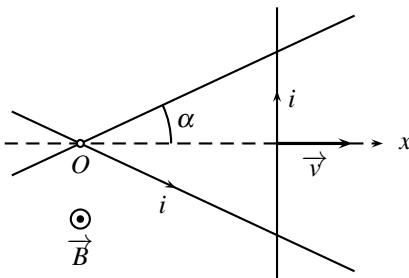
Ils sont placés dans un champ magnétique uniforme et constant $\vec{B} = B\vec{u}_x$ horizontal.

Il n'y a aucun contact électrique entre les cadres, leur courants ne se mélangent pas.

1. Calculer la vitesse angulaire $\omega(t)$ des deux cadres.
2. Former un bilan énergétique du fonctionnement.

30.3 Rails de Laplace croisés (★)

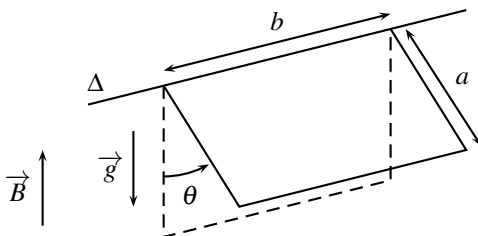
On considère deux rails de Laplace qui se croisent au point O . Ils sont horizontaux et plongés dans un champ magnétique $\vec{B}$ uniforme. Une tige $\mathcal{T}$ glisse sans frottement sur eux, tout en restant parallèle à elle-même ; elle est entraînée à la vitesse constante $\vec{v} = v\vec{u}_x$.



1. Établir, en fonction de la position $x(t)$ de la tige $\mathcal{T}$, l'expression de la f.é.m. induite dans le circuit formé des rails et de la tige.
2. Attendu que la résistance électrique du circuit est proportionnelle à sa longueur, c'est à dire que $R = k\ell$, déterminer l'intensité du courant dans le circuit. Commenter son signe.
3. Établir l'expression de la force qu'un opérateur doit exercer sur la tige $\mathcal{T}$ afin qu'elle garde sa vitesse constante.
4. Établir les expressions de la puissance perdue par effet Joule, de la puissance fournie par l'opérateur. Les comparer.

30.4 Cadre oscillant (★)

Un cadre conducteur tourne sans frottement autour de l'axe Δ . Il est composé de 4 segments, 2 de longueur a , 2 de longueur b . La masse totale du cadre est m , son moment d'inertie par rapport à Δ est J , sa résistance électrique est R et son autoinductance est négligée.



Les champs magnétique et de pesanteur sont uniformes et constants.
On écarte le cadre de sa position d'équilibre verticale (repérée en pointillés sur le schéma) d'un angle θ_0 puis on le lâche sans vitesse initiale.

- 1. Dans quel sens l'axe Δ est-il orienté ? Justifier.
- 2. Établir l'équation différentielle à laquelle obéit $\theta(t)$. La linéariser.
- 3. Tracer l'allure de $\theta(t)$. Discuter en fonction de la valeur du coefficient d'amortissement ξ .

APPROFONDIR

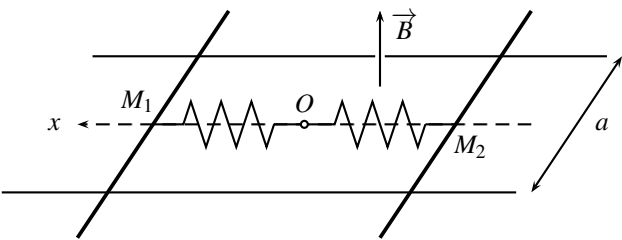
30.5 Deux tiges reliées par des ressorts (★)

Deux barreaux identiques de masse m , parallèles, sont reliés par leurs milieux à deux ressorts identiques de raideur k . Les deux ressorts sont attachés au point O . Les barreaux glissent sans frottement sur deux rails parallèles, horizontaux, distants d'une longueur a . Le tout est plongé dans un champ magnétostatique uniforme orthogonal au plan formé par les rails et les barreaux.

La résistance électrique du circuit filiforme ainsi créé est R , constante quelles que soient les positions des barreaux. De plus, l'autoinductance du circuit sera négligée.

Soient M_1 et M_2 les points d'attache entre les barreaux et les ressorts. On repère la position des barreaux par $x_1(t)$ et $x_2(t)$, distance entre M_1 , ou M_2 , et leur position à vide, où les ressorts ont leur longueur propre.

À l'instant initial, les vitesses des barreaux sont nulles, M_1 est écarté d'une distance b par rapport à sa position au repos et M_2 est au repos.



- 1. Expliquer physiquement ce qu'on observe.
- 2. Quelles sont les deux équations couplées qui régissent le mouvement des barreaux ?
- 3. Calculer $x_1(t)$ et $x_2(t)$ en régime permanent *via* le changement de fonction :

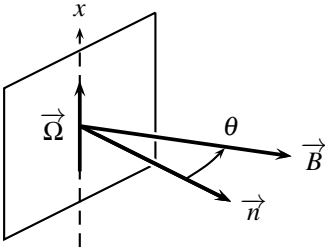
$$\sigma(t) = x_1(t) + x_2(t) \quad \text{et} \quad \delta(t) = x_1(t) - x_2(t).$$

- 4. Dresser et interpréter un bilan énergétique.

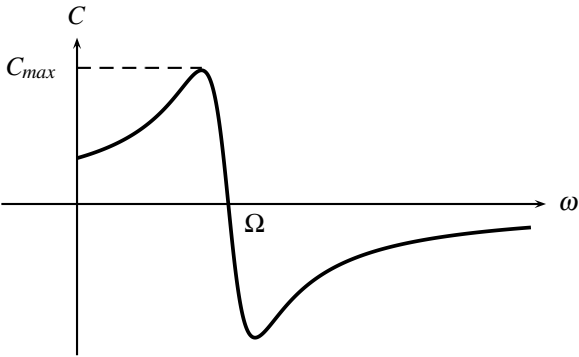
CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

30.6 Moteur asynchrone (MAS) (★)

Une spire plate de résistance R , d'inductance L et de surface S , tourne à vitesse angulaire constante Ω autour d'un de l'axe (Ox) . La normale $\vec{n}$ à la spire est contenue dans le plan (Oyz) . La spire est plongée dans un champ magnétique $\vec{B}$ localement uniforme, contenu dans le plan Oyz , de norme constante, tournant à la vitesse angulaire constante ω autour de (Ox) . Ce dispositif est utilisé en moteur électrique : le champ magnétique entraîne la bobine.



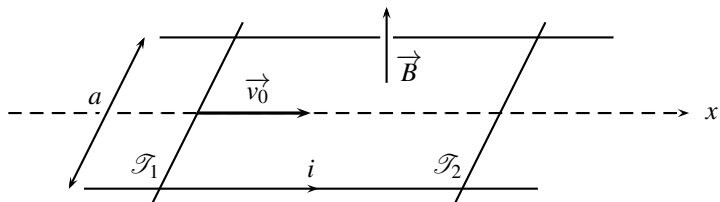
1. Comment réaliser un champ magnétique tournant ?
2. Expliquer sans équation pourquoi la spire tourne. Les deux vitesses ω et Ω peuvent-elles être identiques ?
3. Calculer l'équation différentielle régissant l'évolution du courant dans la bobine en fonction de l'angle instantané θ entre le champ magnétique $\vec{B}$ et la normale $\vec{n}$ à la spire. L'intégrer en régime harmonique permanent grâce au passage aux complexes (on justifiera avec soin la pulsation du courant et on n'oubliera pas de prendre la notation complexe pour les deux membres de l'équation).
4. En considérant le moment magnétique $\vec{\mathcal{M}}$ de la spire, calculer le couple auquel elle est soumise. En déduire le couple moyen au cours du temps C s'exerçant sur la bobine.
5. L'allure de la courbe $C(\omega)$ est donnée sur la figure. Le moteur peut-il démarrer seul ? Étudier graphiquement la stabilité des points de fonctionnement si le moteur entraîne une charge de couple résistant constant connu.



30.7 Interaction entre deux tiges (★)

Deux tiges $\mathcal{T}_1$ et $\mathcal{T}_2$ identiques, de masse m , chacune de résistance électrique $R/2$, sont mobiles sans frottement sur deux rails parallèles horizontaux espacés d'une distance a . La résistance électrique des rails est négligeable devant R ; l'inductance propre du circuit est négligée. L'ensemble est plongé dans un champ magnétique uniforme et constant $\vec{B} = B\vec{e}_z$, et un champ de pesanteur $\vec{g} = -g\vec{u}_z$.

Initialement, $\mathcal{T}_2$ est immobile et $\mathcal{T}_1$ se déplace vers $\mathcal{T}_2$ avec la vitesse $\vec{v}_0 = v_0\vec{u}_x$. Les deux tiges restent parallèles à $\vec{u}_y$ lors de leur mouvement.



1. Expliquer sans calcul pourquoi la tige $\mathcal{T}_1$ ralentit alors que la tige $\mathcal{T}_2$ se met en mouvement.
2. Établir une équation électrique reliant $i(t)$, intensité du courant dans le système, à $v_1(t)$ et $v_2(t)$. $\vec{v}_1 = v_1(t)\vec{u}_x$ et $\vec{v}_2 = v_2(t)\vec{u}_x$ sont les vitesses des tiges $\mathcal{T}_1$ et $\mathcal{T}_2$.
3. Établir deux équations mécaniques.
4. Établir un système d'équations différentielles couplées sur $v_1(t)$ et $v_2(t)$.
5. Découpler et intégrer le système différentiel en utilisant les fonctions somme $\sigma(t) = v_1(t) + v_2(t)$ et différence $\delta(t) = v_1(t) - v_2(t)$. Représenter graphiquement sur un même schéma $v_1(t)$ et $v_2(t)$.
6. Calculer l'intensité du courant $i(t)$ qui circule dans les deux tiges.
7. Calculer la charge totale Q qui circule entre $t = 0$ et $t \rightarrow \infty$.
8. Calculer, en utilisant $i(t)$, l'énergie $\mathcal{E}_J$ dissipée par effet Joule entre $t = 0$ et $t \rightarrow \infty$.
9. Calculer la variation $\Delta\mathcal{E}_m$ d'énergie mécanique du système entre $t = 0$ et $t \rightarrow \infty$. Commenter.

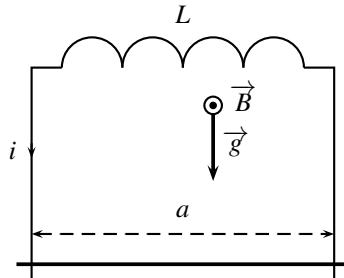
30.8 Tige qui chute (★)

Une tige rectiligne de longueur a , de masse m et de résistance électrique R effectue un mouvement de translation verticale tout en fermant le circuit électrique qui comporte la bobine d'inductance L . On confond la résistance totale du circuit avec R et son inductance totale avec L .

L'ensemble est plongé dans un champ magnétique $\vec{B} = B\vec{u}_z$ et un champ de pesanteur $\vec{g} = g\vec{u}_y$, uniformes et constants.

Le mouvement de la tige est sans frottement. Elle est abandonnée à $t = 0$ sans vitesse.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE



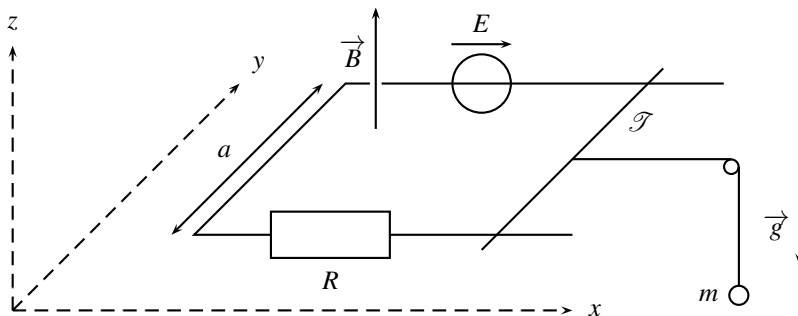
1. Établir l'équation différentielle relative uniquement à l'intensité du courant traversant la tige.
2. Dans l'hypothèse d'une résistance R nulle, calculer explicitement l'intensité du courant puis la vitesse en fonction du temps.
3. Dans le cas d'une résistance « assez grande » (on précisera devant quelle grandeur R doit être grande), décrire qualitativement l'évolution du courant en traçant l'allure de l'intensité du courant $i(t)$. Quelles sont les valeurs de $i(t)$ et $v(t)$ en régime permanent ?

30.9 Tige entraînée par gravité (★)

Une tige conductrice $\mathcal{T}$, de masse m_0 , glisse sans frotter suivant $\vec{u}_x$ sur deux rails distants de a . Elle ferme électriquement un circuit comprenant une résistance R et un générateur de f.é.m. constante E . $\mathcal{T}$ a une résistance électrique négligeable devant R et une masse m_0 . On négligera tout phénomène d'autoinduction.

$\mathcal{T}$ est lié par son centre à un fil sans masse et inextensible portant une masse ponctuelle m dont le mouvement est vertical. Le fil coulisse sans frottement autour d'une poulie immobile. L'ensemble est plongé dans des champ magnétique et de pesanteur uniformes, stationnaires et orthogonaux au plan du circuit électrique.

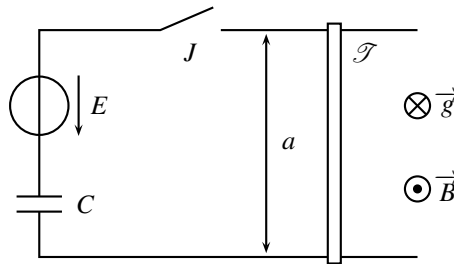
On lâche la tige $\mathcal{T}$ alors qu'elle est initialement immobile. On note $\vec{v}(t) = v(t) \vec{u}_x$ la vitesse de $\mathcal{T}$ et $\vec{v}_m(t) = v_m(t) \vec{u}_z$ celle de la masse m .



1. Quel est le lien entre v et v_m ?
2. Établir l'expression de la f.é.m. induite dans le circuit.

3. Établir l'expression de la force de Laplace qui s'exerce sur $\mathcal{T}$, ainsi que celle de la force du fil sur la tige $\mathcal{T}$.
4. Établir une équation différentielle portant sur $v(t)$.
5. Calculer et tracer $v(t)$. Quels sont les différents comportements possibles suivant les valeurs de E ?

30.10 Tige qui glisse sur un circuit capacitif (★★)



Une tige conductrice $\mathcal{T}$ glisse sur deux rails horizontaux distants de a . Elle ferme électriquement un circuit comprenant un interrupteur J , un condensateur de capacité C et un générateur de f.é.m. constante E . $\mathcal{T}$ a une résistance électrique R et une masse m . L'autoinductance du circuit sera négligée.

L'ensemble est plongé dans des champ magnétique et de pesanteur uniformes et stationnaires. On ferme à l'instant initial l'interrupteur J alors que la tige $\mathcal{T}$ est immobile.

1. Établir une équation électrique et une équation mécanique décrivant le circuit.
2. Établir et intégrer une équation différentielle sur l'intensité du courant dans le circuit pour montrer qu'il s'écrit sous la forme :

$$i(t) = i_0 \exp\left(-\frac{t}{\tau}\right).$$

Identifier les valeurs de i_0 et τ .

3. En déduire que la vitesse $v(t)$ de la tige se met sous la forme :

$$v(t) = v_0 \left(1 - \exp\left(-\frac{t}{\tau}\right)\right).$$

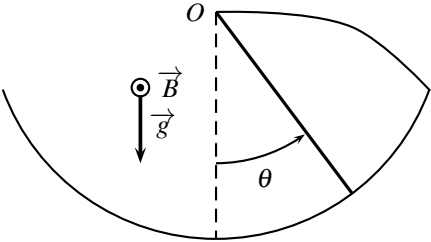
4. Calculer l'énergie $\mathcal{E}_G$ fournie par le générateur entre les instants initial ($t = 0$) et final ($t \rightarrow \infty$), en fonction de E , τ et R .
5. Calculer $u_C(t)$.
6. Calculer l'énergie $\mathcal{E}_C$ emmagasinée par le condensateur entre les instants initial et final.
7. Calculer l'énergie $\mathcal{E}_J$ dissipée par effet Joule entre les instants initial et final.
8. Calculer le travail W des forces de Laplace entre les instants initial et final.
9. Quelle relation existe-t-il entre $\mathcal{E}_G$, $\mathcal{E}_C$, $\mathcal{E}_J$ et W ? L'interpréter.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

30.11 Tige en rotation sur un cercle (PTSI) (★★)

Une barre conductrice, de longueur a , est mobile sur un fil conducteur circulaire. Le circuit est fermé par un fil électrique. La résistance électrique de la barre est R , les résistances des autres éléments du circuit sont négligeables devant R . On ne tiendra pas compte des phénomènes d'autoinduction. L'ensemble est plongé dans un champ magnétique uniforme et constant $\vec{B} = B\vec{u}_z$.

La barre, de longueur a , tourne autour de l'axe (Oz) sans frottement. Son moment d'inertie par rapport à Oz vaut J . La barre est lâchée en $t = 0$ de la position $\theta_0 > 0$ avec une vitesse nulle.



1. Établir puis linéariser l'équation différentielle sur $\theta(t)$. Préciser la pulsation caractéristique ω_0 et le facteur d'amortissement ξ .
2. Mener un bilan énergétique. Le commenter en faisant clairement apparaître les énergies cinétiques et potentielle de la tige.

CORRIGÉS

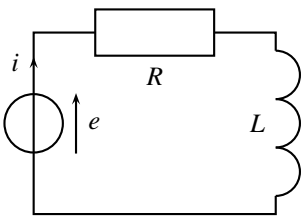
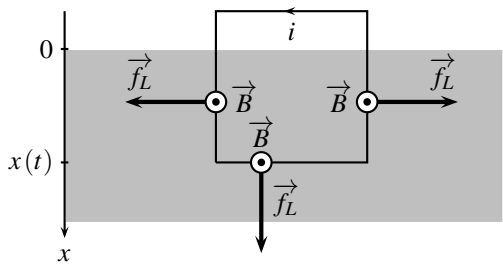
30.1 Cadre qui chute dans un champ localisé

On oriente le cadre dans le sens trigonométrique positif. Ainsi le vecteur surface $\vec{S}$ est colinéaire et de même direction que le champ magnétique $\vec{B}$.

1. $\vec{B} = \vec{0}$: aucun phénomène d'induction. Le cadre est en chute libre. Le principe fondamental de la dynamique appliqué au cadre mène à : $m \frac{d\vec{v}}{dt} = m \vec{g}$ soit $\frac{d\vec{v}}{dt} = \vec{g}$.

2. Il n'y a pas de force de Laplace sur le coté supérieur ; les forces de Laplace sur les cotés verticaux s'annulent deux à deux ; seule reste celle sur le coté inférieur : $\vec{f}_L = iaB\vec{u}_x$.

Le principe fondamental de la dynamique, appliqué au cadre, en projection sur $\vec{u}_x$, mène à : $m \frac{dv}{dt} = iaB + mg$.



La surface du cadre offerte au champ magnétique est $S = ax$. Le flux du champ magnétique à travers le cadre et la f.e.m. induite valent : $\varphi = \vec{B} \cdot \vec{S} = Bax$ donc $e = -\frac{d\varphi}{dt} = -Bav$.

Le schéma électrique du circuit est dessiné ci-dessus : $e = Ri + L \frac{di}{dt} = -Bva$.

On obtient une équation différentielle sur la vitesse seule en prélevant i de l'équation mécanique, $i = \frac{m}{Ba} \left(\frac{dv}{dt} - g \right)$, qu'on injecte dans l'équation électrique : $-Bva = \frac{Rm}{Ba} \left(\frac{dv}{dt} - g \right) + \frac{Lm}{Ba} \frac{d}{dt} \left(\frac{dv}{dt} - g \right)$.

Sous forme canonique $\left(\frac{1}{\omega_0^2} \frac{d^2v}{dt^2} + \frac{2\xi}{\omega_0} \frac{dv}{dt} + v = v_0 \right)$: $\frac{Lm}{(Ba)^2} \frac{d^2v}{dt^2} + \frac{Rm}{(Ba)^2} \frac{dv}{dt} + v(t) = \frac{Rmg}{(Ba)^2}$.

3. Le cadre est complètement immergé dans le champ magnétique. Le flux magnétique à travers le cadre est constant, il n'y a pas de f.é.m. induite, aucun phénomène d'induction. On retrouve donc : $\frac{d\vec{v}}{dt} = \vec{g}$.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

30.2 Deux cadres à angle droit

1. On oriente arbitrairement les cadres comme sur le schéma, afin de faire apparaître les vecteurs surfaces $\vec{S}_1$ et $\vec{S}_2$ d'après la règle de la main droite. Le choix de ces orientations peut différer ; les signes des f.é.m. et des couples changent alors, mais le résultat final reste le même.

Les flux magnétiques et les f.é.m. induites dans les cadres $\mathcal{C}_1$ et $\mathcal{C}_2$ sont :

$$\begin{cases} \phi_1 = BS \cos \left(\theta + \frac{\pi}{2} \right) = -BS \sin \theta \\ \phi_2 = BS \cos (\theta + \pi) = -BS \cos \theta \end{cases} \quad \text{donc} \quad \begin{cases} e_1 = BS \omega \cos \theta = Ri_1 & (EE_1) \\ e_2 = -BS \omega \sin \theta = Ri_2 & (EE_2) \end{cases}$$

Le moment magnétique induit est $\vec{\mathcal{M}} = i_1 \vec{S}_1 + i_2 \vec{S}_2 = i_1 S \begin{pmatrix} -\sin \theta \\ \cos \theta \\ 0 \end{pmatrix} + i_2 S \begin{pmatrix} -\cos \theta \\ -\sin \theta \\ 0 \end{pmatrix}$.

D'où le couple : $\vec{C} = \vec{\mathcal{M}} \wedge \vec{B} = BS (-i_1 \cos \theta + i_2 \sin \theta) \vec{e}_z = -\omega \frac{(BS)^2}{R} \vec{u}_z$.

Les phénomènes d'induction sont équivalents à un couple de frottement fluide qui s'oppose au mouvement.

Appliquons le théorème du moment dynamique à l'ensemble des deux cadres, projeté sur $\vec{u}_z$: $J \frac{d\omega}{dt} = C (EM)$, qui s'intègre, avec la condition initiale $\omega(0)$, en $\omega(t) = \omega_0 \exp \left(-\frac{(SB)^2}{RJ} t \right)$.

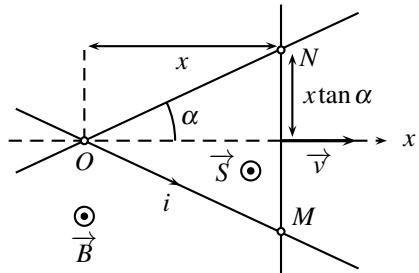
$$2. \begin{cases} (EM) \times \omega \Rightarrow \frac{d}{dt} \left(\frac{1}{2} J \omega^2 \right) = -BSi_1 \omega \cos \theta + BSi_2 \omega \sin \theta \\ (EE_1) \times i_1 \Rightarrow Ri_1^2 = BSi_1 \omega \cos \theta \\ (EE_2) \times i_2 \Rightarrow Ri_2^2 = -BSi_2 \omega \sin \theta \end{cases}$$

On en déduit : $\frac{d}{dt} \left(\frac{1}{2} J \omega^2 \right) = -R(i_x^2 + i_y^2)$. L'énergie cinétique des deux cadres diminue par effet Joule. Les cadres finissent par s'arrêter.

On observe de plus : $\vec{C} \cdot \vec{\omega} + e_x i_x + e_y i_y = \mathcal{P}_{Laplace} + \mathcal{P}_{fém} = 0$.

30.3 Rails de Laplace croisés

L'orientation du circuit avec i impose, avec la règle de la main droite, un vecteur surface dans la même direction que le champ magnétique.



1. La surface du circuit est $S = x \times x \tan \alpha$ et le flux à travers lui est : $\varphi = \vec{B} \cdot \vec{S} = Bx^2 \tan \alpha$. La f.é.m. induit dans le circuit est donc : $e(t) = -\frac{d\varphi}{dt} = B \tan \alpha \frac{d}{dt}(x^2(t)) = -B \tan \alpha 2x(t) \frac{dx}{dt}$, soit $e(t) = -2B \tan \alpha x(t) v$.

2. La longueur ℓ du circuit est : $\ell = 2x \tan \alpha + 2 \frac{x}{\cos \alpha} = 2x \frac{\sin \alpha + 1}{\cos \alpha}$. Si on néglige l'autoinduction dans le circuit, le schéma équivalent est constitué de la f.é.m. e induite et de la résistance $R = k\ell$. Ainsi : $e = Ri$ soit $-2B \tan \alpha xv = 2x \frac{\sin \alpha + 1}{\cos \alpha} ki$. Finalement : $i = -\frac{B}{k} \frac{\sin \alpha}{\sin \alpha + 1} v$.

L'intensité de ce courant est négative ; il crée, d'après la règle de la main droite, un champ magnétique orienté, au niveau du circuit, dans la direction opposée à celle de $\vec{B}$, conformément à la loi de Lenz.

3. La tige $\mathcal{T}$ est parcourue par un courant d'intensité i et est plongée dans un champ magnétique B , une force de Laplace $\vec{f}_L$ s'y exerce et, d'après la loi de Lenz, s'oppose à son mouvement. En notant M et N les deux points extrêmes de la $\mathcal{T}$: $\vec{f}_L = i \overrightarrow{MN} \wedge \vec{B} = iB2x \tan \alpha \vec{u}_x$. Appliquons le principe fondamental de la dynamique à la tige $\mathcal{T}$, en notant $\vec{f}_{op}$ la force exercée par l'opérateur, $m\vec{g}$ son poids et $\vec{R}$ les réactions des rails : $m \frac{d\vec{v}}{dt} = \vec{f}_L + \vec{f}_{op} + m\vec{g} + \vec{R}$.

La vitesse est constante, $\frac{d\vec{v}}{dt} = \vec{0}$. Sans frottements, $\vec{R}$ est orthogonale au plan des rails, tout comme le poids. En projection sur (Ox) : $0 = f_L + f_{op}$ donc $\vec{f}_{op} = -iB2x \tan \alpha \vec{e}_x = 2x \frac{B^2}{k} \frac{\sin \alpha}{\sin \alpha + 1} \tan \alpha v \vec{u}_x$.

4. La puissance P_J perdue par effet Joule a pour expression :

$$P_R = Ri^2 = 2x \frac{\sin \alpha + 1}{\cos \alpha} k \left(-\frac{B}{k} \frac{\sin \alpha}{\sin \alpha + 1} v \right)^2 = 2x \frac{B^2}{k} \frac{\sin \alpha}{\sin \alpha + 1} \tan \alpha v^2.$$

La puissance P_{op} fournie par l'opérateur est : $P_{op} = \vec{f}_{op} \cdot \vec{v} = 2x \frac{B^2}{k} \frac{\sin \alpha}{\sin \alpha + 1} \tan \alpha v^2$.

Les expressions sont identiques. Toute la puissance développée par l'opérateur est dégradée en effet Joule.

30.4 Cadre oscillant

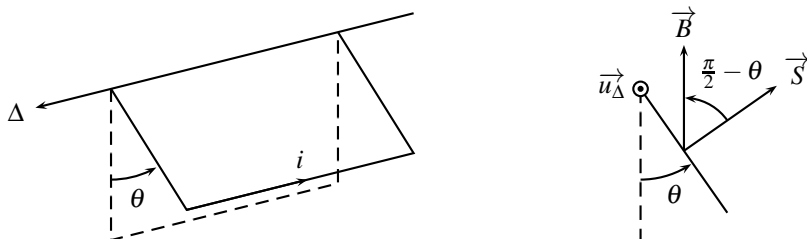
1. L'angle θ est orienté classiquement, d'une position fixe, ici verticale, vers une position variable, celle du cadre. Cette orientation de l'angle impose, d'après la règle de la main droite, celle de l'axe Δ de la droite vers la gauche : les quatre doigts de la main épousent l'arrondi, dans le sens de la flèche, le pouce renseigne sur la direction du vecteur $\vec{u}_\Delta$.

2. On oriente le cadre arbitrairement pour avoir $\vec{S}$ vers le haut sur le schéma, d'après la règle de la main droite.

Le flux magnétique à travers le circuit est : $\varphi = \vec{B} \cdot \vec{S} = Bab \cos \left(\frac{\pi}{2} - \theta \right) = Bab \sin \theta$, et la

f.é.m. induite : $e = -\frac{d\varphi}{dt} = -abB \frac{d\theta}{dt} \cos \theta$.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

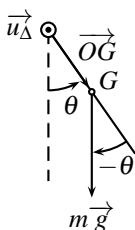


En négligeant l'autoinduction du cadre, le circuit n'est composé que de la f.é.m. e et de la résistance R . Ainsi : $e = Ri$ soit $-abB \frac{d\theta}{dt} \cos \theta = Ri$ (EE).

Le cadre est soumis à son poids $m\vec{g}$, à la force de Laplace $\vec{f}_L$ et la réaction du support.

Le moment des forces de Laplace, par rapport à l'axe de rotation Δ , est : $\vec{\Gamma}_L = \vec{\mathcal{M}} \wedge \vec{B} = i \vec{S} \wedge \vec{B} = iabB \sin\left(\frac{\pi}{2} - \theta\right) \vec{u}_\Delta = iabB \cos \theta \vec{u}_\Delta$.

Soit G le centre de gravité du cadre, en son centre. Le moment du poids par rapport à O est : $\vec{\Gamma}_p = \vec{OG} \wedge m\vec{g} = \frac{a}{2} mg \sin(-\theta) \vec{u}_\Delta = -\frac{a}{2} mg \sin \theta \vec{u}_\Delta$.



La liaison pivot est réputée parfaite, aucun moment n'est à considérer.

La loi du moment cinétique, projetée sur l'axe Δ , mène à :

$$J \frac{d\omega}{dt} = iabB \cos \theta - \frac{a}{2} mg \sin(\theta) = -\frac{(abB)^2}{R} \frac{d\theta}{dt} \cos^2 \theta - \frac{mga}{2} \sin \theta.$$

Au premier ordre en θ , $\cos \theta = 1$ et $\sin \theta = \theta$. Avec $\omega = \frac{d\theta}{dt}$, l'équation différentielle se

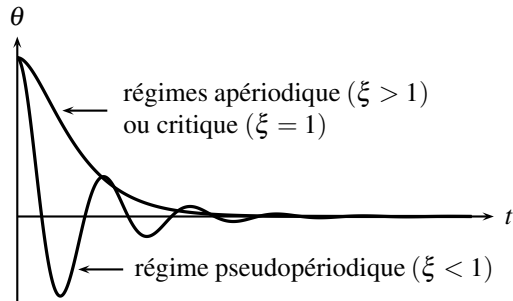
linéarise sous forme canonique : $\frac{2J}{mga} \frac{d^2\theta}{dt^2} + \frac{2ab^2B^2}{mgR} \frac{d\theta}{dt} + \theta(t) = 0$.

Le système est régi par une équation différentielle du deuxième ordre de paramètres canoniques :

$$\begin{cases} \frac{2J}{mga} = \frac{1}{\omega_0^2} \\ \frac{2ab^2B^2}{mgR} = \frac{2\xi}{\omega_0} \end{cases} \text{ donc } \begin{cases} \omega_0 = \sqrt{\frac{mga}{2J}} \\ \xi = \frac{ab^2B^2}{R} \sqrt{\frac{a}{2Jmg}} \end{cases}$$

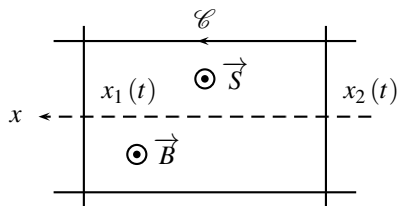
3. Les conditions initiales sont $\theta(0) = \theta_0$ et $\frac{d\theta}{dt}(0) = 0$ (vitesse initiale nulle). L'allure de

$\theta(t)$ dépend de la valeur du coefficient d’amortissement ξ par rapport à 1. Dans tous les cas, θ converge vers 0 (système du deuxième ordre avec une entrée nulle).

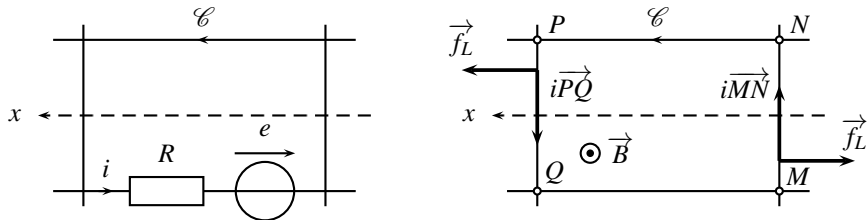


30.5 Deux tiges reliées par des ressorts

- 1. • Un barreau bouge sous la force de rappel du ressort,
 - le circuit composé des rails et des deux barreaux voit le flux magnétique varier et coupe des lignes de champ : f.é.m. induite puis courant induit i , car il est conducteur,
 - $\vec{B}$ et i : force de Laplace qui s’oppose au déplacement premier barreau (loi de Lenz) et met en mouvement le second.
2. Le circuit est orienté de façon à avoir le vecteur surface colinéaire et de même sens que le champ magnétique.



Le flux magnétique à travers le circuit est : $\varphi = \vec{B} \cdot \vec{S} = Ba(x_1(t) - x_2(t))$, et la f.é.m. induite : $e = \frac{d\varphi}{dt} = Ba(v_1(t) - v_2(t))$.



La loi des mailles mène à l’équation électrique : $e = Ri$ soit $Ba(v_1(t) - v_2(t)) = Ri(t)$ (EE). Les forces de Laplace sur les deux barreaux sont $\vec{f}_{L1} = iaB\vec{u}_x$ et $\vec{f}_{L2} = -iaB\vec{u}_x$.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

Attendu que les positions sont repérées par rapport à la position d'équilibre à vide, les forces de rappel des ressorts sont $\vec{f}_{R1} = -kx_1 \vec{u}_x$ et $\vec{f}_{R2} = -kx_2 \vec{u}_x$.

Les réactions des rails sur les barreaux sont orthogonales au déplacement, car celui-ci est sans frottement.

Le principe fondamental de la dynamique, appliqué aux deux barreaux, en projection sur $\vec{u}_x$, mène aux équations mécaniques :

$$m \frac{d^2 x_1}{dt^2} = -kx_1 + iaB \quad (EM_1), \quad m \frac{d^2 x_2}{dt^2} = -kx_2 - iaB \quad (EM_2),$$

soit : $m \frac{d^2 x_1}{dt^2} = -kx_1 - \frac{(aB)^2}{R} \left(\frac{dx_1}{dt} - \frac{dx_2}{dt} \right)$ et $m \frac{d^2 x_2}{dt^2} = -kx_2 + \frac{(aB)^2}{R} \left(\frac{dx_1}{dt} - \frac{dx_2}{dt} \right)$.

3. Sommons les deux dernières équations : $m \frac{d^2 \sigma}{dt^2} = -k\sigma(t)$ soit $\frac{m}{k} \frac{d^2 \sigma}{dt^2} + \sigma(t) = 0$.

Les conditions initiales sur $\sigma(t)$ sont : $\sigma(0) = x_1(0) + x_2(0) = b$ et $\left. \frac{d\sigma}{dt} \right|_0 = v_1(0) + v_2(0) = 0$.

La solution générale de l'équation différentielle est : $\sigma(t) = \alpha \cos\left(\sqrt{\frac{k}{m}}t\right) + \beta \sin\left(\sqrt{\frac{k}{m}}t\right)$.

Les conditions initiales imposent $\sigma(0) = \alpha = b$ et $\left. \frac{d\sigma}{dt} \right|_0 = \sqrt{\frac{k}{m}}\beta = 0$. Ainsi : $\sigma(t) = b \cos\left(\sqrt{\frac{k}{m}}t\right)$.

Effectuons maintenant la différence des deux équations : $m \frac{d^2 \delta}{dt^2} = -2 \frac{(aB)^2}{R} \frac{d\delta}{dt} - k\delta(t)$ soit $\frac{m}{k} \frac{d^2 \delta}{dt^2} + 2 \frac{(aB)^2}{Rk} \frac{d\delta}{dt} + \delta(t) = 0$.

Les conditions initiales sur $\delta(t)$ sont : $\delta(0) = x_1(0) - x_2(0) = b$ et $\left. \frac{d\delta}{dt} \right|_0 = v_1(0) - v_2(0) = 0$.

L'équation différentielle sur $\delta(t)$ modélise un système qui tend vers zéro en l'infini, c'est à dire $\delta(t) \xrightarrow[t \rightarrow \infty]{} 0$, soit $x_1(\infty) = x_2(\infty)$. Attendu qu'on ne demande que la solution de $x_1(t)$ et $x_2(t)$ qu'en régime permanent, c'est à dire mathématiquement pour $t \rightarrow \infty$, la solution en régime permanent de l'équation différentielle sur $\delta(t)$ est $x_1(t) = x_2(t)$.

Dans ce cas $\sigma(t) = 2x_1(t) = 2x_2(t)$ en régime permanent : $x_1(t) = x_2(t) = \frac{b}{2} \cos\left(\sqrt{\frac{k}{m}}t\right)$.

$$4. \begin{cases} (EM_1) \times v_1 \\ (EM_2) \times v_2 \\ (EE) \times i \end{cases} \text{ donc } \begin{cases} \frac{d}{dt} \left(\frac{1}{2} m \dot{x}_1^2 + \frac{1}{2} k x_1^2 \right) = iaB \frac{dx_1}{dt} \\ \frac{d}{dt} \left(\frac{1}{2} m \dot{x}_2^2 + \frac{1}{2} k x_2^2 \right) = -iaB \frac{dx_2}{dt} \\ ei = Ri^2 = -aBi \left(\frac{dx_1}{dt} - \frac{dx_2}{dt} \right) \end{cases}$$

En sommant ces trois équations : $\frac{d}{dt} \left(\frac{1}{2}mv_1^2 + \frac{1}{2}mv_2^2 + \frac{1}{2}kx_1^2 + \frac{1}{2}kx_2^2 \right) = -Rt^2$.

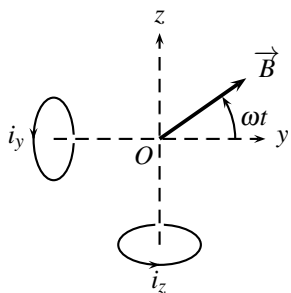
L'énergie mécanique des deux barreaux, sommes des énergies cinétique et potentielle élastique, diminue par effet Joule. Mais l'état final n'est pas un état d'énergie mécanique nulle. En effet, en régime permanent, les deux barreaux oscillent en phase (un ressort est comprimé, l'autre est allongé, et vice-versa), bien qu'il se déforme, la surface du circuit est une constante, il n'y a plus de phénomènes d'induction.

30.6 Machine asynchrone (MAS)

1. Il faut deux spires (de rayon a , situées à une distance d de O) en quadrature spatiale alimentées par des courants en quadrature temporelle :

$$\begin{cases} i_y = i_0 \cos \omega t \\ i_z = i_0 \sin \omega t \end{cases} \quad \text{ainsi} \quad \begin{cases} B_y = B_0 \cos \omega t \\ B_z = B_0 \sin \omega t \end{cases}$$

où B_0 est un coefficient qui dépend de la géométrie des deux spires et de leur position par rapport à O .



- 2. • La spire est initialement au repos et $\vec{B}$ tourne,
- le flux de $\vec{B}$ varie à travers la spire,
- f.é.m. induite dans la spire puis courant induit i ,
- couple de Laplace qui entraîne la spire en rotation (ou bien : la spire tourne pour orienter son moment magnétique $\vec{M}$ parallèlement à $\vec{B}$),
- la spire essaie de rattraper $\vec{B}$ pour arrêter la variation de flux (loi de Lenz).

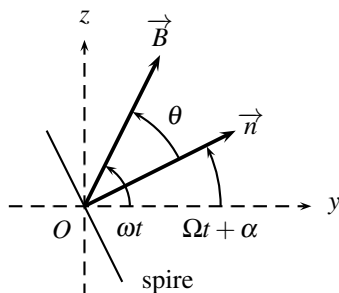
La spire peut-elle rattraper $\vec{B}$?

Si elle le rattrapait ($\omega = \Omega$) alors elle ne verrait plus de flux variable, il n'y aurait plus de f.é.m. induite ni de couple de Laplace. Elle s'arrêterait et serait de nouveau distancée par $\vec{B}$. On retrouverait $\omega \neq \Omega$.

Il y a donc toujours un léger **glissement** entre $\vec{B}$ et la spire, qui ne tourne pas au synchronisme avec $\vec{B}$, d'où le nom de machine asynchrone.

3. On rajoute l'angle α car $\vec{n}$ et $\vec{B}$ ne sont pas nécessairement colinéaires à $t = 0$: l'angle que fait $\vec{B}$ avec (Oy) est ωt , celui que fait $\vec{n}$ $\Omega t + \alpha$. Ainsi $\theta = (\omega - \Omega)t - \alpha$. Alors : $\phi = BS \cos \theta$ donc $e = -\frac{d\phi}{dt} = BS \dot{\theta} \sin \theta$ où $\dot{\theta} = \omega - \Omega$.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE



D'où l'équation électrique (EE) : $L \frac{di}{dt} + Ri = BS(\omega - \Omega) \sin((\omega - \Omega)t - \alpha)$.

En régime permanent, i est harmonique de même pulsation que le terme de droite, c'est à dire $\omega - \Omega$. Ainsi : $\underline{i} = \underline{i}_0 \exp(j(\omega - \Omega)t)$, et (EE) devient : $(jL(\omega - \Omega) + R)\underline{i} = BS(\omega - \Omega) \exp(j(\theta - \frac{\pi}{2}))$.

D'où l'intensité du courant dans la spire :

$$\begin{aligned} \underline{i} &= \frac{BS(\omega - \Omega)}{jL(\omega - \Omega) + R} \exp(j(\theta - \frac{\pi}{2})) \\ &= \frac{BS(\omega - \Omega)(R - jL(\omega - \Omega))}{R^2 + L^2(\omega - \Omega)^2} \exp(j(\theta - \frac{\pi}{2})) \\ &= \frac{BS(\omega - \Omega)(R - jL(\omega - \Omega))}{R^2 + L^2(\omega - \Omega)^2} (\cos(\theta - \frac{\pi}{2}) + j \sin(\theta - \frac{\pi}{2})) \\ \underline{i} &= \frac{BS(\omega - \Omega)}{R^2 + L^2(\omega - \Omega)^2} (R \sin \theta - L(\omega - \Omega) \cos \theta - jR \cos \theta - jL(\omega - \Omega) \sin \theta) \\ i(t) = \text{Re}(\underline{i}) &= \frac{BS(\omega - \Omega)}{R^2 + L^2(\omega - \Omega)^2} (R \sin \theta(t) - L(\omega - \Omega) \cos \theta(t)). \end{aligned}$$

4. Le couple s'exerçant sur la spire est : $\vec{\Gamma}(t) = \vec{\mathcal{M}} \wedge \vec{B} = i(t) S \vec{n} \wedge \vec{B} = i(t) SB \sin \theta(t) \vec{u}_x$.

En moyenne, sachant $\langle \sin^2 \theta(t) \rangle = \frac{1}{2}$ et $\langle \cos \theta(t) \sin \theta(t) \rangle = 0$:

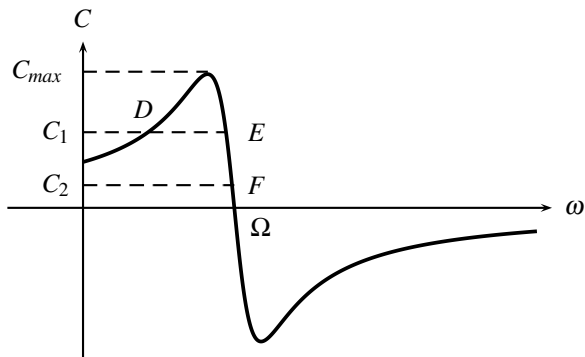
$$\vec{C} = \frac{R(BS)^2(\omega - \Omega)}{2[R^2 + L^2(\omega - \Omega)^2]} \vec{u}_x.$$

5. Le moteur ne peut pas développer un couple supérieur à C_{max} (cf figure suivante). Il peut démarrer seul car $C(\omega) > 0$.

Si le couple résistant vaut C_2 alors le point de fonctionnement du moteur est en F . Il tourne avec une vitesse angulaire proche de Ω (tout en lui restant inférieure).

Si le couple résistant vaut C_1 , deux points de fonctionnement sont possibles :

- D qui est instable : si le couple résistant augmente, alors la vitesse ω du moteur diminue (théorème du moment cinétique) et donc son couple aussi. Il ne peut plus entraîner la charge et décroche. Si le couple résistant diminue alors ω augmente et C aussi ; le moteur tourne de plus en plus vite et arrive en E .



- E qui est stable : si le couple résistant augmente, ω diminue mais C augmente. Le moteur continue à tourner normalement.

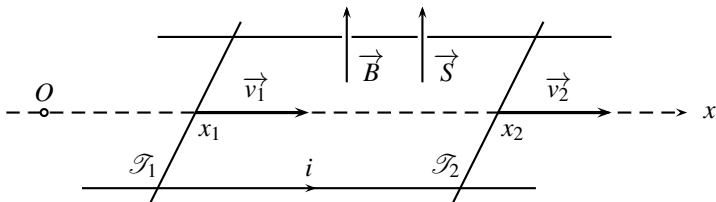
Si on entraîne le moteur par un couple extérieur $C_3 < 0$, alors $C < 0$, c'est à dire que la MAS fournit de la puissance. C'est un générateur (alternateur asynchrone).

30.7 Interaction entre deux tiges

1. • $\mathcal{T}_1$ bouge.
 - Le circuit formé des rails, de $\mathcal{T}_1$ et $\mathcal{T}_2$ coupe des lignes de champ et voit un flux variable : une f.é.m. e y est induite.
 - Le circuit est fermé et conducteur, un courant induit i circule.
 - i et $\vec{B}$ créent une force de Laplace $\vec{F}_1$ sur $\mathcal{T}_1$ qui s'oppose à son mouvement (loi de Lenz) : $\mathcal{T}_1$ **ralentit**.
 - i et $\vec{B}$ créent une force de Laplace $\vec{F}_2$ sur $\mathcal{T}_2$ qui la met en **mouvement**.

2. Le circuit constitué par les rails, $\mathcal{T}_1$ et $\mathcal{T}_2$ est orienté dans le même que le courant d'intensité i . D'après la règle de la main droite, $\vec{S} = S\vec{u}_z$.

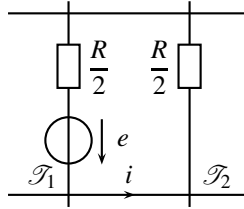
L'origine de l'axe Ox est arbitrairement choisie en O . Les positions des tiges sont notées x_1 et x_2 .



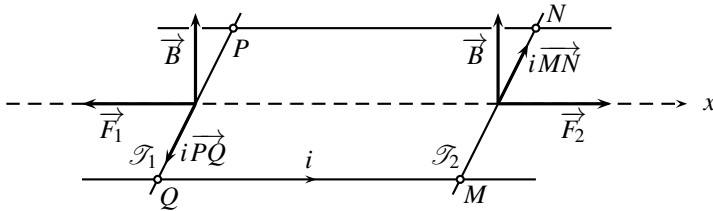
Le flux du champ magnétique à travers le circuit est : $\varphi = \vec{B} \cdot \vec{S} = BS = B(x_2 - x_1)a$. La f.é.m. induite vaut alors : $e = -\frac{d\varphi}{dt} = -B(v_2 - v_1)a$.

La résistance totale du circuit est R . La loi des mailles impose : $e = Ri = -B(v_2 - v_1)a$.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE



3. $\vec{F}_1 = i\vec{PQ} \wedge \vec{B} = -iaBu_x^{\rightarrow}$ et $\vec{F}_2 = iaBu_x^{\rightarrow}$.



Les lois de la quantité de mouvement, appliquées aux tiges $\mathcal{T}_1$ et $\mathcal{T}_2$, en projection sur $\vec{u}_x^{\rightarrow}$, mènent à : $m \frac{dv_1}{dt} = -iaB$ et $m \frac{dv_2}{dt} = iaB$.

4. Avec l'équation électrique : $m \frac{dv_1}{dt} = -\frac{(aB)^2}{R} (v_1 - v_2)$ et $m \frac{dv_2}{dt} = \frac{(aB)^2}{R} (v_1 - v_2)$.

5. Sommons les deux dernières équations : $m \frac{d}{dt} (v_1 + v_2) = 0$ soit $\frac{d\sigma}{dt} = 0$.

La condition initiale est : $\sigma(0) = v_1(0) + v_2(0) = v_0$.

La solution du couple {équation différentielle, condition initiale} est donc : $\sigma(t) = v_1(t) + v_2(t) = v_0$.

Faisons la différence des équations mécaniques : $m \frac{d}{dt} (v_1 - v_2) = m \frac{d\delta}{dt} = -\frac{2(aB)^2}{R} \delta(t)$.

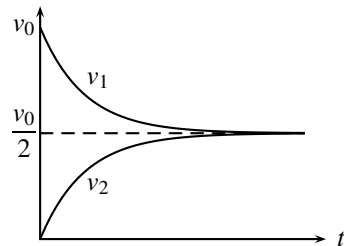
On reconnaît une équation différentielle du premier ordre, de constante de temps $\tau = \frac{Rm}{2(aB)^2}$:

$\tau \frac{d\delta}{dt} + \delta(t) = 0$. La condition initiale est : $\delta(0) = v_1(0) - v_2(0) = v_0$. La solution du couple {équation différentielle, condition initiale} est donc : $\delta(t) = v_1(t) - v_2(t) = v_0 \exp\left(-\frac{t}{\tau}\right)$.

Repassons à $v_1(t)$ et $v_2(t)$:

$$v_1(t) = \frac{\sigma(t) + \delta(t)}{2} = \frac{v_0}{2} \left(1 + \exp\left(-\frac{t}{\tau}\right)\right)$$

$$v_2(t) = \frac{\sigma(t) - \delta(t)}{2} = \frac{v_0}{2} \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$$



6. D'après l'équation électrique : $i(t) = \frac{Ba}{R}(v_1 - v_2) = \frac{Bav_0}{R} \exp\left(-\frac{t}{\tau}\right)$.

7. $i = \frac{dQ}{dt}$ ainsi $Q = \int_{t=0}^{\infty} i(t) dt = -\tau \frac{Bav_0}{R} \left[\exp\left(-\frac{t}{\tau}\right) \right]_0^{\infty} = \frac{mv_0}{2aB}$.

8. La puissance Joule est $p = Ri^2$. L'énergie dissipée par effet Joule est :

$$\begin{aligned} \mathcal{E}_J &= \int_{t=0}^{\infty} p dt = \frac{(Bav_0)^2}{R} \int_0^{\infty} \exp\left(-2\frac{t}{\tau}\right) dt \\ &= -\frac{\tau (Bav_0)^2}{2R} \left[\exp\left(-2\frac{t}{\tau}\right) \right]_0^{\infty} = \frac{1}{4} mv_0^2. \end{aligned}$$

9. Les rails sont horizontaux, l'énergie potentielle des tiges reste nulle :

$$\begin{aligned} \Delta \mathcal{E}_m &= \mathcal{E}_{C1}(\infty) + \mathcal{E}_{C2}(\infty) - \mathcal{E}_{C1}(0) - \mathcal{E}_{C2}(0) \\ \Delta \mathcal{E}_m &= \frac{1}{2} m \left(\left(\frac{v_0}{2}\right)^2 + \left(\frac{v_0}{2}\right)^2 - v_0^2 - 0 \right) = -\frac{mv_0^2}{4}. \end{aligned}$$

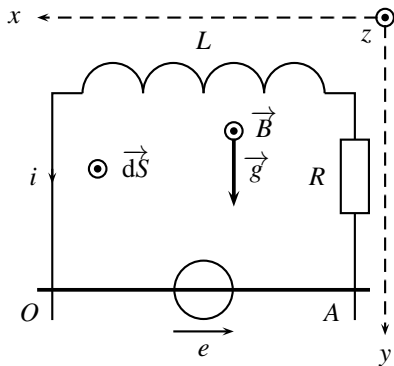
La moitié de l'énergie initiale du système est dissipée en effet Joule : $\Delta \mathcal{E}_m = -\mathcal{E}_J$. C'est la version intégrée par rapport au temps, entre $t = 0$ et $t \rightarrow \infty$, du bilan énergétique :

$$\frac{d}{dt} \left(\frac{1}{2} mv_1^2 + \frac{1}{2} mv_2^2 \right) = -Ri^2.$$

30.8 Tige qui chute

1. Comme la tige coupe des lignes de $\vec{B}$ et qu'on peut définir un flux magnétique φ variable, on peut aussi utiliser la loi de Faraday : $e = -\frac{d\varphi}{dt} = -\frac{d}{dt}(Bay) = -Ba\frac{dy}{dt} = -Bav$.

D'où l'équation électrique : $Ri + L\frac{di}{dt} = -Bva$ (EE).



La résultante des forces de Laplace est : $\vec{f}_L = i\vec{OA} \wedge \vec{B} = iaB\vec{u}_y$.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

D'où l'équation mécanique : $m \frac{dv}{dt} = mg + iaB$ (EM).

On obtient l'équation uniquement sur l'intensité du courant en dérivant l'équation électrique :

$$R \frac{di}{dt} + L \frac{d^2i}{dt^2} = -Ba \frac{dv}{dt}, \text{ puis en remplaçant } \frac{dv}{dt} \text{ avec l'équation mécanique : } R \frac{di}{dt} + L \frac{d^2i}{dt^2} = -Ba \left(g + \frac{Ba}{m} i(t) \right).$$

Écrivons le sous forme canonique $\left(\frac{1}{\omega_0^2} \frac{d^2i}{dt^2} + \frac{2\xi}{\omega_0} \frac{di}{dt} + i(t) = I_0 \right)$: $\frac{mL}{(aB)^2} \frac{d^2i}{dt^2} + \frac{mR}{(aB)^2} \frac{di}{dt} + i = -\frac{mg}{aB}$.

2. Si $R = 0$ alors : $\frac{mL}{(Ba)^2} \frac{d^2i}{dt^2} + i(t) = -\frac{mg}{Ba}$.

La pulsation caractéristique du système est $\omega_0 = \frac{aB}{\sqrt{mL}}$. La solution générale de l'équation différentielle du second ordre est : $i(t) = \alpha \cos(\omega_0 t) + \beta \sin(\omega_0 t) - \frac{mg}{aB}$.

La continuité du courant dans l'inductance impose $i(0) = 0$, soit $\alpha = \frac{mg}{Ba}$.

Avec, $R = 0$, (EE) écrite en $t = 0$ est : $L \left(\frac{di}{dt} \right)_{t=0} = -\underbrace{Bv(0)}_0 a$ donc $\left(\frac{di}{dt} \right)_{t=0} = 0$.

Attendu que $\frac{di}{dt} = -\omega_0 \alpha \sin(\omega_0 t) + \omega_0 \beta \cos(\omega_0 t)$, on a alors $\beta = 0$.

D'où : $i(t) = \frac{mg}{aB} (\cos \omega t - 1)$.

On remplace $i(t)$ dans (EM) : $\frac{dv}{dt} = g + g(\cos \omega t - 1) = g \cos \omega t$, qui s'intègre, sachant $v(0) = 0$, en $v(t) = \frac{g}{\omega} \sin \omega t$.

Ce mouvement harmonique surprenant est dû à $R = 0$: il n'y a pas de pertes d'énergie, donc des échanges permanents d'énergie entre les formes cinétique, potentielle et magnétique.

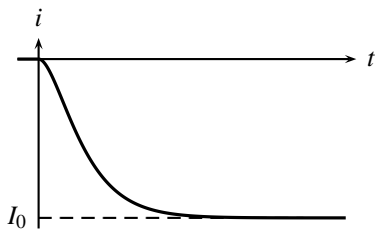
3. Si R est « grand » alors l'équation différentielle sur i décrit un système du second ordre très amorti. Identifions le facteur d'amortissement :

$$\left\{ \begin{array}{l} \omega_0 = \frac{aB}{\sqrt{mL}} \\ \frac{2\xi}{\omega_0} = \frac{mR}{(aB)^2} \end{array} \right. \text{ donc } \xi = \frac{R}{2aB} \sqrt{\frac{m}{L}}.$$

Le mouvement est très amorti si $\xi \gg 1$, soit : $R \gg 2aB \sqrt{\frac{L}{m}}$.

La réponse graphique est celle d'un système du second ordre très amorti soumis à un échelon $I_0 = -\frac{mg}{aB} < 0$:

Le système est alimenté par un terme constant $I_0 = -\frac{mg}{aB}$. Le régime permanent, ou établi, correspond donc à des valeurs constantes de l'intensité du courant et de la vitesse. Ainsi, en



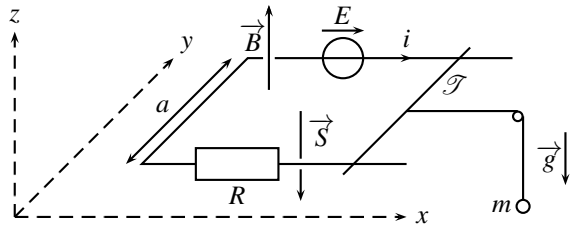
régime permanent, $\frac{di}{dt} = 0$, donc *via* (EE), v est une constante, donc $\frac{dv}{dt} = 0$. Alors (EE) et (EM) imposent les valeurs en régime permanent :

$$\begin{cases} 0 + Ri_{\infty} = -Bav_{\infty} \\ 0 = mg + Bai_{\infty} \end{cases} \quad \text{donc} \quad \begin{cases} v_{\infty} = \frac{mgR}{(aB)^2} \\ i_{\infty} = -\frac{mg}{aB} \end{cases}$$

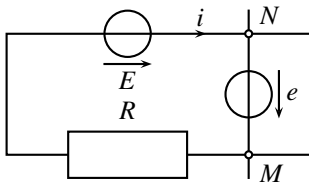
La vitesse de chute est limitée, on a freinage par induction.

30.9 Tige entraînée par gravité

- 1. Le fil reste tendu donc $v = -v_m$. Attention au signe : lorsque la tige bouge vers la droite ($v > 0$), la masse descend ($v_m < 0$).
- 2. On oriente arbitrairement le circuit dans le sens de E .



Le flux magnétique à travers le circuit est : $\varphi = \vec{B} \cdot \vec{S} = B\vec{e}_z \cdot S(-\vec{e}_z) = -BS = -Bax$.
La f.é.m. induite est alors, avec $v = \frac{dx}{dt}$: $e = -\frac{d\varphi}{dt} = Ba$. L'équation électrique est $E + e = Ri$ avec l'équivalent électrique du circuit :



- 3. La force de Laplace est : $\vec{f} = i\overrightarrow{MN} \wedge \vec{B} = -iaB\vec{u}_x$.

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

Pour établir la tension du fil, $\vec{T}$, appliquons le PFD à la masse m , de vitesse $\vec{v}_m = v_m \vec{u}_z$:
 $m \frac{d\vec{v}_m}{dt} = m \vec{g} + \vec{T}_m$ donc $\vec{T}_m = m \left(g + \frac{dv_m}{dt} \right) \vec{u}_z$.

Attention au signe $\oplus$ dans $\vec{v}_m = v_m \vec{u}_z$: le vecteur vitesse est orienté suivant l'axe (Oz) , quel que soit le mouvement. Si la masse chute, alors v_m est négatif et la formule reste la même.

Le fil est considéré comme sans masse. Ainsi la tension $\vec{T}$ est alors : $\vec{T} = m \left(g + \frac{dv_m}{dt} \right) \vec{u}_x$.

Or $v = -v_m$ donc $\vec{T} = m \left(g - \frac{dv}{dt} \right) \vec{u}_x$.

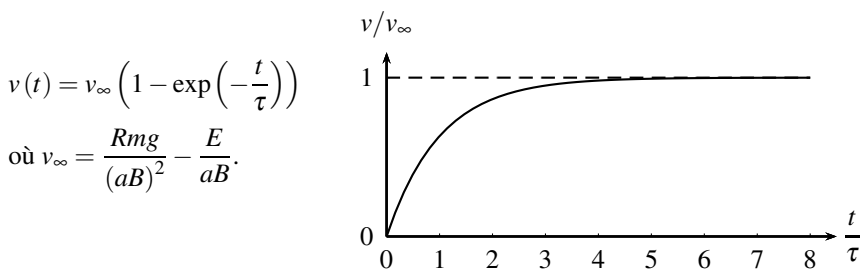
4. S'exercent sur $\mathcal{S}$: le poids $-m_0 g \vec{u}_z$, la réaction des rails $N \vec{u}_z$ (normale car pas de frottement), la force de Laplace $-iaB \vec{u}_x$ et la tension du fil $\vec{T} = T \vec{u}_x$.

Dès lors, le PFD, appliqué à la tige, projeté sur $\vec{u}_y$, mène à : $m_0 \frac{dv}{dt} = -iaB + m \left(g - \frac{dv}{dt} \right)$,

soit : $(m_0 + m) \frac{dv}{dt} = mg - iaB = mg - \frac{(aB)^2}{R} v - \frac{EaB}{R}$.

Introduisons le temps caractéristique τ du système : $\left[\frac{R(m_0 + m)}{(aB)^2} \right] \frac{dv}{dt} + v(t) = \frac{Rmg}{(aB)^2} - \frac{E}{aB}$.

5. Avec $v(0) = 0$, la solution du couple {équation différentielle, condition initiale} est :

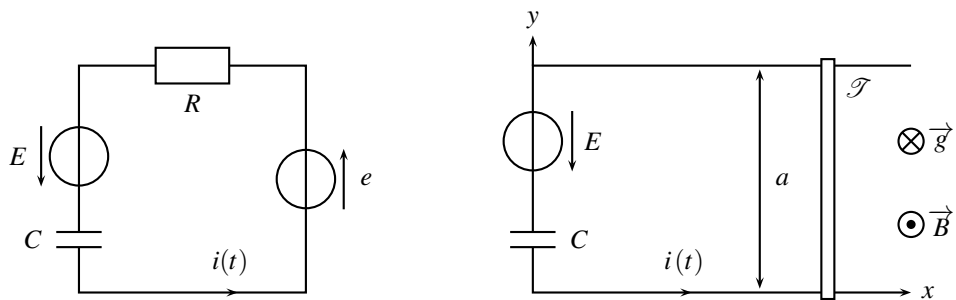


La masse m descend si $E < \frac{Rmg}{aB}$ ($v_\infty > 0$), reste immobile si $E = \frac{Rmg}{aB}$ et monte si $E > \frac{Rmg}{aB}$.

Remarquons que cette modélisation n'est valable que si le fil est tendu, c'est à dire si : $g - \frac{dv}{dt} > 0$ soit $g > \frac{v_\infty}{\tau} \exp\left(-\frac{t}{\tau}\right)$. Si cette relation est vérifiée à $t = 0$, alors elle le reste dans la suite, car $t \mapsto \exp\left(-\frac{t}{\tau}\right)$ est strictement décroissante.

30.10 Tige qui glisse sur un circuit capacitif

1. L'axe (Ox) est arbitrairement choisi vers la droite et l'axe (Oz) tel que $\vec{B} = B \vec{u}_z$ et $\vec{g} = -g \vec{u}_z$. Le circuit $\mathcal{C}$ est orienté par E , en convention générateur. Le vecteur surface du circuit est donc $\vec{S} = S \vec{u}_z$, d'après la règle de la main droite. Le flux magnétique à travers $\mathcal{C}$ est $\varphi = \vec{B} \cdot \vec{S} = Bax$; d'où une f.é.m. induite $e = -Bva$.



D'où l'équation électrique (EE), avec $i = C \frac{du_C}{dt}$: $E - u_C + e - Ri = 0$ soit $E - Bva = Ri + u_C$.

La force de Laplace $\vec{f}_L$ sur $\mathcal{T}$ est : $\vec{f}_L = iaB\vec{u}_x$.

Attendu que ni le poids, ni la réaction des rails, n'ont de composante suivant $\vec{u}_x$, le principe fondamental de la dynamique, appliqué à $\mathcal{T}$ en projection sur $\vec{u}_x$, mène à l'équation mécanique (EM) : $m \frac{dv}{dt} = iaB$ (EM).

2. Il faut maintenant trouver une équation sur $i(t)$ depuis (EE) et (EM). Attendu que $i = C \frac{du_C}{dt}$, on pense à dériver (EE) par rapport au temps : $0 - Ba \frac{dv}{dt} = R \frac{di}{dt} + \frac{i}{C} = -Ba \frac{iaB}{m}$, qui s'écrit sous forme canonique : $\tau \frac{di}{dt} + i(t) = 0$ où $\tau = \frac{mRC}{m + C(aB)^2}$. Ainsi $i(t) = \lambda \exp\left(-\frac{t}{\tau}\right)$.

La tension u_C est continue donc nulle en $t = 0$, la tige initialement immobile ($e = 0$) donc (EE) devient en $t = 0$: $E - 0 = Ri(0) + 0$ donc $\lambda = \frac{E}{R}$, ainsi $i(t) = \frac{E}{R} \exp\left(-\frac{t}{\tau}\right)$.

3. On en déduit avec (EM) : $\frac{dv}{dt} = \frac{aB}{m} \frac{E}{R} \exp\left(-\frac{t}{\tau}\right)$ donc $v(t) = -\frac{aB}{m\tau} \frac{E}{R} \exp\left(-\frac{t}{\tau}\right) + \mu$, où la constante μ est trouvée avec $v(0) = 0$. Au final : $v(t) = \frac{aB}{m\tau} \frac{E}{R} \left[1 - \exp\left(-\frac{t}{\tau}\right)\right]$.

4. La puissance délivrée par le générateur est $Ei(t)$. Ainsi, l'énergie délivrée entre $t = 0$ et $t \rightarrow \infty$ est :

$$\mathcal{E}_G = \int_0^\infty E i dt = \frac{E^2}{R} \int_0^\infty \exp\left(-\frac{t}{\tau}\right) dt = \frac{E^2 \tau}{R}.$$

5. $\begin{cases} i(t) = C \frac{du_C}{dt} = \frac{E}{R} \exp\left(-\frac{t}{\tau}\right) \\ u_C(0) = 0 \end{cases}$ donc $u_C(t) = \frac{E\tau}{RC} \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$.

6. $\mathcal{E}_G = \int_0^\infty u_C i dt = C \int_0^\infty u_C \frac{du_C}{dt} dt = C \int_0^\infty \frac{d}{dt} \left(\frac{u_C^2}{2}\right) dt = \frac{1}{2} u_C(\infty)^2 - \frac{1}{2} u_C(0)^2 = \frac{E^2 \tau^2}{2R^2 C}.$

7. $\mathcal{E}_J = \int_0^\infty Ri^2 dt = \frac{E^2}{R} \int_0^\infty \exp\left(-2\frac{t}{\tau}\right) dt = \frac{E^2 \tau}{2R}.$

8. La puissance de la force de Laplace est $\vec{f}_L \cdot \vec{v} = iaBv$. Avec l'équation mécanique, cette

CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

puissance devient $m \frac{dv}{dt} v = \frac{d}{dt} \left(\frac{1}{2} m v^2 \right)$. D'où le travail demandé :

$$W = \int_0^\infty \frac{d}{dt} \left(\frac{1}{2} m v^2 \right) dt = \frac{1}{2} m v^2 (\infty) - \frac{1}{2} m v^2 (0) = \frac{a^2 B^2 E^2}{2 m \tau R^2}.$$

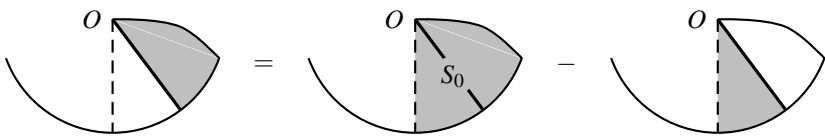
Ce dernier résultat constitue le théorème de l'énergie cinétique.

9. $\mathcal{E}_G = \mathcal{E}_C + \mathcal{E}_J + W$: l'énergie fournie par le générateur sert à charger le condensateur, à mettre en mouvement la tige et à compenser les pertes Joules.

30.11 Tige en rotation sur un cercle (PTSI)

1. Il convient tout d'abord d'orienter le circuit. On choisit arbitrairement d'orienter la tige du centre vers la périphérie.

Quelle est la surface du circuit ? L'aire grisée est l'aire S_0 moins celle du secteur angulaire d'angle θ :



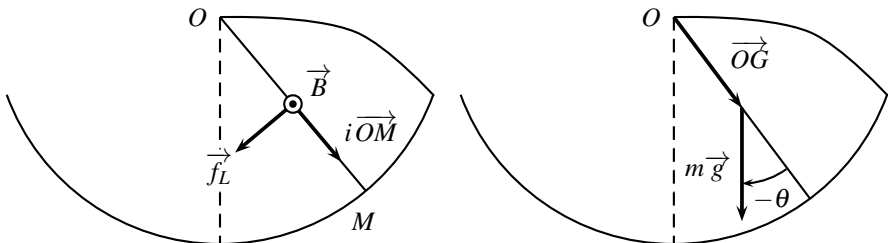
Quelle est la surface du secteur angulaire d'angle θ ? Pour un angle de 2π , la surface est πa^2 , pour un angle $\theta < 2\pi$, elle vaut $\theta \frac{a^2}{2}$ (règle de proportionnalité). Ainsi, la surface du circuit est : $S = S_0 - \theta \frac{a^2}{2}$.

Avec une tige orientée du centre vers la périphérie, le vecteur surface est, d'après la règle de la main droite, de même direction que le champ magnétique. Le flux à travers le circuit est donc : $\varphi = \vec{B} \cdot \vec{S} = B \left(S_0 - \theta \frac{a^2}{2} \right)$.

La f.é.m. induite est donc : $e = - \frac{d\varphi}{dt} = \frac{d\theta}{dt} \frac{B a^2}{2}$.

En négligeant l'autoinductance du circuit, le schéma électrique se résume à la f.é.m. et à la résistance R de la tige d'où l'équation électrique (EE) : $e = Ri$ soit $\frac{d\theta}{dt} \frac{B a^2}{2} = Ri$.

L'équation mécanique s'obtient en appliquant le théorème du moment dynamique à la tige en O . Il faut tenir compte du poids et de la force de Laplace.



Le poids s'exerce au centre de gravité au milieu de la tige :

$$\vec{\mathcal{M}}_{\text{poids}} = \vec{OG} \wedge m \vec{g} = -\frac{a}{2} mg \sin \theta \vec{u}_z.$$

Le moment du poids est affecté du signe $-$ car on va de $\vec{OG}$ vers $m \vec{g}$, le poids tend bien à remettre la barre dans sa position verticale, donc exerce un moment négatif en projection sur $\vec{u}_z$.

La force de Laplace est $\vec{f}_L = i \vec{OP} \wedge \vec{B} = -iaB \vec{u}_\theta$. On admet qu'elle s'applique au milieu G du segment $[OM]$. Son moment par rapport à O est donc :

$$\vec{\mathcal{M}}_L = \vec{OG} \wedge \vec{f}_L = \frac{a}{2} \vec{u}_r \wedge (-iaB \vec{u}_\theta) = -\frac{iBa^2}{2} \vec{u}_z.$$

Avec une liaison pivot parfaite, le théorème du moment dynamique, appliqué à la tige, en O , en projection sur $\vec{u}_z$, mène à l'équation mécanique (EM) : $J \frac{d^2 \theta}{dt^2} = -\frac{iBa^2}{2} - \frac{a}{2} mg \sin \theta$.

On élimine i via (EE) : $J \frac{d^2 \theta}{dt^2} = -\frac{1}{R} \left(\frac{Ba^2}{2} \right)^2 \frac{d\theta}{dt} - \frac{a}{2} mg \sin \theta$.

On reconnaît un système du deuxième ordre qu'on écrit sous forme canonique dans l'hypothèse de petits angles ($\sin \theta = \theta$ au deuxième ordre en θ) : $\frac{2J}{amg} \frac{d^2 \theta}{dt^2} + \frac{B^2 a^3}{2Rmg} \frac{d\theta}{dt} + \theta(t) = 0$, d'où :

$$\left\{ \begin{array}{l} \frac{2J}{amg} = \frac{1}{\omega_0^2} \\ \frac{B^2 a^3}{2Rmg} = \frac{2\xi}{\omega_0} \end{array} \right. \text{ donc } \left\{ \begin{array}{l} \omega_0 = \sqrt{\frac{amg}{2J}} \\ \xi = \frac{B^2 a^3}{4R} \sqrt{\frac{a}{2Jmg}} \end{array} \right.$$

2. On note $\omega = \frac{d\theta}{dt}$ la vitesse angulaire de la tige :

$$\left\{ \begin{array}{l} (EM) \times \omega \Rightarrow \frac{d}{dt} \left(\frac{1}{2} J \omega^2 \right) = -\frac{iBa^2 \omega}{2} - \frac{a}{2} mg \frac{d\theta}{dt} \sin \theta \\ (EE) \times i \Rightarrow Ri^2 = \frac{iBa^2 \omega}{2} \end{array} \right.$$

En sommant les deux équations et en reconnaissant $\frac{d}{dt} (\cos \theta(t)) = -\frac{d\theta}{dt} \sin \theta(t)$:

$$\frac{d}{dt} \left(\frac{1}{2} J \omega^2 - mg \frac{a \cos \theta}{2} \right) = -Ri^2.$$

Notons $\vec{u}_x$ l'axe cartésien orienté dans la direction de la verticale ascendante. Alors la position du centre de gravité G de la tige est en $x_G = -\frac{a}{2} \cos \theta$ et le bilan énergétique devient :

$$\frac{d}{dt} \left(\frac{1}{2} J \omega^2 + mg x_G \right) = -Ri^2.$$

L'énergie mécanique de la tige ne peut que décroître par effet Joule. La tige finit par s'arrêter.

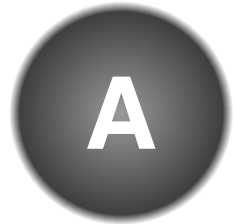
CHAPITRE 30 – CIRCUIT MOBILE DANS UN CHAMP MAGNÉTIQUE STATIONNAIRE

Sixième partie

Appendices

Trouver la bibliothèque des livres ici : <https://1024terabox.com/s/1mg0wxl22G8NjM8MA2RzgFw>

Mesures et incertitudes



Dans ce chapitre, on s'intéresse à la mesure d'une grandeur physique, à sa précision et au niveau de confiance que l'on peut lui accorder. La science de la mesure s'appelle la **métrologie**. C'est une discipline à part entière de la physique. Elle est codifiée par des normes internationales qui définissent un vocabulaire spécifique et des usages précis. Un des textes de référence est le *Guide pour l'expression de l'incertitude de mesure* publié par le *Bureau International des Poids et Mesures*. Le but de ce chapitre est d'explicitier une méthode permettant de déterminer une incertitude de mesure lors d'une séance de travaux pratiques sans surcharger le lecteur de subtilités.

1 Mesure d'une grandeur physique

La compréhension des phénomènes physiques est construite par un va-et-vient entre des modèles théoriques et des observations expérimentales. La validation des modèles théoriques est assurée par la comparaison de ses prévisions à des résultats expérimentaux. Comme on ne sait comparer que des nombres, il faut chiffrer les grandeurs physiques donc les mesurer.

1.1 Représentation d'une grandeur physique

Pour connaître et manipuler une grandeur physique, il faut la représenter. En pratique, on la représente par une lettre (représentation littérale) ou par un nombre (représentation numérique) que l'on obtient en la mesurant ou en la calculant avec un modèle.

a) Qu'est-ce qu'une mesure physique ?

Mesurer une grandeur physique, c'est la chiffrer par comparaison à une grandeur de même nature choisie comme unité. La représentation numérique d'une grandeur physique doit systématiquement être accompagnée de son unité.

CHAPITRE A – MESURES ET INCERTITUDES

Exemple

La largeur l (représentation littérale) d'une feuille de papier A4 est une grandeur physique que l'on représente par les nombres 0,210 m dans le système métrique et 0,869 feet dans le système anglo-saxon.

Il est possible d'écrire $l = 0,210 \text{ m} = 0,869 \text{ feet}$ mais absurde d'écrire $0,210 = 0,869 \dots$

La représentation numérique d'une grandeur physique n'est unique que si elle est accompagnée d'une unité. Une représentation numérique sans précision de l'unité n'a aucun sens.

b) Nécessité d'un système d'unités cohérent au niveau international

La première norme internationale consiste en un ensemble complet d'unités appelé **Système International d'unités**. Ces unités sont définies par une convention internationale dont l'organe scientifique est le *Bureau International des Poids et Mesures* basé à Sèvres en banlieue parisienne. Cet organisme scientifique international est chargé de définir les unités et de préciser les méthodes qui permettent d'exploiter ces définitions. Le choix des unités est arbitraire. Avant l'avènement du Système International d'unités, chaque pays utilisait son propre système. Cela posait d'énormes problèmes de communication.

Exemple

Pour vous faire une idée du casse-tête, essayez de parler de la consommation d'essence d'une voiture avec un Américain : une voiture qui consomme 8 L d'essence aux 100 km pour un Français permet à un Américain de rouler 30 miles par gallon d'essence.

c) Vocabulaire de la métrologie

Tout comme les unités, le vocabulaire de métrologie a été codifié au niveau international pour permettre une communication scientifique précise. Dans ce texte, on indique les termes de ce vocabulaire à connaître en les écrivant en gras lorsqu'ils apparaissent pour la première fois.

1.2 Mesure d'une grandeur physique

Le processus d'attribution d'une valeur expérimentale à une grandeur physique s'appelle le **mesurage**. Il est réalisé à l'aide d'un instrument de mesure qui fournit une **valeur mesurée**. La grandeur soumise au mesurage s'appelle le **mesurande**.

a) Erreur de mesure

Aucune valeur mesurée n'est rigoureusement exacte : une **erreur de mesure** qu'on espère minimale y est systématiquement associée. Cette erreur dépend des instruments de mesure et du protocole utilisé. Considérant une grandeur à mesurer, on note :

- x_v la **valeur vraie** de cette grandeur ;
- x la **valeur mesurée**, résultat du mesurage effectué.

La valeur vraie est généralement inconnue puisqu'on cherche à la mesurer. L'**erreur de mesure**, défini comme l'écart entre la valeur mesurée x et la valeur vraie x_v est donc également inconnue. Suivant la manière dont elles évoluent lorsque l'on répète plusieurs fois le même protocole de mesure, on classe les erreurs de mesure en deux catégories :

- les **erreurs systématiques** qui ne varient pas lors de mesures répétées ;
- les **erreurs aléatoires**, dont le signe et la valeur varient aléatoirement.

On illustre souvent ces deux types d'erreurs par une analogie avec un tir sur cible répété :

- lorsque le fusil est mal réglé, les impacts sont systématiquement décalés dans une direction donnée. Il s'agit d'une erreur systématique ;
- lorsque le tireur bouge au moment où il déclenche son tir, les impacts sont dispersés de manière aléatoire. Il s'agit d'une erreur aléatoire.

Sur la figure A.1, le centre de la cible est matérialisé par une étoile. Les impacts forment un nuage de point dont le centre est matérialisé par un disque gris. Les quatre cas (a), (b), (c) et (d) présentent une dispersion des impacts autour du centre d'impact. Cette dispersion est la caractéristique des erreurs aléatoires qui sont faibles sur les figures (a) et (c) et plus importantes sur les figures (b) et (d). Par ailleurs, le centre du nuage de point des figures (a) et (b) est confondu avec le centre de la cible (de sorte que le disque gris est confondu avec l'étoile) alors que celui des figures (c) et (d) est décalé par rapport au centre de la cible. Sur les figures (c) et (d), les impacts présentent une erreur systématique.

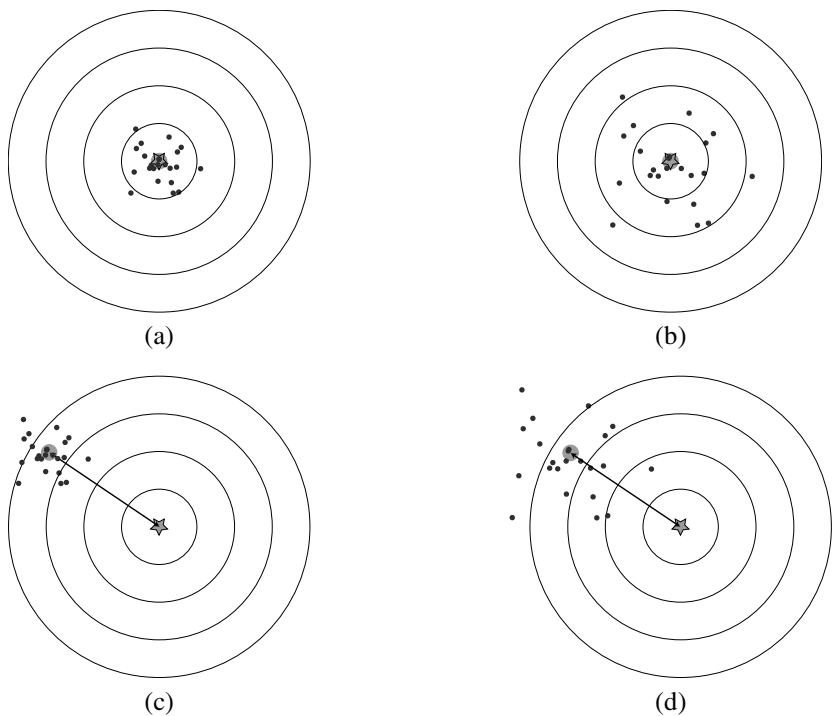


Figure A.1 – Impact sur une cible. Les figures (a) et (b) présentent uniquement des erreurs aléatoires. Les figures (c) et (d) présentent des erreurs aléatoires et une erreur systématique.

CHAPITRE A – MESURES ET INCERTITUDES

b) Erreurs systématiques

Il est très difficile de détecter une erreur systématique puisque le décalage entre la valeur mesurée et la valeur vraie (inconnue) n'est pas connu.

Quelles sont les causes des erreurs systématiques ? Les erreurs systématiques sont généralement dues à des instruments de mesure mal étalonnés (erreurs matérielles) ou à une modélisation théorique de l'expérience insuffisamment précise (erreurs théoriques).

Exemple

Si l'on cherche à mesurer la largeur l d'une feuille de papier A4 avec un mètre ruban en tissu qui mesure 1,01 m au lieu de 1,00 m, on va systématiquement trouver une longueur inférieure à sa valeur vraie qui vaut 0,210 m. Il s'agit d'une erreur matérielle.

Lorsque l'on détermine une résistance par le rapport d'une tension à une intensité mesurées, on fait une erreur systématique si le modèle ne tient pas compte des résistances de l'ampèremètre et du voltmètre. Il s'agit d'une erreur théorique.

Les problèmes d'étalonnage d'un instrument de mesure ont des causes variées, comme :

- les défauts de conception de l'instrument (dans l'exemple précédent, dès son achat, le mètre ruban est déjà trop long) ;
- son usure (dans l'exemple précédent, le mètre ruban s'est allongé à force d'être utilisé) ;
- l'influence de paramètres extérieurs sur son comportement (dans l'exemple précédent, c'est l'humidité de l'air qui a provoqué la distension du mètre en tissu).

Dans ce dernier cas, les métrologues parlent de **grandeur d'influence**, qui sont les paramètres expérimentaux qui influent sur la réponse de l'instrument de mesure alors qu'ils n'ont pas d'incidence sur la grandeur mesurée. Les plus courantes sont la température et la fréquence d'échantillonnage des instruments numériques mais il en existe de nombreuses autres.

Remarque

Pour pouvoir déterminer les grandeurs d'influence d'un instrument de mesure, il faut avoir une idée assez précise de son mode de fonctionnement.

Comment peut-on remédier à une erreur systématique ? Si l'erreur systématique est une erreur théorique, il faut reprendre et corriger le modèle. Si l'erreur systématique est une erreur matérielle, il faut reprendre l'étalonnage de la chaîne de mesure. Pour cela, on utilise un étalon de mesure dont on connaît la valeur vraie. On mesure cet étalon, on obtient sa valeur mesurée et on détermine l'erreur systématique par différence entre ces deux valeurs. On peut ensuite corriger cette erreur sur les mesures déjà effectuées.

Remarque

L'étalonnage de la chaîne de mesure est le seul cas où l'on connaît la valeur vraie.

c) Erreurs aléatoires

Les erreurs aléatoires dispersent les valeurs mesurées autour d'une valeur centrale. Elles peuvent être détectées en répétant plusieurs fois la mesure mais ne peuvent en aucun cas

être corrigées. Sur la figure A.1, les mesures sont plus dispersées dans le cas (b) que dans le cas (a), ce qui signifie que les erreurs aléatoires sont plus importantes en (b) qu'en (a). Les erreurs aléatoires rendent incertaines les valeurs mesurées. Dans le paragraphe suivant, on montre comment estimer cette **incertitude de mesure**.

2 Incertitudes et intervalle de confiance

Étant donné que l'on ne peut pas faire une mesure parfaite, on ne peut accéder ni à la valeur vraie x_v ni à l'erreur de mesure $x - x_v$. Le mieux que l'on puisse faire est de proposer un **intervalle de confiance** dans lequel la valeur vraie x_v se trouve avec une probabilité donnée. Dans ce paragraphe, on cherche à estimer les bornes de l'intervalle de confiance, c'est-à-dire à évaluer l'incertitude de mesure. Pour cela, on suppose que les erreurs systématiques ont été corrigées et on évalue l'importance des erreurs aléatoires.

2.1 Notion d'intervalle de confiance

La détermination de l'intervalle de confiance est basée sur l'hypothèse que, en l'absence d'erreur systématique, la mesure de x suit une loi normale centrée sur la valeur vraie x_v et d'écart-type σ . Cette hypothèse est justifiée par un théorème mathématique appelé *Théorème Central Limite*. On utilise alors un résultat du cours de mathématique de Terminale qui stipule que, lorsqu'une variable aléatoire suit une loi normale $\mathcal{N}(x_v, \sigma^2)$:

- la probabilité d'un tirage dans l'intervalle $[x_v - \sigma, x_v + \sigma]$ est d'environ 68% ;
- la probabilité d'un tirage dans l'intervalle $[x_v - 2\sigma, x_v + 2\sigma]$ est d'environ 95% ;
- la probabilité d'un tirage dans l'intervalle $[x_v - 3\sigma, x_v + 3\sigma]$ est d'environ 99,7%.

Tout le problème est alors d'évaluer expérimentalement x_v et σ . Dans un premier temps, on suppose connue $\bar{x}$ la meilleure estimation de x_v et $u(x)$ la meilleure estimation de σ . Dans le paragraphe 2.2, on montre comment évaluer concrètement les valeurs de $\bar{x}$ et $u(x)$.

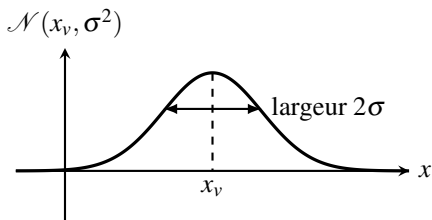


Figure A.2 – Loi de probabilité d'un mesurage x .

a) Incertitude-type, intervalle de confiance à 68%

L'écart-type σ mesure la largeur de la loi de probabilité $\mathcal{N}(x_v, \sigma^2)$. Sa meilleure estimation $u(x)$ est liée à la dispersion des données autour de x_v et on l'appelle l'**incertitude-type de mesure**. En remplaçant x_v et σ par leurs meilleures estimations $\bar{x}$ et $u(x)$, la valeur vraie x_v est dans l'intervalle $[\bar{x} - u(x), \bar{x} + u(x)]$ avec un niveau de confiance approximativement égal à 68% ce que l'on note :

$$x = \bar{x} \pm u(x) \text{ à } 68\%.$$

L'incertitude-type est notées $u(x)$ en référence au mot anglais *uncertainty*.

On définit également l'**incertitude-type relative** $u_r(x) = \frac{u(x)}{\bar{x}}$ que l'on exprime généralement à l'aide d'un pourcentage.

CHAPITRE A – MESURES ET INCERTITUDES

b) Incertitude élargie, intervalle de confiance à 95%

Si l'on veut augmenter le niveau de confiance, il faut agrandir l'intervalle de confiance en introduisant une **incertitude élargie** $k \times u(x)$ avec $k > 1$. La valeur de k dépend du niveau de confiance que l'on désire atteindre. Sous les mêmes hypothèses que précédemment, k est proche de 2 pour atteindre un niveau de confiance de 95%. La valeur vraie x_v est alors dans l'intervalle $[\bar{x} - 2u(x), \bar{x} + 2u(x)]$ avec un niveau de confiance approximatif de 95% ce que l'on note :

$$x = \bar{x} \pm 2u(x) \text{ à } 95\%.$$

Remarque

On peut avoir besoin d'un niveau de confiance encore plus grand. Dans ce cas, il faut encore élargir l'incertitude en augmentant la valeur de k .

2.2 Évaluation d'une incertitude-type

On cherche maintenant à évaluer concrètement les valeurs de $\bar{x}$ et $u(x)$ qui sont les meilleures estimations de x_v et σ .

a) Évaluation de type A

L'évaluation de type A est une évaluation statistique de ces deux grandeurs. Pour la pratiquer, il faut réaliser au minimum 5 à 10 mesures et elle sera d'autant plus précise que le nombre de mesures sera grand. Elle nécessite donc d'avoir le temps. Elle est particulièrement utilisée lorsqu'il est difficile de déterminer les sources d'erreurs.

On note $x_1, x_2, \dots, x_n$ les valeurs mesurées au cours de n prises de mesures indépendantes. On définit la moyenne $\bar{x}$ et l'écart-type $s(x)$ de l'échantillon de n mesures :

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i \quad \text{et} \quad s(x) = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2}.$$

En l'absence d'erreur systématique, la valeur moyenne $\bar{x}$ est la meilleure estimation de x_v et l'écart-type expérimental $s(x)$ mesure la dispersion des données autour de leur valeur moyenne. Les calculatrices permettent de calculer ces deux valeurs et les notent généralement « $\bar{x}$ » pour la moyenne de l'échantillon et « s_x » ou « $x\sigma_{n-1}$ » pour son écart-type expérimental.

Remarque

Les calculatrices donnent aussi un autre écart-type : $\sigma(x) = \sqrt{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2}$ noté « σ_x » ou « $x\sigma_n$ ». Pour distinguer $s(x)$ de $\sigma(x)$, on peut remarquer que $s(x) > \sigma(x)$.

On estime x_v grâce à la valeur moyenne des n mesures. C'est donc l'écart-type sur la moyenne qui caractérise l'incertitude de la détermination de x_v . **L'écart-type expérimental de la moyenne** des n mesures est relié à l'écart-type expérimental d'une seule mesure par la relation :

$$u(x) = \frac{s(x)}{\sqrt{n}}.$$

L'incertitude-type décroît avec n car en calculant la moyenne $\bar{x}$, on additionne des erreurs aléatoires qui se compensent statistiquement. Plus le nombre de mesure est grand, plus la compensation est bonne et plus l'estimation de x_v par $\bar{x}$ est précise.

L'évaluation de type A d'une incertitude de mesure basée sur n mesures $x_1, x_2, \dots, x_n$ indépendantes et dépourvues d'erreurs systématiques conduit à l'**incertitude-type** :

$$u(x) = \frac{s(x)}{\sqrt{n}},$$

où $s(x) = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2}$ est l'écart-type expérimental d'une mesure qui caractérise la dispersion des valeurs mesurées autour de leur valeur moyenne.

On peut être tenté d'augmenter le nombre de mesures pour diminuer l'incertitude. Cependant, la dépendance en $\frac{1}{\sqrt{n}}$ rend cette augmentation peu efficace au delà de $n = 10$ mesures : il faut multiplier par 100 le nombre de mesures pour gagner un facteur 10 dans la précision. Ceci n'est donc envisageable que lorsque l'on peut automatiser les mesures.

b) Évaluation de type B

Une **évaluation de type B** est réalisée par tout autre moyen que le traitement statistique. Elle est utilisée lorsqu'il est impossible ou trop long de faire une évaluation de type A. Une connaissance générale de l'expérience est nécessaire pour rechercher et évaluer les sources d'erreurs. On détermine alors l'*intervalle des valeurs mesurées raisonnablement acceptables*. Cet intervalle a pour demi-largeur la précision de la mesure Δ qui dépend du type d'instrument et de la méthode de mesure utilisés. La difficulté est que la notion de valeur « acceptable » est relative et dépend du vécu de l'expérimentateur. . .

La précision Δ des instruments de mesure gradués est égale à une demi graduation.

La précision Δ des autres instruments notamment numériques est à chercher sur la notice du constructeur.

La précision Δ de la méthode de mesure est à déterminer expérimentalement.

Exemple

- Sur une règle graduée au mm, la précision vaut : $\Delta = 0,5 \text{ mm}$.
- Sur la notice d'un multimètre numérique utilisé en voltmètre DC sur le calibre 5V, on lit : « Accuracy : 0.3% rdg + 2 digits ». La précision de l'appareil est 0,3% de la valeur lue (rdg signifie **reading**, soit valeur lue) auquel on ajoute 2 fois la valeur du dernier chiffre affiché. Pour une valeur lue de 2,5462V, la précision vaut :
$$\Delta = \frac{0,3}{100} \times 2,5462 + 0,0002 = 0,0076 \text{ V}.$$
- Lors d'une expérience sur un banc d'optique, on place un objet lumineux, une lentille et un écran et on déplace l'écran pour obtenir une image nette. La latitude de mise

CHAPITRE A – MESURES ET INCERTITUDES

au point permet de positionner l'écran sur ± 5 mm autour d'une position centrale. La précision vaut : $\Delta = 5$ mm.

L'évaluation de type B de l'incertitude-type d'une mesure effectuée sur un instrument de précision Δ est :

$$u(x) = \frac{\Delta}{\sqrt{3}}.$$

Le facteur $\frac{1}{\sqrt{3}}$, préconisé par une norme internationale, est obtenu en considérant une loi de probabilité uniforme sur l'intervalle de largeur 2Δ .

Exemple

- L'incertitude-type d'une mesure effectuée sur une règle graduée au mm vaut : $u(x) = \frac{\Delta}{\sqrt{3}} = \frac{0,5}{1,732} = 0,29$ mm.
- L'incertitude-type de la mesure de la tension effectuée sur le voltmètre précédent vaut : $u(x) = \frac{\Delta}{\sqrt{3}} = \frac{0,0076}{1,732} = 0,0044$ mV.
- L'incertitude-type de la mesure de la position de l'écran sur l'exemple précédent vaut : $u(x) = \frac{\Delta}{\sqrt{3}} = \frac{5}{1,732} = 2,9$ mm.

c) Que faire quand on est limité par le temps ?

Une évaluation de type A requiert de répéter chaque mesure 5 à 10 fois et il est très rare que l'on puisse la réaliser pour chaque point de mesures. Dans ce cas, on évalue l'incertitude-type par une méthode de type A pour un point de mesure « typique », et on utilise une évaluation de type B pour les autres. L'évaluation de type B est alors basée sur l'incertitude-type obtenue pour la première mesure. Pour reporter intelligemment l'incertitude, il faut se poser la question suivante :

- l'incertitude est-elle constante ? Dans ce cas, on reporte l'incertitude-type ;
- l'incertitude augmente-t-elle proportionnellement à la valeur mesurée ? Dans ce cas, on reporte l'incertitude relative.

Exemple

Les incertitudes de mesure sur un instrument gradué sont constantes.

Les incertitudes de mesure sur un multimètre sont proportionnelles à la valeur mesurée.

d) Différence entre erreur et incertitude

Maintenant que l'on a défini les notions d'erreur et d'incertitude, on peut revenir sur la différence entre ces deux termes :

- l'erreur est la différence entre la moyenne des valeurs mesurées et la « valeur vraie » ;
- l'incertitude caractérise la dispersion des valeurs mesurées et quantifie le doute que l'on a sur le résultat de mesure.

Les erreurs, si elles sont systématiques et connues, peuvent être corrigées. L'incertitude ne peut qu'être affichée.

2.3 Incertitude-type composée

Lorsque l'on réalise une mesure indirecte, on calcule une grandeur physique à partir de grandeurs mesurées. Les incertitudes de détermination des grandeurs mesurées se propagent sur la grandeur calculée et on doit déterminer l'incertitude induite sur cette dernière. Cette compétence est particulièrement utile lorsque l'on fait une évaluation de type B.

Exemple

Pour déterminer la résistance d'un dipôle, on mesure la différence de potentiel ΔV à ses bornes et le courant I qui le traverse. On calcule alors la résistance du dipôle comme le rapport $R = \frac{\Delta V}{I}$. Les incertitudes sur la mesure de ΔV et de I se propagent sur leur rapport R .

a) Méthode de propagation des erreurs

Pour pouvoir estimer l'erreur sur une grandeur g calculée à partir des valeurs d'entrée $x, y, \dots$, on fait deux hypothèses :

- les grandeurs d'entrées sont indépendantes ;
- on connaît la relation algébrique $g = f(x, y, \dots)$.

Pour simplifier les notations, on se restreint au cas où g ne dépend que de deux grandeurs d'entrée x et y qui sont déterminées avec une incertitude-type $u(x)$ et $u(y)$ et une incertitude-type relative $u_r(x) = \frac{u(x)}{x}$ et $u_r(y) = \frac{u(y)}{y}$. Le lecteur généralisera sans problème aux cas où les grandeurs d'entrées sont plus nombreuses. On se restreint également aux cas simples pour lesquels la fonction f est une somme, une différence, un produit ou un quotient. Dans les cas plus compliqués, on utilise un logiciel dédié.

b) Comparaison de l'importance des sources d'erreur

La première étape de l'évaluation d'une incertitude composée est de déterminer chaque source d'erreur, d'évaluer les incertitudes associées et de comparer les incertitudes sur g induites par chaque source d'erreur. Pour cela, on calcule numériquement à la calculatrice :

- l'incertitude sur g due à x : $u_x(g) = |f(x + u(x), y) - f(x, y)|$;
- l'incertitude sur g due à y : $u_y(g) = |f(x, y + u(y)) - f(x, y)|$.

On compare ensuite ces grandeurs.

CHAPITRE A – MESURES ET INCERTITUDES

c) Cas où une source d’erreur est dominante

Si $u_x(g) \gg u_y(g)$, la grandeur g est principalement sensible aux erreurs sur x et l’incertitude-type composée sur g vaut :

$$u_c(g) \simeq u_x(g) = |f(x + u(x), y) - f(x, y)|.$$

De manière générale, il faut simplifier le problème en négligeant toutes les sources d’erreur qui ont une influence négligeable sur la valeur de g et on est très souvent ramené au cas où une seule source d’erreur est prédominante.

d) Cas où plusieurs sources d’erreur interviennent

On note $u_c(g)$ l’incertitude-type composée de g et $u_{r,c}(g) = \frac{u_c(g)}{g}$ son incertitude relative.

On regroupe les cas où deux sources d’erreurs influent sur l’incertitude de g dans le tableau A.1 :

- dans le cas d’une somme ou d’une différence, les carrés des incertitudes s’ajoutent ;
- dans le cas d’une produit ou d’un rapport, ce sont les carrés des incertitudes relatives.

Tableau A.1 – Incertitudes composées

$g(x, y)$	$x + y$	$x - y$	$x \times y$	$\frac{x}{y}$
$u_c(g)$	$\sqrt{u(x)^2 + u(y)^2}$	$\sqrt{u(x)^2 + u(y)^2}$		
$u_{r,c}(g)$			$\sqrt{u_r(x)^2 + u_r(y)^2}$	$\sqrt{u_r(x)^2 + u_r(y)^2}$

3 Présentation d’un résultat expérimental

3.1 Notation d’un résultat

On présente toujours un résultat expérimental accompagné d’une incertitude et d’une unité, avec un nombre de chiffres significatifs cohérent et en indiquant le niveau de confiance.

On présente une mesure physique sous la forme de l’intervalle dimensionné :

$$x = \left(\bar{x} \pm k \times u(x) \right) \text{ unités}$$

où $u(x)$ représente l’incertitude-type de mesure ou l’incertitude-type composée. Le niveau de confiance est approximativement de 68% pour $k = 1$ et de 95% pour $k = 2$.

L’incertitude élargie est souvent notée $U(x) = ku(x)$. On peut noter qu’elle n’apporte pas plus d’information que l’incertitude-type et n’est donc qu’un outil de communication. On la présente également sous la forme d’une **incertitude élargie relative** $U_r(x) = \frac{k \times u(x)}{|\bar{x}|}$ que l’on exprime par un pourcentage.

Remarque

Lorsque l'on fait une évaluation de type A, on peut affiner l'analyse des incertitudes si nécessaire en prenant en compte le fait que le facteur d'élargissement dépend du nombre de mesures n effectuées et du niveau de confiance c désiré. Ce facteur, noté $t(c, n)$, est une fonction de c et de n . Il est donné par la loi de Student (voir tableau A.2).

Tableau A.2 – Table de Student

$c \backslash n$	2	4	6	8	10	20	∞
68%	1,84	1,20	1,11	1,08	1,06	1,03	1,00
95%	12,71	3,18	2,57	2,36	2,26	2,09	1,96

3.2 Chiffres significatifs et arrondis

a) Chiffres significatifs

Précision des incertitudes L'estimation de l'incertitude est elle-même incertaine. Lors d'une évaluation de type A, elle est déterminée avec une incertitude de 35% pour 5 mesures, de 25% pour 10 mesures et de 10% pour 50 mesures. Lors des travaux pratiques, seul le premier chiffre de l'incertitude a généralement un sens.

Chiffres significatifs affichés

On présente l'incertitude-type $u(x)$ avec un ou deux chiffres significatifs.
Le dernier chiffre significatif de $\bar{x}$ doit correspondre au dernier chiffre significatif de l'incertitude-type $u(x)$.

b) Arrondis

Pour respecter les règles de présentation précédentes, il est nécessaire d'arrondir les nombres obtenus. On utilise généralement l'arrondi au plus proche. Attention, arrondir un résultat numérique fixe la précision par défaut, il faut donc réfléchir soigneusement aux nombres de chiffres significatifs que l'on conserve et avoir en tête les ordres de grandeur du tableau A.3.

Tableau A.3 – Nombre de chiffres significatifs et précision relative.

nombre de chiffres significatifs	précision relative
1	$\simeq 100\%$
2	$\simeq 10\%$
3	$\simeq 1\%$

4 Validité d'un résultat expérimental

4.1 Comparaison entre une valeur mesurée et une valeur de référence

Il est fréquent que le but d'un travail expérimental fait en classe soit de retrouver une valeur de référence établie précisément par des expérimentateurs chevronnés. Dans ce cas, on valide la démarche expérimentale par une comparaison entre la valeur de référence et l'intervalle de confiance obtenu précédemment.

Lorsque l'on utilise un niveau de confiance de 68% en utilisant l'incertitude-type, la valeur de référence a une probabilité de 32% de ne pas se situer dans l'intervalle de confiance. Il ne faut donc pas trop s'en inquiéter.

Lorsque l'on utilise un niveau de confiance de 95% en utilisant l'incertitude élargie par un facteur $k = 2$, la valeur de référence a une probabilité de 5% de ne pas se situer dans l'intervalle de confiance. C'est plus inquiétant et on peut commencer à douter de la mesure effectuée ou de l'estimation des incertitudes.

On convient généralement que la valeur mesurée et la valeur de référence sont incompatibles lorsque leur écart excède 3 fois l'écart-type (la probabilité qu'un tel événement arrive fortuitement n'est que de 0,3%). Dans ce cas, on recherche et on discute des causes de cet écart qui peuvent être des erreurs systématiques, des erreurs de manipulations, une estimation des incertitudes trop optimistes, ...

4.2 Vérification d'une relation linéaire entre des données

Il arrive également souvent que le but d'un travail expérimental soit de vérifier une loi physique, c'est-à-dire de vérifier que deux grandeurs x et y sont liées par une relation $y = f(x)$.

Dans cette partie, on cherche à vérifier si des données expérimentales sont compatibles avec une relation linéaire du type $y = ax + b$ et, en cas de réponse positive, à déterminer les valeurs de a et b ainsi que leurs incertitudes. Pour cela, on mesure une série de $n \geq 5$ données $x_1, x_2, \dots, x_n$ et les valeurs correspondantes $y_1, y_2, \dots, y_n$ accompagnées de leurs incertitudes. On recueille l'ensemble de ces données dans un tableau qui permet d'en faire une représentation graphique et d'effectuer une **régression linéaire**.

Remarque

Même si la relation qui lie y à x n'est pas linéaire, on peut souvent s'y ramener.

a) Détermination des paramètres de la régression linéaire

Dans un premier temps, on réalise une représentation graphique des données. Le tableau permet d'effectuer une régression linéaire par la **méthode des moindres carrés** et de représenter la droite modèle $y = ax + b$ sur le même graphique. Le logiciel fournit généralement les valeurs de a , b et du **coefficient de corrélation** linéaire r .

b) Validité du modèle linéaire

Pour juger de la validité d'un modèle linéaire, il faut vérifier deux points :

- les données doivent être compatibles avec le modèle linéaire, ce qui implique que les points

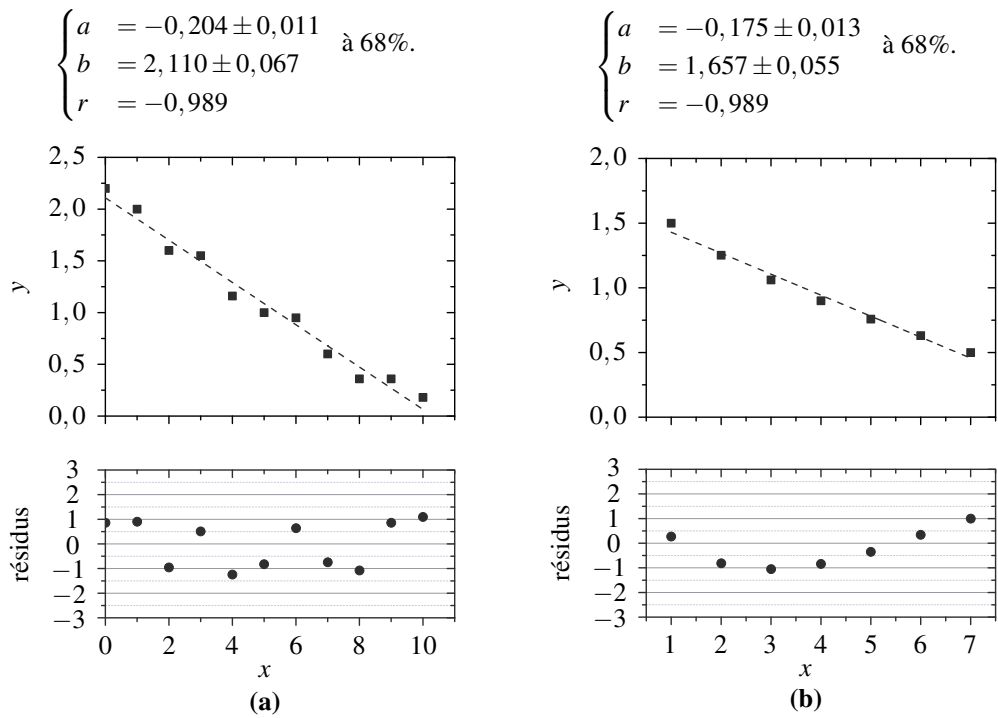


Figure A.3 – Exemple de régression linéaire avec tracé des résidus normalisés.

- de mesure soient approximativement alignés sur une droite ;
- les incertitudes-types de mesure doivent être compatibles avec le modèle linéaire.

Les données sont-elles compatibles avec le modèle linéaire ? Pour tester ce point, on vérifie que **les points sont répartis de manière aléatoire de part et d'autre de la droite de régression**. On peut le vérifier « à l'œil » directement sur la courbe ou s'aider en traçant les écarts des points de mesure y_i à la droite de régression. De nombreux tableurs et les logiciels de traitement de données proposent le tracé de ces écarts appelés résidus. Ils sont généralement normalisés pour prendre des valeurs entre -3 et 3 (voir figure A.3).

Si les résidus forment une courbe et ne sont pas aléatoirement répartis, c'est que le modèle linéaire n'est pas le meilleur modèle et qu'un paramètre physique qui a une influence visible sur les mesures expérimentales a été négligé dans le modèle (voir figure A.3 (b)). Le modèle n'est pas forcément invalidé mais on sait qu'il n'est pas complet.

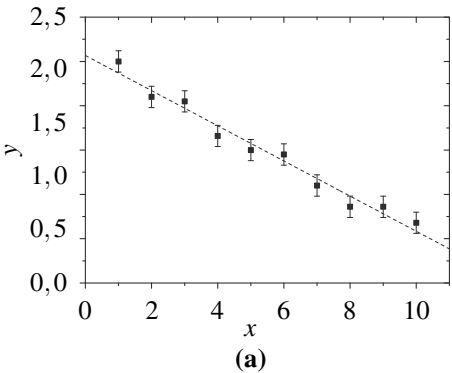
Si les points sont répartis aléatoirement autour de la droite de régression linéaire, les résidus sont répartis aléatoirement de part et d'autre de 0 et le modèle linéaire est le meilleur modèle permettant d'expliquer les données (voir figure A.3 (a)). On doit alors pousser l'analyse un peu plus loin et vérifier que les incertitudes-types le sont également.

CHAPITRE A – MESURES ET INCERTITUDES

Les incertitudes-types sont-elles compatibles avec le modèle linéaire ? Pour tester ce point, on ajoute les incertitudes-types des données sur le graphe précédent. Trois situations sont envisageables :

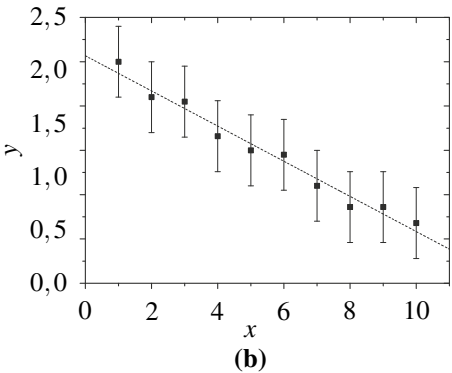
- les écarts à la droite de régression sont comparables aux incertitudes-types. Dans ce cas, le modèle linéaire est validé (voir figure A.4 (a)) et le logiciel de traitement de données donne :

$$\begin{cases} a = -0,204 \pm 0,011 \\ b = 2,110 \pm 0,067 \\ r = -0,989 \end{cases} \quad \text{à } 68\%.$$



- les écarts à la droite de régression sont petits devant les incertitudes-types. Dans ce cas, le modèle linéaire est validé mais on a probablement sur-estimé les incertitudes (voir figure A.4 (b)) et le logiciel de traitement de données donne :

$$\begin{cases} a = -0,204 \pm 0,038 \\ b = 2,11 \pm 0,23 \\ r = -0,989 \end{cases} \quad \text{à } 68\%.$$



- les écarts à la droite de régression sont grands devant les incertitudes-types. Dans ce cas, le modèle linéaire est invalidé et il y a deux possibilités : soit le modèle est réellement réfuté par les mesures, soit on a sous-estimé les incertitudes (voir figure A.4 (c)). Les logiciels de traitements de données donnent néanmoins :

$$\begin{cases} a = -0,204 \pm 0,005 \\ b = 2,110 \pm 0,028 \\ r = -0,989 \end{cases} \quad \text{à } 68\%.$$

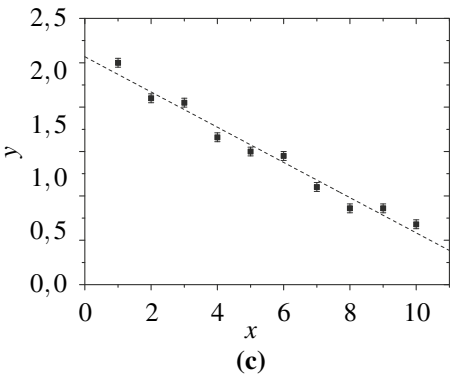


Figure A.4 – Exemples de régressions linéaires avec incertitudes.

Au final, les logiciels de traitement de données permettent d’obtenir les valeurs de a et b , ainsi que leurs incertitudes-types.

Que signifie le coefficient de corrélation r ? Le coefficient de corrélation linéaire r est un paramètre statistique qui mesure si y est lié par une relation linéaire à x . Il est égal à ± 1 lorsque y varie de manière linéaire avec x et s'en éloigne si ce n'est pas le cas. Malheureusement, ce paramètre souffre d'un grave défaut, il ne permet pas de déterminer si l'écart à 1 est dû aux incertitudes de mesures affectant x et y ou si cet écart est dû au fait que le modèle linéaire n'est pas valide. Par exemple, les coefficients de corrélation linéaire r sont approximativement les mêmes sur les figures A.3(a) et A.3(b) alors que la modélisation linéaire n'est valide que sur la figure A.3(a).

Par ailleurs, on voit sur la figure A.4 que ce coefficient donne toujours la même valeur, quelques soient les incertitudes pesant sur x et y . Il ne permet donc pas d'évaluer les incertitudes sur a et b .

Au final, en pratique, r n'est réellement utile que pour réaliser la vérification suivante : si l'on s'attend à une relation linéaire et que $|r|$ est inférieur à 0,9, il y a vraisemblablement un problème dans les mesures ou une faute de frappe dans le tableau de données. Il est alors profitable de relire et éventuellement de reprendre quelques mesures.

Conclusion En pratique, le tracé d'un graphe est le moyen le plus efficace pour juger immédiatement de la qualité des mesures réalisées et ainsi contrôler l'absence de points aberrants.

CHAPITRE A – MESURES ET INCERTITUDES

Outils mathématiques



1 Équations algébriques

1.1 Système linéaire de n équations à p inconnues

a) Cas général

On appelle système linéaire de n équations à p inconnues un jeu de n équations de la forme :

$$\begin{cases} a_{1,1}x_1 + a_{1,2}x_2 + \dots + a_{1,p}x_p = C_1 \\ a_{2,1}x_1 + a_{2,2}x_2 + \dots + a_{2,p}x_p = C_2 \\ \dots \\ a_{n,1}x_1 + a_{n,2}x_2 + \dots + a_{n,p}x_p = C_n \end{cases}$$

En physique ces équations proviennent de l'écriture de n lois physiques différentes. Les $n \times p$ nombres, $a_{i,j}$ s'expriment en fonction des données du modèle, de même que les n nombres C_i . Les p nombres x_j représentent des grandeurs physiques que l'on cherche à calculer. Suivant le cas tous ces nombres sont des réels ou des complexes (si l'on travaille avec une notation complexe comme, par exemple, en électrocinétique).

Il ne peut y avoir plus d'équations que d'inconnues, ainsi dans tous les cas : $n \leq p$. Dans les cas où $n = p$ les équations permettent de calculer toutes les inconnues. On dit que le système est fermé. Dans le cas où $n < p$, le système ne permet pas de calculer toutes les inconnues mais seulement d'exprimer n inconnues en fonction des $p - n$ autres.

b) Cas d'un système de 2 équations à 2 inconnues

Le programme fixe comme compétence exigible la résolution d'un système linéaire de la forme précédente, uniquement dans le cas où $n = p = 2$.

On peut écrire un tel système sous la forme plus simple :

$$\begin{cases} ax + by = C \\ cx + dy = D \end{cases}$$

en notant x et y les nombres inconnus et a, b, c, d, C et D les nombres connus (c'est-à-dire les nombres exprimés en fonction des paramètres du modèle). Les deux équations provenant de deux lois physiques différentes, elles sont indépendantes. On supposera dans la suite que leurs coefficients ne sont pas proportionnels, soit que $ad \neq bc$.

CHAPITRE B – OUTILS MATHÉMATIQUES

Pour trouver les expressions des inconnues, la méthode à suivre est la suivante :

- exprimer y en fonction de x et des nombres connus à partir d'une des équations ;
- remplacer y par cette expression dans l'autre équation, on arrive alors à une équation linéaire pour x , contenant seulement x et les nombres connus ;
- en déduire l'expression de x en fonction des nombres connus (donc des données du modèle) ;
- en déduire, à l'aide de l'expression de y obtenue à la première étape, y en fonction des nombres connus (soit des données du modèle).

Ainsi, en supposant $b \neq 0$:

- la première équation donne $y = \frac{C - ax}{b}$;
- donc la deuxième équation s'écrit $cx + d\frac{C - ax}{b} = D$;
- d'où l'on tire : $x = \frac{dC - bD}{ad - bc}$;
- d'où, en utilisant la première équation $y = \frac{1}{b} \left(C - a\frac{dC - bD}{ad - bc} \right) = \frac{aD - cC}{ad - bc}$.

Ces formules ne sont pas à retenir, mais il faut savoir appliquer la méthode ci-dessus pour résoudre le système.

Remarque

Il arrive que l'une des équations ne contienne que l'une des deux inconnues : par exemple, si $b = 0$, la première équation ne contient que x . Dans ce cas, cette équation nous fournit directement une expression de cette inconnue en fonction des nombres connus (donc des données) : $x = \frac{C}{a}$. En injectant cette expression dans l'autre équation, on trouve une équation pour l'autre inconnue : $c\frac{C}{a} + dy = D$, d'où l'on tire facilement l'expression de y en fonction des nombres connus (donc des données) : $y = \frac{aD - cC}{ad}$.

c) Autres cas

Dans les autres cas, les logiciels de calcul formel permettent d'obtenir la solution. Certaines calculatrices ont un tel logiciel implanté.

1.2 Équation non linéaire

Une équation non linéaire à une inconnue x est une relation de la forme :

$$f(x) = g(x)$$

vérifiée par un nombre x . En physique, x représente une grandeur que l'on souhaite calculer et les fonctions f et g font intervenir les paramètres d'un modèle que l'on étudie. Mis à part dans quelques cas particuliers (comme le cas de l'équation du second degré où $f(x) = ax^2 + bx + c$ et $g(x) = 0$) les équations non linéaires ne peuvent se résoudre formellement. On doit donc avoir recours au calcul numérique d'une solution approchée.

La méthode graphique consiste à représenter sur un même graphe les fonctions $f(x)$ et $g(x)$ sur un domaine correspondant aux valeurs possibles de x (vraisemblables physiquement). On trouve alors la ou les solutions en lisant l'abscisse du ou des points d'intersection des deux courbes. Pour affiner la détermination, on peut "faire un zoom" autour de la valeur recherchée. Cette méthode est illustrée sur la figure B.1 avec l'équation $\cos(x) = \sin(x)$ pour x compris entre 0 et π dont la solution $x = \frac{\pi}{4}$ est bien connue.

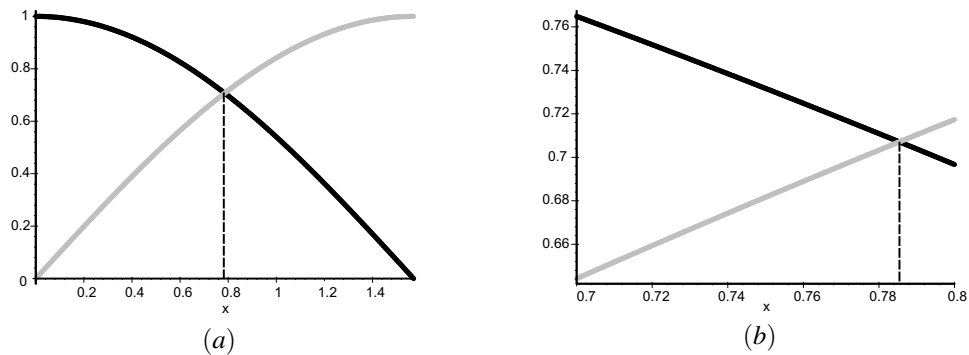


Figure B.1 – Les fonctions $f(x) = \sin(x)$ et $g(x) = \cos(x)$. Le graphe (a) permet de situer la solution de $f(x) = g(x)$ entre 0,7 et 0,8. La figure « zoomée » (b) permet de trouver l'approximation $x \simeq 0,785$, valeur voisine de la solution exacte : $x = \frac{\pi}{4}$.

Remarque

Il peut y avoir plusieurs solutions... ou aucune. L'existence de solution est d'ailleurs souvent discutée graphiquement, lorsque l'une des deux fonctions dépend d'un paramètre.

Il existe aussi des algorithmes mathématiques pour le calcul de solutions approchées. Ces algorithmes sont implantés sur la plupart des calculettes ou dans les logiciels de calcul formel.

CHAPITRE B – OUTILS MATHÉMATIQUES

2 Équations différentielles

On appelle équation différentielle une relation entre une fonction inconnue d’une variable et au moins l’une de ses dérivées. Résoudre l’équation différentielle consiste à trouver la fonction inconnue.

Le plus souvent, en physique, la variable dont dépend la fonction est le temps t . On notera dans la suite $x(t)$ la fonction inconnue.

On notera $\frac{dx}{dt}$ sa dérivée par rapport au temps, appelée dérivée première, et $\frac{d^2x}{dt^2} = \frac{d}{dt} \left(\frac{dx}{dt} \right)$ sa dérivée seconde qui est la dérivée par rapport au temps de la dérivée première. On peut ainsi définir des dérivées successives : la dérivée d’ordre n s’obtient en dérivant n fois la fonction.

Remarque

La notation utilisée en mathématique est $x'(t)$ pour la dérivée première, et $x''(t)$ pour la dérivée seconde.

On appelle ordre de l’équation différentielle l’ordre de dérivation le plus grand qui apparaît dans l’équation. Dans le cadre du programme on considérera des équations différentielles d’ordre 1 ou 2 uniquement.

2.1 Équation différentielle linéaire d’ordre 1 à coefficients constants

Dans ce paragraphe on considère une équation différentielle linéaire du **premier ordre** à coefficient constants. Cette équation différentielle peut se mettre sous la forme :

$$\frac{dx}{dt} + ax(t) = f(t), \tag{B.1}$$

où a est une constante non nulle (en physique cette constante s’exprime en fonction de paramètres connus du modèle). $f(t)$ est une fonction d’expression connue, appelée « second membre de l’équation ».

a) Solution générale de l’équation homogène : cas où $f(t) = 0$

L’équation est dite homogène si son second membre est nul, soit si $f(t) = 0$. Elle se réécrit dans ce cas :

$$\frac{dx}{dt} = -ax(t).$$

Or, on sait que la dérivée de la fonction exponentielle $g(t) = \exp(rt)$ est $g'(t) = r \exp(rt) = rg(t)$. Il apparaît donc que la fonction $x(t) = \exp(-at)$ est une solution de cette équation.

Cette solution ne peut convenir, car c’est une pure fonction mathématique, sans dimension physique (autrement dit, sans unité) et on cherche une grandeur physique $x(t)$. Mais on vérifie facilement que $x(t) = \lambda \exp(-at)$ où λ est une constante (indépendante du temps) est aussi une solution. Ceci est dû à la linéarité de l’équation. En donnant à λ la dimension physique de la grandeur x , on obtient une solution convenable. On admettra qu’il s’agit de la solution la plus générale :

L'équation différentielle linéaire homogène à coefficients constants

$$\frac{dx}{dt} + ax(t) = 0 \tag{B.2}$$

(avec a non nul) a pour solution générale :

$$x(t) = \lambda \exp(-at),$$

où λ est une constante quelconque.

Comment détermine-t-on la constante λ correspondant à la solution d'un problème donné ? On utilise la condition initiale. Dans le cas d'une équation du premier ordre, la condition initiale est une valeur donnée x_0 de $x(0)$, valeur de la fonction $x(t)$ à l'instant initial. La relation déterminant λ est ainsi :

$$x_0 = x(0) = \lambda \exp(-a \times 0) = \lambda,$$

de sorte que :

La solution de l'équation différentielle (B.2) vérifiant la condition initiale $x(0) = x_0$ est

$$x(t) = x_0 \exp(-at).$$

b) Méthode générale de résolution dans le cas où $f(t)$ n'est pas la fonction nulle

Pour une équation différentielle **linéaire** avec un second membre non nul, la recherche de la solution se fait en trois étapes :

1. On cherche la solution $x_h(t)$ de l'équation homogène associée à l'équation différentielle, c'est-à-dire celle que l'on obtient en remplaçant le second membre $f(t)$ par 0. Cette solution s'exprime en fonction de constantes qui restent indéterminées à ce stade.
2. On cherche une solution particulière $x_p(t)$ de l'équation complète (avec son second membre), de la forme mathématique la plus simple possible, sans se préoccuper des conditions initiales.
3. On cherche la solution $x(t)$ sous la forme $x(t) = x_p(t) + x_h(t)$ où $x_h(t)$. On détermine alors les constantes en imposant les conditions initiales à $x(t)$.

Cette méthode s'applique notamment à l'équation différentielle (B.1). L'étape 1 fait l'objet du paragraphe précédent et on a trouvé :

$$x_h(t) = \lambda \exp(-at),$$

où λ est une constante.

L'étape 2, fait l'objet des paragraphes suivants pour deux formes mathématiques différentes de la fonction $f(t)$.

CHAPITRE B – OUTILS MATHÉMATIQUES

Une fois une solution particulière $x_p(t)$ déterminée, on connaît toutes les solutions de (B.1), qui sont de la forme :

$$x(t) = x_p(t) + \lambda \exp(-at),$$

où λ est une constante. L'étape 3 consiste à trouver, parmi ces solutions, celle qui vérifie la condition initiale. Dans le cas d'une équation du premier ordre, la condition initiale est une valeur donnée x_0 de $x(0)$. La relation déterminant λ est ainsi :

$$x_0 = x(0) = x_p(0) + \lambda.$$

Ainsi, la solution de l'équation différentielle (B.1) vérifiant $x(0) = x_0$ est :

$$x(t) = x_p(t) + (x_0 - x_p(0)) \exp(-at).$$

c) Solution particulière avec $f(t) = C$, constante

Dans ce cas, la forme la plus simple envisageable pour la solution particulière est une fonction constante $x_p(t) = K$. Alors, $\frac{dx_p}{dt} = 0$ et l'équation (B.1) s'écrit : $0 + aK = C$ d'où l'on tire :

$$K = \frac{C}{a}.$$

Ainsi, une solution particulière de (B.1), dans le cas où $f(t) = C$, est :

$$x_p(t) = \frac{C}{a}.$$

d) Solution particulière avec $f(t) = C \cos(\omega t + \varphi)$

On considère le cas où le second membre est un signal sinusoïdal : $f(t) = C \cos(\omega t + \varphi)$, où C , ω et φ sont des constantes

On cherche une solution particulière de la forme :

$$x_p(t) = K \cos(\omega t + \psi),$$

où K et ψ sont des constantes. Pour cela, il y a deux méthodes efficaces :

- la méthode des grandeurs complexes
- la méthodes des vecteurs de Fresnel

Méthode des grandeurs complexes : Cette méthode consiste à représenter les signaux sinusoïdaux $x_p(t)$ et $f(t)$ par les signaux complexes :

$$\underline{x_p} = K \exp(i(\omega t + \psi)) \quad \text{et} \quad \underline{f} = C \exp(i(\omega t + \varphi)) \quad (\text{B.3})$$

tels que : $x_p(t) = \text{Re}(\underline{x_p})$ et $f(t) = \text{Re}(\underline{f})$. Avec les signaux complexes, l'opérateur de dérivation par rapport au temps est une simple multiplication par le nombre complexe $i\omega$, par exemple :

$$\frac{d\underline{x_p}}{dt} = i\omega \underline{x_p}.$$

L'équation différentielle (B.1) se traduit pour les signaux complexes par la relation :

$$i\omega \underline{x_p} + a \underline{x_p} = \underline{f}$$

d'où l'on tire immédiatement :

$$\underline{x_p} = \frac{\underline{f}}{i\omega + a} = \frac{C}{i\omega + a} \exp(i(\omega t + \varphi)).$$

On en déduit facilement l'amplitude K de $x_p(t)$ par :

$$K = |\underline{x_p}| = \frac{C}{\sqrt{\omega^2 + a^2}},$$

Pour obtenir une expression du signal $x_p(t)$ il faut prendre la partie réelle de $\underline{x_p}$. Or :

$$\begin{aligned} \underline{x_p} &= \frac{C}{\omega^2 + a^2} (-i\omega + a) \exp(i(\omega t + \varphi)) \\ &= \frac{C}{\omega^2 + a^2} (-i\omega + a) (\cos(\omega t + \varphi) + i \sin(\omega t + \varphi)) \\ &= \frac{C}{\omega^2 + a^2} (-i\omega \cos(\omega t + \varphi) + a \cos(\omega t + \varphi) + \omega \sin(\omega t + \varphi) + ia \sin(\omega t + \varphi)) \end{aligned}$$

Donc :

$$x_p(t) = \frac{C}{\omega^2 + a^2} (a \cos(\omega t + \varphi) + \omega \sin(\omega t + \varphi))$$

Méthode des vecteurs de Fresnel : La méthode des vecteurs de Fresnel consiste à traduire la relation :

$$\frac{dx_p}{dt} + ax_p(t) = f(t) = C \cos(\omega t + \varphi)$$

par la relation vectorielle :

$$\overrightarrow{DX_p} + a\overrightarrow{X_p} = \overrightarrow{F}$$

entre les vecteurs de Fresnel $\overrightarrow{DX_p}$, $\overrightarrow{X_p}$ et $\overrightarrow{F}$ respectivement associés aux signaux sinusoïdaux $\frac{dx_p}{dt}$, $x_p(t)$ et $f(t)$. Ces vecteurs ont les caractéristiques suivantes :

- $\overrightarrow{F}$ a une norme égale à C et fait avec l'axe OX un angle égal à φ ,
- $\overrightarrow{X_p}$ a une norme égale à K et fait avec l'axe OX un angle égal à ψ ,
- $\overrightarrow{DX_p}$ a une norme égale à ωK et fait avec l'axe OX un angle égal à $\psi + \frac{\pi}{2}$.

On représente cette relation graphiquement en procédant de la manière suivante : on trace le vecteur $a\overrightarrow{X_p}$ à partir de l'origine avec un angle ψ quelconque (on ne connaît pas encore cet angle), puis, en partant de l'extrémité de ce vecteur, le vecteur $\overrightarrow{DX_p}$ qui fait un angle $+\frac{\pi}{2}$ avec $\overrightarrow{X_p}$. Le vecteur entre l'origine et l'extrémité de $\overrightarrow{DX_p}$ est $\overrightarrow{F}$. Ceci est réalisé sur la figure B.2 dans le cas $a > 0$. On détermine K et ψ en exploitant la figure :

CHAPITRE B – OUTILS MATHÉMATIQUES

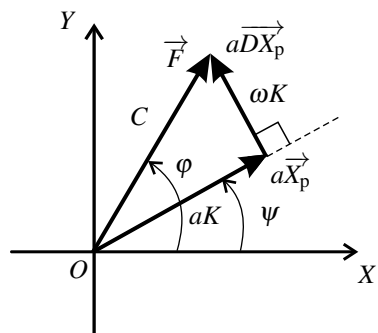


Figure B.2 – Résolution de l'équation (B.1) par les vecteurs de Fresnel, dans le cas $a > 0$. On a indiqué sur chaque vecteur sa norme.

- d'après le théorème de Pythagore :

$$C^2 = (aK)^2 + (\omega K)^2 \quad \text{d'où} \quad K = \frac{C}{\sqrt{\omega^2 + a^2}}.$$

- on voit que $\varphi - \psi$ est compris entre 0 et $\frac{\pi}{2}$ et vérifie :

$$\cos(\varphi - \psi) = \frac{aK}{C} = \frac{a}{\sqrt{\omega^2 + a^2}}, \quad \text{et} \quad \sin(\varphi - \psi) = \frac{\omega K}{C} = \frac{\omega}{\sqrt{\omega^2 + a^2}}.$$

Il vient finalement :

$$\begin{aligned} x_p(t) &= \frac{C}{\sqrt{\omega^2 + a^2}} \cos(\omega t + \varphi - (\varphi - \psi)) \\ &= \frac{C}{\sqrt{\omega^2 + a^2}} (\cos(\varphi - \psi) \cos(\omega t + \varphi) + \sin(\varphi - \psi) \sin(\omega t + \varphi)). \\ &= \frac{C}{\omega^2 + a^2} (a \cos(\omega t + \varphi) + \omega \sin(\omega t + \varphi)). \end{aligned}$$

2.2 Équation différentielle linéaire homogène d'ordre 2 à coefficients constants

Dans ce paragraphe on résout une équation différentielle linéaire et homogène du **deuxième ordre** à coefficients constant. Cette équation différentielle est de la forme :

$$\frac{d^2x}{dt^2} + a \frac{dx}{dt} + bx(t) = 0 \tag{B.4}$$

où a et b sont des constantes connues et $x(t)$ la fonction inconnue recherchée.

a) Solutions d’une équation homogène, équation caractéristique

Par analogie avec le cas de l’équation du premier ordre, on peut chercher des solutions de la forme $x(t) = \exp(rt)$ où r est une constante. D’après les propriétés de l’exponentielle :

$$\frac{dx}{dt} = r \exp(rt), \quad \frac{d^2x}{dt^2} = \frac{d}{dt}(r \exp(rt)) = r^2 \exp(rt)$$

Ainsi $x(t) = \exp(rt)$ est solution de (B.4) si et seulement si :

$$r^2 \exp(rt) + a \exp(rt) + b \exp(rt) = 0 \quad \text{soit} \quad r^2 + ar + b = 0.$$

L’ équation caractéristique associée à l’équation différentielle homogène

$$\frac{d^2x}{dt^2} + a \frac{dx}{dt} + bx(t) = 0$$

est l’équation du second degré, d’inconnue r :

$$r^2 + ar + b = 0. \tag{B.5}$$

La résolution de l’équation caractéristique dépend du signe de son discriminant :

$$\Delta = a^2 - 4b.$$

Premier cas : $\Delta > 0$. L’équation caractéristique admet deux solutions réelles :

$$r_1 = \frac{-a + \sqrt{\Delta}}{2} \quad \text{et} \quad r_2 = \frac{-a - \sqrt{\Delta}}{2}.$$

On retiendra que :

Lorsque l’équation caractéristique a un discriminant positif, elle admet deux racines réelles r_1 et r_2 , et la solution générale de l’équation différentielle homogène (B.4) est :

$$x(t) = \lambda \exp(r_1 t) + \mu \exp(r_2 t),$$

où λ et μ sont des constantes quelconques.

Deuxième cas : $\Delta = 0$. Dans ce cas l’équation caractéristique admet une seule racine (racine double) :

$$r_1 = -\frac{a}{2}.$$

On peut alors montrer que la fonction $x(t) = t \exp(r_1 t)$ vérifie aussi (B.4). En effet, on a alors :

$$\frac{dx(t)}{dt} = \exp(r_1 t) + r_1 t \exp(r_1 t) \quad \text{et} \quad \frac{d^2x(t)}{dt^2} = 2r_1 \exp(r_1 t) + r_1^2 t \exp(r_1 t),$$

CHAPITRE B – OUTILS MATHÉMATIQUES

de sorte que :

$$\frac{d^2x(t)}{dt^2} + a \frac{dx(t)}{dt} + bx(t) = (2r_1 + a) \exp(r_1 t) + (r_1^2 + ar_1 + b) t \exp(r_1 t) = 0.$$

Finalement :

Lorsque l'équation caractéristique a un discriminant nul, elle admet une racine double $r_1 = -\frac{a}{2}$ et la solution générale de l'équation différentielle homogène est :

$$x(t) = (\lambda t + \mu) \exp(r_1 t),$$

où λ et μ sont des constantes quelconques.

Troisième cas : $\Delta < 0$. Dans ce cas , l'équation caractéristique n'a pas de racines réelle mais deux racines complexes conjuguées :

$$\underline{r_1} = -\frac{a}{2} + i\omega \quad \text{et} \quad \underline{r_2} = -\frac{a}{2} - i\omega,$$

en notant : $\omega = \frac{\sqrt{-\Delta}}{2}$. L'équation différentielle (B.4) admet les fonctions complexes $\exp(\underline{r_1}t)$ et $\exp(\underline{r_2}t)$ comme solutions, mais aussi, puisqu'elle est linéaire, les fonctions :

$$\frac{1}{2} (\exp(\underline{r_1}t) + \exp(\underline{r_2}t)) = \frac{1}{2} \exp\left(-\frac{a}{2}t\right) (\exp(i\omega t) + \exp(-i\omega t)) = \exp\left(-\frac{a}{2}t\right) \cos(\omega t),$$

et :

$$\frac{1}{2i} (\exp(\underline{r_1}t) - \exp(\underline{r_2}t)) = \frac{1}{2i} \exp\left(-\frac{a}{2}t\right) (\exp(i\omega t) - \exp(-i\omega t)) = \exp\left(-\frac{a}{2}t\right) \sin(\omega t).$$

Finalement :

Lorsque l'équation caractéristique a discriminant négatif, elle admet deux solutions complexes conjuguées $-\frac{a}{2} \pm i\omega$ où $\omega = \frac{\sqrt{-\Delta}}{2}$ et la solution générale de l'équation différentielle homogène est :

$$x(t) = \exp\left(-\frac{a}{2}t\right) (\lambda \cos(\omega t) + \mu \sin(\omega t)),$$

où λ et μ sont des constantes quelconques.

b) Critère de stabilité

Une équation différentielle est dite stable si les solutions de l'équation différentielle homogène associée tendent toutes vers 0 pour t tendant vers l'infini.

Peut-on trouver un critère simple de stabilité pour l'équation (B.4) ?

On sait que la fonction $\exp(rt)$ tend vers 0 pour t tendant vers l'infini si $r < 0$ et tend vers l'infini si $r > 0$. Ainsi, dans le premier cas ($\Delta > 0$) l'équation est stable si et seulement si les deux racines r_1 et r_2 sont négatives. Il faut et il suffit donc que :

- leur somme soit négative soit : $r_1 + r_2 = -a < 0$,
- leur produit soit positif soit : $r_1 r_2 = b > 0$.

Dans le deuxième et le troisième cas, l'équation est stable si et seulement si $-\frac{a}{2} < 0$ soit

$a > 0$. Et comme : $b = \frac{a^2 - \Delta}{4} \geq \frac{a^2}{4}$, on a aussi $b > 0$.

En conclusion :

L'équation différentielle (B.4) est stable si et seulement si $a > 0$ et $b > 0$.

c) Formes canoniques de l'équation homogène

Dans le cas où l'équation différentielle (B.4) est stable, on a l'habitude de l'écrire sous une forme particulière, appelée forme canonique.

La forme canonique de l'équation différentielle (B.4), lorsqu'elle est stable est :

$$\frac{dx^2}{dt^2} + 2\xi \omega_0 \frac{dx}{dt} + \omega_0^2 x(t) = 0 \quad (\text{B.6})$$

ou

$$\frac{dx^2}{dt^2} + \frac{\omega_0}{Q} \frac{dx}{dt} + \omega_0^2 x(t) = 0. \quad (\text{B.7})$$

Solutions de l'équation $\frac{dx^2}{dt^2} + 2\xi \omega_0 \frac{dx}{dt} + \omega_0^2 x(t) = 0$ Les solutions trouvées ci-dessus sont :

$$\begin{aligned} \text{si } \xi > 1, \quad x(t) &= \lambda \exp\left(\left(-\xi + \sqrt{\xi^2 - 1}\right) \omega_0 t\right) + \mu \exp\left(\left(-\xi - \sqrt{\xi^2 - 1}\right) \omega_0 t\right), \\ \text{si } \xi = 1, \quad x(t) &= (\lambda t + \mu) \exp(-\omega_0 t), \\ \text{si } \xi < 1, \quad x(t) &= \exp(-\xi \omega_0 t) \left(\lambda \cos\left(\sqrt{1 - \xi^2} \omega_0 t\right) + \mu \sin\left(\sqrt{1 - \xi^2} \omega_0 t\right) \right), \end{aligned}$$

où λ et μ des constantes quelconques.

CHAPITRE B – OUTILS MATHÉMATIQUES

Solutions de l'équation $\frac{dx^2}{dt^2} + \frac{\omega_0}{Q} \frac{dx}{dt} + \omega_0^2 x(t) = 0$ Les solutions trouvées ci-dessus sont :

si $Q < \frac{1}{2}$,

$$x(t) = \lambda \exp \left(\left(-\frac{1}{2Q} + \sqrt{\frac{1}{4Q^2} - 1} \right) \omega_0 t \right) + \mu \exp \left(\left(-\frac{1}{2Q} - \sqrt{\frac{1}{4Q^2} - 1} \right) \omega_0 t \right),$$

si $Q = \frac{1}{2}$, $x(t) = (\lambda t + \mu) \exp(-\omega_0 t)$,

si $Q > \frac{1}{2}$, $x(t) = \exp \left(-\frac{\omega_0}{2Q} t \right) \left(\lambda \cos \left(\sqrt{1 - \frac{1}{4Q^2}} \omega_0 t \right) + \mu \sin \left(\sqrt{1 - \frac{1}{4Q^2}} \omega_0 t \right) \right),$

où λ et μ des constantes quelconques.

d) Équation différentielle de l'oscillateur harmonique

Dans le cas où $\xi = 0$ ou $Q = \infty$, l'équation différentielle s'écrit :

$$\frac{d^2x}{dt^2} + \omega_0^2 x(t) = 0, \tag{B.8}$$

et a pour solution :

$$x(t) = \lambda \cos(\omega_0 t) + \mu \sin(\omega_0 t),$$

où λ et μ sont des constantes quelconques.

2.3 Équation différentielle linéaire d'ordre 2 à coefficients constants, avec un second membre non nul

Dans ce paragraphe on va résoudre une équation différentielle de la forme :

$$\frac{d^2x}{dt^2} + a \frac{dx}{dt} + bx(t) = f(t) \tag{B.9}$$

où a et b sont des constantes connues et $x(t)$ la fonction inconnue recherchée. $f(t)$ est une fonction connue, appelée second membre de l'équation.

a) Méthode générale de résolution

La méthode est la même que pour l'équation du premier ordre. La solution la plus générale de l'équation avec second membre (B.9) est :

$$x(t) = x_p(t) + x_h(t),$$

où x_p est une solution particulière de l'équation complète (avec son second membre non nul) et $x_h(t)$ une solution quelconque de l'équation homogène associée (B.4), c'est-à-dire la même avec un second membre nul. Dans l'expression de $x_h(t)$ apparaissent deux constantes λ et μ qui peuvent avoir *a priori* n'importe quelle valeur.

La solution correspondant à un problème donné est celle qui vérifie les conditions initiales qui sont, pour une équation du deuxième ordre :

- une valeur initiale donnée de la fonction : $x(0) = x_0$,
- une valeur initiale donnée de la dérivée première de la fonction : $\left(\frac{dx}{dt}\right)_{t=0} = v_0$.

Ces conditions permettent de déterminer les constantes λ et μ pour la solution recherchée, en résolvant un système linéaire de deux équations à deux inconnues.

Par exemple, en utilisant la forme (B.6) pour l'équation homogène, le système à résoudre dans le cas où $\xi > 1$ est :

$$\begin{cases} \lambda + \mu = x_0 - x_p(0) \\ (-\xi + \sqrt{\xi^2 - 1})\lambda + (-\xi - \sqrt{\xi^2 - 1})\mu = v_0 - \left(\frac{dx_p}{dt}\right)_{t=0} \end{cases}$$

Dans les paragraphes suivants, on donne la méthode pour obtenir une solution particulière $x_p(t)$ de (B.9) pour différentes formes du second membre $f(t)$.

b) Solution particulière avec $f(t) = C$, constante

Dans ce cas, on cherche une solution particulière constante : $x_p(t) = K$. Alors, les dérivées de $x_p(t)$ sont nulles et (B.9) s'écrit : $bK = C$ d'où l'on tire : $K = \frac{C}{b}$.

Ainsi, une solution particulière de (B.9), dans le cas où $f(t) = C$, est :

$$x_p(t) = \frac{C}{b}.$$

c) Solution particulière avec $f(t) = C \exp(\gamma t)$

Dans le cas où le second membre est de la forme $f(t) = C \exp(\gamma t)$ avec C et γ constantes, on cherche une solution particulière de la forme :

$$x_p(t) = K \exp(\gamma t),$$

où K est une constante. Alors : $\frac{dx_p}{dt} = \gamma K \exp(\gamma t)$ et $\frac{d^2x_p}{dt^2} = \gamma^2 K \exp(\gamma t)$. L'équation (B.9) s'écrit :

$$\gamma^2 K \exp(\gamma t) + a\gamma K \exp(\gamma t) + bK \exp(\gamma t) = C \exp(\gamma t)$$

d'où l'on tire : $K = \frac{C}{\gamma^2 + a\gamma + b}$. Ainsi, une solution particulière de (B.1) dans ce cas est :

$$x_p(t) = \frac{C}{\gamma^2 + a\gamma + b} \exp(\gamma t).$$

CHAPITRE B – OUTILS MATHÉMATIQUES

d) **Solution particulière avec** $f(t) = C \cos(\omega t + \varphi)$

Dans le cas où le second membre est de la forme : $f(t) = C \cos(\omega t + \varphi)$ avec C , ω et φ constantes, on cherche une solution particulière de la forme :

$$x_p(t) = K \cos(\omega t + \psi),$$

où K et ψ sont des constantes. Pour cela, on peut utiliser :

- la méthode des grandeurs complexes,
- la méthodes des vecteurs de Fresnel.

Méthode des grandeurs complexes Cette méthode consiste à représenter les signaux sinusoïdaux $x_p(t)$ et $f(t)$ par les signaux complexes $\underline{x_p}$ et $\underline{f}$ (voir paragraphe correspondant dans l'étude de l'équation différentielle d'ordre 1).

L'équation différentielle (B.9) se traduit pour les signaux complexes par la relation :

$$(i\omega)^2 \underline{x_p} + a(i\omega) \underline{x_p} + b \underline{x_p} = \underline{f}$$

d'où l'on tire immédiatement :

$$\underline{x_p} = \frac{\underline{f}}{-\omega^2 + ia\omega + b} = \frac{C}{-\omega^2 + ia\omega + b} \exp(i(\omega t + \varphi)).$$

Remarque

Noter l'analogie avec l'expression de $x_p(t)$ du paragraphe précédent, pour $\gamma = i\omega$.

On en déduit facilement l'amplitude K de $x_p(t)$ par :

$$K = |\underline{x_p}| = \frac{C}{\sqrt{(b - \omega^2)^2 + (a\omega)^2}},$$

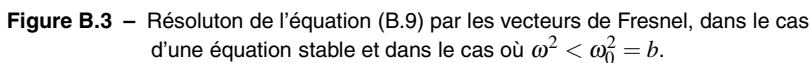
La détermination de la phase initiale ψ de $x_p(t)$ est plus délicate et ne sera pas traitée ici. Pour obtenir une expression du signal $x_p(t)$ il faut prendre la partie réelle de $\underline{x_p}$. L'expression, assez lourde, ne peut être donnée ici.

Méthode des vecteurs de Fresnel On peut traduire (B.9) par la relation vectorielle :

$$\overrightarrow{D_2 X_p} + a \overrightarrow{DX_p} + b \overrightarrow{X_p} = \overrightarrow{F}$$

où $\overrightarrow{D_2 X_p}$, $\overrightarrow{DX_p}$, $\overrightarrow{X_p}$ et $\overrightarrow{F}$ sont les vecteurs de Fresnel respectivement associés aux signaux sinusoïdaux $\frac{d^2 x_p}{dt^2}$, $\frac{dx_p}{dt}$, $x_p(t)$ et $f(t)$.

Le vecteur $\overrightarrow{D_2 X_p}$ a une norme égale à $\omega^2 K$ et il a la direction de $\overrightarrow{DX_p}$ tournée d'un angle $\frac{\pi}{2}$: il fait avec l'axe OX un angle égal à $\psi + \pi$.



Le théorème de Pythagore donne dans ce cas : $C^2 = (bK - \omega^2 K)^2 + (a\omega K)^2$, d'où :

De plus on voit que, dans le cas représenté, le retard de phase de $x_p(t)$ par rapport à $f(t)$ est compris entre 0 et $\frac{\pi}{2}$.

a) Intégration numérique

Le plus souvent, ces équations ne peuvent être résolues formellement sur le papier. Il faut faire appel à une technique d'intégration numérique approchée.

Ces techniques utilisent des méthodes « pas-à-pas » pour calculer petit à petit une valeur approchée de $x(t)$ à des instants croissants. Elles ne fonctionnent qu'avec des conditions initiales qu'il faut fournir au programme et qui sont :

- 1077

CHAPITRE B – OUTILS MATHÉMATIQUES

b) Intégrale première d'une équation de la forme $\frac{d^2x}{dt^2} = F(x)$

Une équation différentielle d'ordre 2 de la forme :

$$\frac{d^2x}{dt^2} = F(x), \quad (\text{B.10})$$

où $F(x)$ est une fonction quelconque est appelée équation de Newton. On obtient fréquemment une telle équation en appliquant le principe fondamental de la dynamique. Pour une fonction $F(x)$ non linéaire, cette équation n'est, en général, pas soluble formellement. On peut néanmoins dans tous les cas obtenir une **intégrale première**, c'est-à-dire une équation du premier ordre qui lui est équivalente.

Pour cela, on multiplie les deux membres de l'équation par la dérivée première de la fonction inconnue :

$$\frac{dx}{dt} \frac{d^2x}{dt^2} = F(x) \frac{dx}{dt}.$$

On reconnaît au membre de gauche la dérivée par rapport au temps de $\frac{1}{2} \left(\frac{dx}{dt} \right)^2$. En effet,

pour toute fonction $f(t)$ on a : $\frac{d}{dt} (f(t)^2) = 2f(t) \frac{df}{dt}$.

Au membre de droite, on a la dérivée par rapport au temps de $G(x)$ où la fonction G est une primitive de la fonction F , c'est-à-dire que $\frac{dG}{dx} = F(x)$. En effet, $G(x)$ dépend du temps

puisque x dépend du temps et : $\frac{dG(x)}{dt} = \frac{dG}{dx} \frac{dx}{dt} = F(x) \frac{dx}{dt}$.

On a ainsi :

$$\frac{d}{dt} \left(\frac{1}{2} \left(\frac{dx}{dt} \right)^2 \right) = \frac{dG(x)}{dt}$$

ce qui est équivalent à :

$$\frac{1}{2} \left(\frac{dx}{dt} \right)^2 = G(x) + K \quad (\text{B.11})$$

où K est une constante (indépendante du temps). La relation (B.11) est l'intégrale première associée à l'équation différentielle (B.10).

La valeur de K dépend des conditions initiales $x_0 = x(0)$ et $v_0 = \left(\frac{dx}{dt} \right)_{t=0}$ puisque la relation (B.11) appliquée à l'instant $t = 0$ nous donne :

$$K = \frac{1}{2} v_0^2 - G(x_0).$$



Dans ce calcul il faut bien faire attention à la variable par rapport à laquelle on dérive : il s'agit parfois de x et parfois de t .

On peut d'ailleurs écrire l'équation (B.10) de la manière suivante :

$$\frac{d}{dt} \left(\frac{dx}{dt} \right) = \frac{dG}{dx}$$

mais cela ne la rend pas intégrable car il y a d'un côté du signe égal une dérivée par rapport à t et de l'autre côté une dérivée par rapport à x .

c) Méthode de séparation des variables

La méthode de *séparation des variables* s'applique à toute équation différentielle du premier ordre de la forme :

$$\frac{dx}{dt} = F(x) \tag{B.12}$$

où $F(x)$ est une fonction de x .

La méthode consiste à séparer les variables de chaque côté du signe égal, de la manière suivante :

$$\frac{dx}{F(x)} = dt.$$

On fait ainsi apparaître des différentielles que l'on intègre, pour x entre la valeur initiale x_0 et $x(t)$, et pour t entre 0 et t :

$$\int_{x_0}^{x(t)} \frac{dx}{F(x)} = \int_0^t dt = t.$$

Si l'on sait calculer l'intégrale au membre de gauche de cette équation, on a obtenu une expression de t en fonction de $x(t)$. Cette expression peut parfois être inversée pour avoir l'expression de la fonction $x(t)$.

CHAPITRE B – OUTILS MATHÉMATIQUES

3 Fonctions

3.1 Fonctions usuelles

a) Exponentielle

Certaines fonctions sont très employées en Sciences Physiques. Il faut savoir en tracer l’allure et connaître leurs variations.

L’**exponentielle** $\exp(x)$, ou e^x , est définie sur tout $\mathbb{R}$. Elle converge très rapidement vers 0 en $-\infty$ et diverge vers $+\infty$ en $+\infty$. De plus, $\exp(0) = e^0 = 1$. L’exponentielle d’une somme ou d’un produit est à connaître :

$$\exp(x+y) = e^{x+y} = e^x \times e^y$$

et

$$\exp(xy) = e^{xy} = (e^x)^y = (e^y)^x.$$

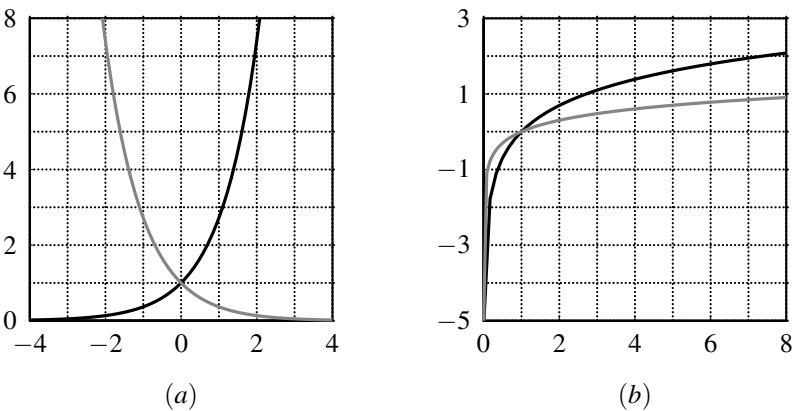


Figure B.4 – Graphes des fonctions : (a) $x \mapsto \exp(x)$ en noir, $x \mapsto \exp(-x)$ en gris, (b) $x \mapsto \ln(x)$ en noir, $x \mapsto \log(x)$ en gris.

b) Logarithme

Le **logarithme népérien** $\ln(x)$ est la fonction réciproque de l’exponentielle :

$$\exp(\ln(x)) = x \quad \text{et} \quad \ln(\exp(x)) = x.$$

Symboliquement, $\exp \circ \ln = \ln \circ \exp = \text{identité}$.

Le logarithme est défini sur $\mathbb{R}_+^* =]0, +\infty[$. Il diverge vers $-\infty$ en 0 et vers $+\infty$ en $+\infty$. De plus, $\ln(1) = 0$. La logarithme d’un produit doit être connu :

$$\ln(xy) = \ln(x) + \ln(y)$$

d’où

$$\ln\left(\frac{x}{y}\right) = \ln(x) - \ln(y)$$

et

$$\ln(x^\alpha) = \alpha \ln(x).$$

Le **logarithme décimal** $\log(x)$ est la fonction réciproque de la fonction $x \mapsto 10^x$:

$$10^{\log(x)} = \log(10^x) = x.$$

Le lien entre les logarithmes népérien et décimal est :

$$\log(x) = \frac{\ln(x)}{\ln(10)} \simeq \frac{\ln(x)}{2,3} \simeq 0,43 \ln(x) \quad \text{et} \quad \ln(x) = \frac{\log(x)}{\log(e)} = \frac{\log(x)}{0,43} \simeq 2,3 \log(x).$$

On retrouve :

$$\log(xy) = \log(x) + \log(y), \quad \log\left(\frac{x}{y}\right) = \log(x) - \log(y), \quad \log(x^\alpha) = \alpha \log(x).$$

c) Fonction puissance

La fonction puissance $x \mapsto x^a$ est définie sur $\mathbb{R}$ si $a > 0$ et sur $\mathbb{R}^*$ si $a < 0$. Elle n'est d'aucune utilité si $a = 0$:

$$x^0 = 1.$$

On retiendra :

$x^{a+b} = x^a \times x^b$

et

$x^{a-b} = \frac{x^a}{x^b}.$

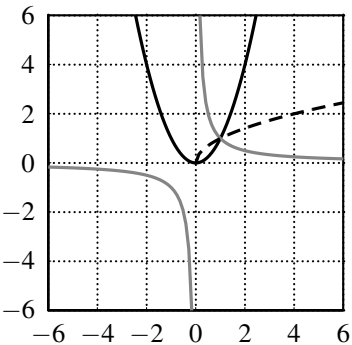


Figure B.5 – Graphes de quelques fonctions puissance : $x \mapsto x^2$ (en noir), $x \mapsto 1/x$ (en gris) et $x \mapsto \sqrt{x}$ (en pointillés).

d) Fonctions trigonométriques

Les fonctions trigonométriques sont approfondies dans la suite. Le **cosinus** et le **sinus** sont définis sur $\mathbb{R}$, leur valeur est comprise dans l'intervalle $[-1, 1]$. La **tangente** est définie par :

$$\tan(x) = \frac{\sin(x)}{\cos(x)}.$$

Elle n'est donc pas définie lorsque le cosinus s'annule, c'est à dire en $\frac{\pi}{2} + n\pi, n \in \mathbb{Z}$.

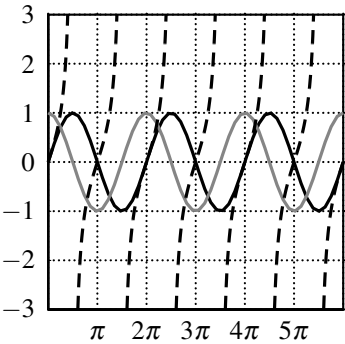


Figure B.6 – Graphes des fonctions sinus (en noir), cosinus (en gris), tangente (en pointillés).

Leurs fonctions réciproques sont notées $\arccos(x)$, $\arcsin(x)$ et $\arctan(x)$. Par exemple :

$$\sin(x) = 0,3 \Leftrightarrow x = \arcsin(0,3).$$

Les fonctions trigonométriques s'expriment en fonction d'exponentielles complexes. Avec le

CHAPITRE B – OUTILS MATHÉMATIQUES

complexe j tel que $j^2 = -1$:

$$\cos(x) = \frac{\exp(jx) + \exp(-jx)}{2} \quad \text{et} \quad \sin(x) = \frac{\exp(jx) - \exp(-jx)}{2j}.$$

e) Fonctions trigonométriques hyperboliques

Ces relations peuvent être étendues afin de définir le **cosinus hyperbolique**, noté $\cosh(x)$ ou $\text{ch}(x)$, et le **sinus hyperbolique**, noté $\sinh(x)$ ou $\text{sh}(x)$:

$$\cosh(x) = \frac{\exp(x) + \exp(-x)}{2}$$

et

$$\sinh(x) = \frac{\exp(x) - \exp(-x)}{2}.$$

Le cosinus hyperbolique, comme le sinus hyperbolique, sont définies sur $\mathbb{R}$ et divergent en $-\infty$ et $+\infty$.

Les fonctions réciproques sont nommées **argument cosinus hyperbolique** et **argument sinus hyperbolique**, et notées $\text{argch}(x)$, $\text{argsh}(x)$:

$$\text{ch}(x) = 4 \Leftrightarrow x = \text{argch}(4).$$

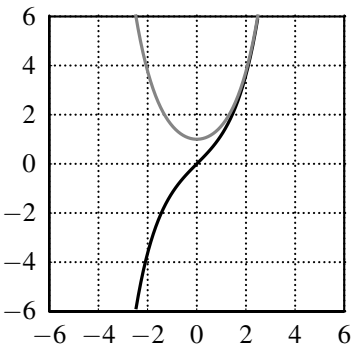


Figure B.7 – Graphes des fonctions sinus hyperbolique (en noir), cosinus hyperbolique (en gris).

3.2 Dérivée

a) Définition

La dérivée d’une fonction f est notée f' et est définie par : $f'(x) = \lim_{\delta x \rightarrow 0} \frac{f(x + \delta x) - f(x)}{\delta x}$.

On la note usuellement :

$$f' = \frac{df}{dx} \begin{matrix} \longleftarrow \text{dérivée de } f \\ \longleftarrow \text{par rapport à la variable } x \end{matrix}$$

Par exemple, la position x d’un mobile, en fonction du temps t , est une fonction $x(t)$, dont la dérivée est notée $\frac{dx}{dt}$. On rencontre aussi la notation $\dot{x}$ pour la dérivée par rapport au temps.

b) Homogénéité

La définition de la dérivée montre clairement que sa dimension est celle de f divisée par celle de x : $\left[\frac{df}{dx} \right] = \frac{[f]}{[x]}$.

Par exemple, pour la position en fonction du temps, $x(t)$, la dérivée $\frac{dx}{dt}$ s’exprime en m.s^{-1} .

c) Dérivées usuelles

$f(x)$	$f'(x)$	$f(x)$	$f'(x)$
$\cos(x)$	$-\sin(x)$	$\exp(ax)$	$a \exp(ax)$
$\sin(x)$	$\cos(x)$	$\ln(x)$	$1/x$
$\cosh(x)$	$\sinh(x)$	x^a	ax^{a-1}
$\sinh(x)$	$\cosh(x)$	constante	0

d) Dérivée d'un produit

La dérivée du produit de deux fonctions de x est : $\frac{d}{dx}(fg) = \frac{df}{dx}g + f\frac{dg}{dx}$.

e) Fonction composée

Une fonction composée s'écrit sous la forme $f \circ g$, par exemple $\ln(1 + 3x^2)$, où on applique d'abord la fonction $g : x \mapsto 1 + 3x^2$ puis le logarithme népérien f . La formule de dérivation d'une fonction composée est : $(f \circ g)' = g' \times f' \circ g$.

Avec l'exemple précédent :

$$(\ln(1 + 3x^2))' = \frac{d}{dx}(\ln(1 + 3x^2)) = \frac{d}{dx}(1 + 3x^2) \times \frac{1}{1 + 3x^2} = \frac{6x}{1 + 3x^2}.$$

Trois exemples usuels en physique :

$$\begin{aligned} \frac{d}{dt}(\sin(\theta(t))) &= \frac{d\theta}{dt} \cos(\theta(t)) \\ \frac{d}{dx}(f^a(x)) &= af^{a-1}(x) \frac{df}{dx} \quad (a \text{ positif ou négatif}) \\ \frac{d \exp(f(t))}{dt} &= f'(t) \exp(f(t)). \end{aligned}$$

f) Dérivée seconde

La dérivée seconde f'' d'une fonction f est la dérivée de la dérivée : $f'' = (f')'$. Ainsi :

$$f'' = \frac{d}{dx} \left(\frac{df}{dx} \right) = \left(\frac{d}{dx} \right)^2 f = \frac{d^2 f}{dx^2}.$$

L'opérateur dérivée par rapport à x , $\frac{d}{dx}$, apparaît donc deux fois ; il est symboliquement écrit au carré.



Attention à la place des carrés dans la notation $\frac{d^2 f}{dx^2}$.

CHAPITRE B – OUTILS MATHÉMATIQUES

Quelle est la dimension de la dérivée seconde ? La notation $\frac{d^2f}{dx^2}$ est claire ; sa dimension est celle de f divisée par celle de x au carré : $\left[\frac{d^2f}{dx^2}\right] = \frac{[f]}{[x]^2}$.

Par exemple, la dérivée seconde $\frac{d^2x}{dt^2}$ de la position x en fonction du temps t , est en $m.s^{-2}$. On rencontre aussi la dérivée $\ddot{x}$ pour la dérivée seconde par rapport au temps.

3.3 Développements limités

a) Série de Taylor

Pour une fonction de classe C^∞ sur $\mathbb{R}$, on montre la **formule de Taylor**, qui relie la valeur de la fonction en x à celle en a :

$$f(x) = f(a) + (x-a)f'(a) + \frac{(x-a)^2}{2!}f''(a) + \frac{(x-a)^3}{3!}f'''(a) + \dots + \frac{(x-a)^n}{n!}f^{(n)}(a) + \dots$$

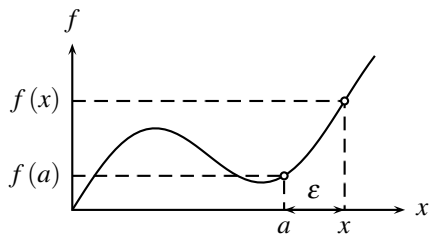


Figure B.8 – Exemple de fonction étudiée.

En posant $\varepsilon = x - a$, elle se met aussi sous la forme, très utile quand on évalue comment la fonction f varie au voisinage de a , dans le cas où ε est très inférieur à a :

$$f(a + \varepsilon) = f(a) + \varepsilon f'(a) + \frac{\varepsilon^2}{2}f''(a) + \frac{\varepsilon^3}{3!}f'''(a) + \dots + \frac{\varepsilon^n}{n!}f^{(n)}(a) + \dots$$

Attendu que ε est très inférieur à a , les termes en ε sont d'autant plus faibles qu'ils figurent à une puissance élevée. En ne gardant que les termes prépondérants d'ordre zéro, un et deux :

$$f(a + \varepsilon) = f(a) + \varepsilon f'(a) + \frac{\varepsilon^2}{2}f''(a) + o(\varepsilon^2),$$

où la notation $o(\varepsilon^2)$ symbolise des termes négligeables devant le terme en ε^2 .

b) Approximation locale d'une fonction

Lorsqu'on utilise une somme tronquée de Taylor, on obtient une approximation locale d'une fonction, en un point, d'autant plus précise que la somme contient de nombreux termes.

Par exemple, pour la fonction f , représentée sur la figure B.8, on trace les sommes partielles d'ordre un et deux :

$$f_1(x) = f(a) + (x-a)f'(a)$$
$$f_2(x) = f(a) + (x-a)f'(a) + \frac{(x-a)^2}{2!}f''(a)$$

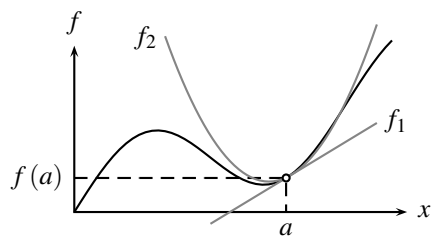


Figure B.9 – Approximations d'une fonction en un point.

On reconnaît en f_1 la tangente à f en a . Quant à f_2 , elle représente la parabole qui épouse le mieux la courbe en a .

c) Développements limités

On déduit de la formule de Taylor les **développements limités** suivants à connaître :

au premier ordre :	au deuxième ordre :
$(1 \pm x)^\alpha = 1 \pm \alpha x + o(x)$	$\cos(x) = 1 - \frac{x^2}{2} + o(x^2)$
$e^x = 1 + x + o(x)$	$\sin(x) = x + o(x^2)$
$\ln(1+x) = x + o(x)$	$\tan(x) = x + o(x^2)$

3.4 Primitive et intégrale

a) Primitive

Soit une fonction f . Une fonction F qui a pour dérivée la fonction f , est appelée **primitive** de f : $F' = f$. Par exemple, une primitive de $f(x) = x$ est $F(x) = \frac{x^2}{2}$.

Attendu que la dérivée d'une constante C est nulle :

$$(F(x) + C)' = F'(x) = f(x),$$

toute fonction $F(x) + C$ est aussi une primitive de $f(x)$.

b) Intégrale définie

Soit une fonction f admettant une primitive F . L'intégrale définie de f , entre la valeur a et la valeur b , est :

$$\int_a^b f(x) \, dx = [F(x)]_a^b = F(b) - F(a).$$

CHAPITRE B – OUTILS MATHÉMATIQUES

c) Intégrale définie et aire

On montre que l'intégrale $\int_a^b f(x) dx$ représente l'aire comprise entre la courbe qui représente f et l'axe des abscisses.

Cette aire est algébrique, c'est à dire qu'elle peut être positive comme négative. Dans l'exemple suivant, l'aire comprise entre a et 0 est négative, car en-dessous de l'axe des abscisses, et celle entre 0 et b est positive, car au-dessus.

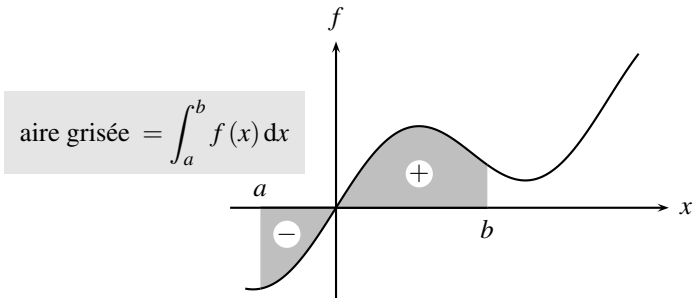


Figure B.10 – Intégrale définie.

d) Intégrale définie comme limite d'une somme

Découpons l'aire comprise entre la courbe qui représente f et l'axe des abscisses, en bandes verticales de même largeur. Chaque bande est comprise entre les abscisses x_k et x_{k+1} et a comme hauteur $f(x_k)$ ou $f(x_{k+1})$ suivant la bande. Cette opération peut être menée avec les rectangles inscrits, de surface totale S_1 , ou exinscrits, de surface totale S_2 , comme le montre la figure B.11.

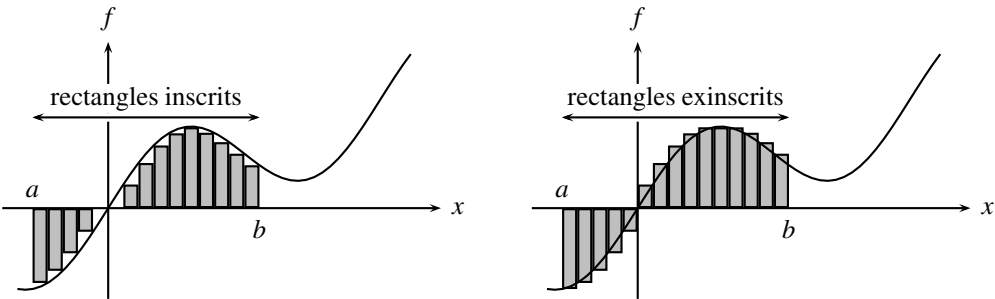


Figure B.11 – Intégrale définie et somme de rectangles.

L'aire sous la courbe qui représente f est alors comprise entre S_1 et S_2 : $S_1 < \int_a^b f(x) dx < S_2$.

On observe intuitivement que si la base des rectangles tend vers zéro, c'est à dire que la distance $x_{k+1} - x_k$ tend vers zéro, les deux sommes S_1 et S_2 convergent vers l'intégrale définie de f entre a et b .

On peut alors interpréter l'intégrale définie, comme une somme de contributions infinitésimales, qui correspondent chacune à l'aire d'un rectangle de base dx , infiniment faible, et de hauteur $f(x)$.

e) Valeur moyenne

La valeur moyenne d'une fonction s , de période T est : $\langle s(t) \rangle = \frac{1}{T} \int_{t_0}^{t_0+T} s(t) dt$,
où t_0 est un instant quelconque, très souvent pris à nul afin de simplifier les calculs.

Attendu que pour une fonction harmonique $\omega = \frac{2\pi}{T}$, soit $\omega T = 2\pi$, les valeurs moyennes d'un sinus ou d'un cosinus sont nulles :

$$\begin{aligned} \langle \sin(\omega t) \rangle &= \frac{1}{T} \int_0^T \sin(\omega t) dt = \frac{1}{T} \left[-\frac{\cos(\omega t)}{\omega} \right]_0^T = \frac{1}{T} \frac{1 - \cos(\omega T) + \cos(0)}{\omega} \\ &= \frac{1 - \cos(2\pi) + \cos(0)}{T \omega} = 0, \end{aligned}$$

$$\begin{aligned} \langle \cos(\omega t) \rangle &= \frac{1}{T} \int_0^T \cos(\omega t) dt = \frac{1}{T} \left[\frac{\sin(\omega t)}{\omega} \right]_0^T = \frac{1}{T} \frac{\sin(\omega T) - \sin(0)}{\omega} \\ &= \frac{1 \sin(2\pi) - \sin(0)}{T \omega} = 0. \end{aligned}$$

Ces résultats sont immédiats si on interprète l'intégrale comme l'aire sous la courbe. Par exemple pour le sinus ou le sinus au carré :

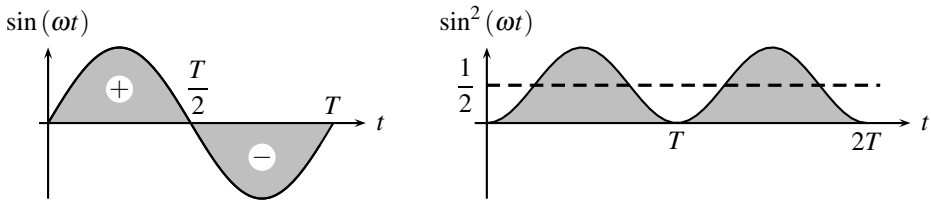


Figure B.12 – Valeur moyenne nulle d'un sinus, $\frac{1}{2}$ d'un sinus carré.

La valeur moyenne d'un $\sin^2$ ou d'un $\cos^2$ vaut $\frac{1}{2}$:

$$\begin{aligned} \langle \sin^2(\omega t) \rangle &= \frac{1}{T} \int_0^T \sin^2(\omega t) dt = \frac{1}{T} \int_0^T \frac{1 - \cos(2\omega t)}{2} dt \\ &= \frac{1}{2T} \left[t - \frac{\sin(2\omega t)}{2\omega} \right]_0^T = \frac{T - 0}{2T} - \frac{\sin(2\omega T) - \sin(0)}{4T\omega} = \frac{1}{2}. \end{aligned}$$

Idem pour celle d'un $\cos^2$ car $\sin^2(\omega t) = \frac{1 + \cos(2\omega t)}{2}$.

Le produit $\sin \times \cos$ a une valeur moyenne nulle :

$$\langle \sin(\omega t) \cos(\omega t) \rangle = \frac{1}{T} \int_0^T \sin(\omega t) \cos(\omega t) dt = \frac{1}{T} \int_0^T \frac{\sin(2\omega t)}{2} dt = 0.$$

CHAPITRE B – OUTILS MATHÉMATIQUES

Le cas d’un produit de cosinus, dont l’un est déphasé par rapport à l’autre, est à savoir traiter :

$$\begin{aligned}\langle \cos(\omega t - \varphi) \cos(\omega t) \rangle &= \langle (\cos(\omega t) \cos \varphi + \sin(\omega t) \sin \varphi) \cos(\omega t) \rangle \\ &= \cos \varphi \langle \cos^2(\omega t) \rangle + \sin \varphi \langle \cos(\omega t) \sin(\omega t) \rangle \\ &= \cos \varphi \times \frac{1}{2} + \sin \varphi \times 0 = \frac{\cos \varphi}{2}.\end{aligned}$$

Récapitulatif à connaître :

fonction	valeur moyenne
$\sin(\omega t)$	0
$\cos(\omega t)$	0
$\sin^2(\omega t)$	$\frac{1}{2}$
$\cos^2(\omega t)$	$\frac{1}{2}$
$\sin(\omega t) \cos(\omega t)$	0

3.5 Représentation graphique d’une fonction

a) Dérivée et extremum local

Une fonction f possède un extremum local, minimum ou maximum, en un point où sa dérivée est nulle.

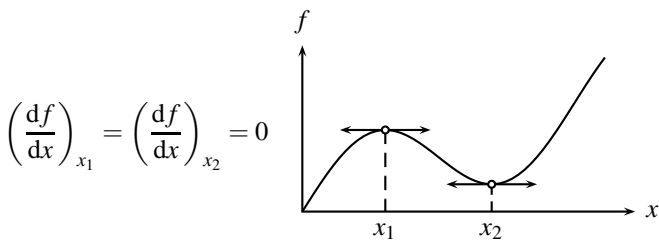


Figure B.13 – Extrema locaux.

b) Échelles logarithmiques

Soit une fonction f qui dépend de la variable x . On peut tracer f en fonction de x . En **échelles logarithmiques**, on trace $\log(f)$ en fonction de $\log(x)$.

Le quadrillage en échelles logarithmiques laisse apparaître les puissances de 10 sur chaque axe. De plus, les subdivisions qui marquent les entiers successifs, 1, 2, 3... 8, et 9, ne sont pas séparées de la même distance, car c’est le log de cet entier qui intervient :

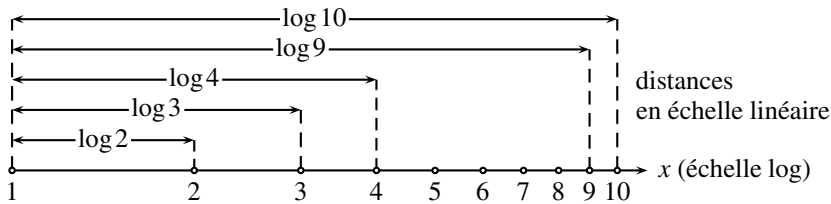


Figure B.14 – Graduations logarithmiques

L’allure d’une fonction est profondément modifiée en échelles logarithmiques, par exemple pour la fonction puissance $f(x) = \alpha x^\beta$, dont le logarithme décimal est :

$$\log(f(x)) = \log(\alpha x^\beta) = \log(\alpha) + \beta \log(x).$$

La représentation graphique de $\log(f)$ en fonction de $\log(x)$ est donc une droite.

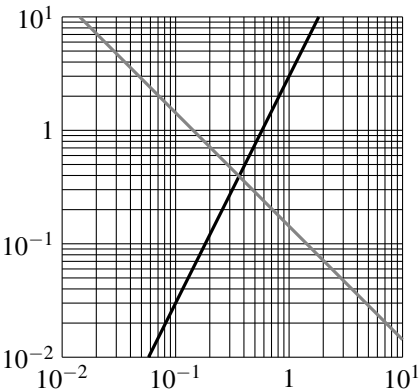


Figure B.15 – Graphe des fonctions $f(x) = 3x^2$ (en noir) et $g(x) = \frac{1}{7x} \simeq 0,14x^{-1}$ (en gris).

Une loi de puissance $x \mapsto \alpha x^\beta$ est représenté par une droite de pente β en échelles logarithmiques.

3.6 Développement en série de Fourier

a) Théorème de Dirichlet

Toute $f(t)$ une fonction de période T , continue sauf en un nombre fini de points admet un développement en série de Fourier, appelé par son acronyme DSF, tel que :

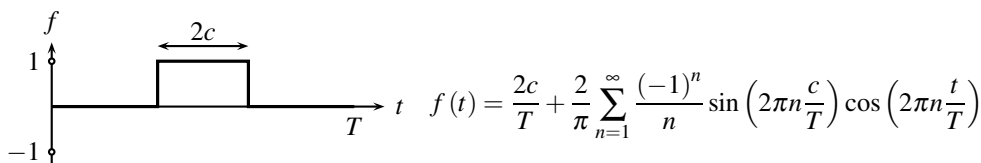
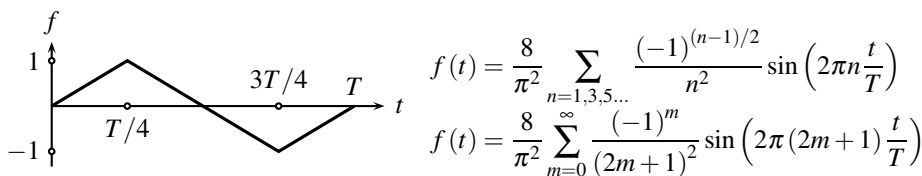
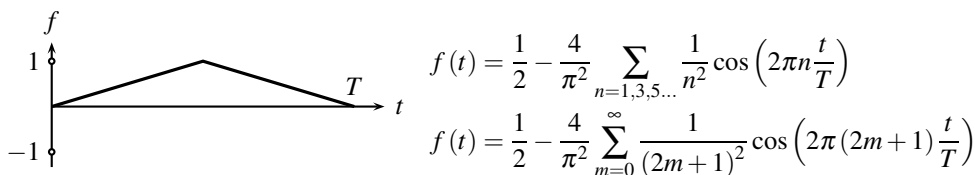
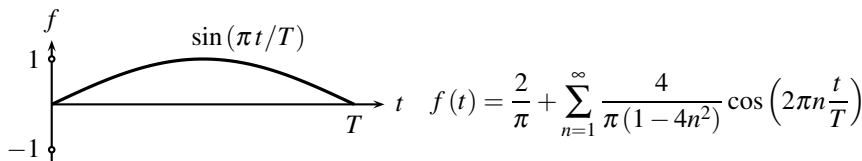
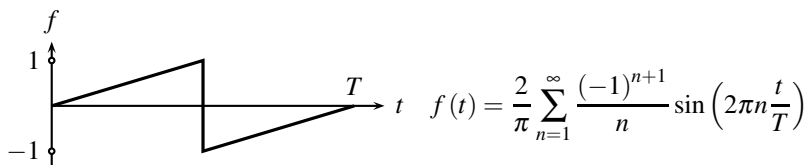
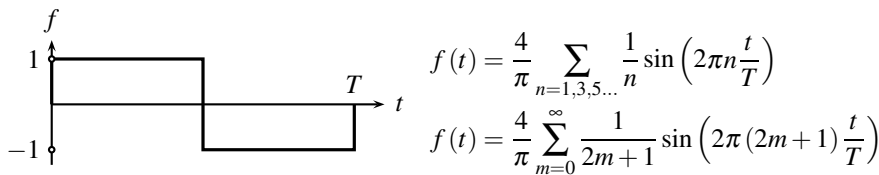
$$f(t) = A_0 + \sum_{n=0}^{\infty} A_n \cos\left(2\pi n \frac{t}{T} + \varphi_n\right),$$

en tout point où f est continue. On rencontre aussi dans la littérature la forme développée suivante :

$$f(t) = A_0 + \sum_{n=0}^{\infty} a_n \cos\left(2\pi n \frac{t}{T}\right) + b_n \sin\left(2\pi n \frac{t}{T}\right).$$

CHAPITRE B – OUTILS MATHÉMATIQUES

b) Exemples



4 Géométrie

4.1 Projection d'un vecteur, produit scalaire

a) Composantes d'un vecteur

Les systèmes de repérage d'un point de l'espace ont été étudiés dans la partie *Mécanique 1*. Dans ce qui suit on utilisera la base orthonormée $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ associée aux coordonnées cartésiennes. Mais les résultats se généralisent pour toute autre base orthonormée, par exemple la base $(\vec{u}_r, \vec{u}_\theta, \vec{u}_z)$ associée en un point M aux coordonnées cylindriques, ou la base $(\vec{u}_r, \vec{u}_\theta, \vec{u}_\varphi)$ associée en un point M aux coordonnées sphériques. Il suffira simplement de remplacer dans les formules les indices x, y , et z , par les indices r, θ et z ou bien les indices r, θ et φ .

Tout vecteur $\vec{V}$ peut s'écrire sous la forme :

$$\vec{V} = V_x \vec{u}_x + V_y \vec{u}_y + V_z \vec{u}_z,$$

où V_x, V_y et V_z sont des scalaires appelés **composantes du vecteur** sur la base $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$. On écrit de manière résumée :

$$\vec{V} = \begin{pmatrix} V_x \\ V_y \\ V_z \end{pmatrix}.$$

Attention, l'usage de cette écriture suppose que l'on a précisé clairement la base orthonormée utilisée.

b) Produit scalaire de deux vecteurs

Le **produit scalaire** des vecteurs $\vec{V}$ et $\vec{W}$ est un scalaire (c'est-à-dire un nombre) noté $\vec{V} \cdot \vec{W}$ et défini par :

$$\vec{V} \cdot \vec{W} = \begin{pmatrix} V_x \\ V_y \\ V_z \end{pmatrix} \cdot \begin{pmatrix} W_x \\ W_y \\ W_z \end{pmatrix} = V_x W_x + V_y W_y + V_z W_z. \quad (\text{B.13})$$



Dans la pratique, il faut faire attention à exprimer les vecteurs $\vec{V}$ et $\vec{W}$ sur la même base.

Le produit scalaire est symétrique :

$$\vec{V} \cdot \vec{W} = \vec{W} \cdot \vec{V},$$

et il est bilinéaire :

$$\vec{V} \cdot (\vec{W} + \vec{W}') = \vec{V} \cdot \vec{W} + \vec{V} \cdot \vec{W}' \quad \text{et} \quad \vec{V} \cdot (\lambda \vec{W}) = \lambda (\vec{V} \cdot \vec{W}),$$

$$(\vec{V} + \vec{V}') \cdot \vec{W} = \vec{V} \cdot \vec{W} + \vec{V}' \cdot \vec{W} \quad \text{et} \quad (\lambda \vec{V}) \cdot \vec{W} = \lambda (\vec{V} \cdot \vec{W}).$$

CHAPITRE B – OUTILS MATHÉMATIQUES

c) Carré scalaire et norme d'un vecteur

Le carré scalaire d'un vecteur $\vec{V}$ est : $\vec{V}^2 = \vec{V} \cdot \vec{V} = V_x^2 + V_y^2 + V_z^2$. Il est souvent noté V^2 .

La norme du vecteur $\vec{V}$ est la racine carrée de son carré scalaire :

$$\|\vec{V}\| = \sqrt{\vec{V}^2} = \sqrt{V_x^2 + V_y^2 + V_z^2}. \tag{B.14}$$

C'est un scalaire positif.

De la bilinéarité du produit scalaire découlent les formules du carré scalaire d'une somme :

$$(\vec{V} + \vec{W})^2 = \vec{V}^2 + \vec{W}^2 + 2\vec{V} \cdot \vec{W} = \|\vec{V}\|^2 + \|\vec{W}\|^2 + 2\vec{V} \cdot \vec{W},$$

et d'une différence :

$$(\vec{V} - \vec{W})^2 = \vec{V}^2 + \vec{W}^2 - 2\vec{V} \cdot \vec{W} = \|\vec{V}\|^2 + \|\vec{W}\|^2 - 2\vec{V} \cdot \vec{W}.$$

d) Projection d'un vecteur

Les produits scalaires des vecteurs de la base orthonormée sont :

$$\vec{u}_x \cdot \vec{u}_x = \vec{u}_y \cdot \vec{u}_y = \vec{u}_z \cdot \vec{u}_z = 1 \quad \text{et} \quad \vec{u}_x \cdot \vec{u}_y = \vec{u}_y \cdot \vec{u}_z = \vec{u}_z \cdot \vec{u}_x = 0.$$

Pour tout vecteur, on a donc :

$$\vec{V} \cdot \vec{u}_x = (V_x \vec{u}_x + V_y \vec{u}_y + V_z \vec{u}_z) \cdot \vec{u}_x = V_x \vec{u}_x \cdot \vec{u}_x + V_y \vec{u}_y \cdot \vec{u}_x + V_z \vec{u}_z \cdot \vec{u}_x = V_x.$$

Ceci est général : la composante d'un vecteur $\vec{V}$ sur un vecteur unitaire $\vec{u}_\alpha$ (α est un indice parmi x, y, z, r, θ ou φ) est :

$$V_\alpha = \vec{V} \cdot \vec{u}_\alpha. \tag{B.15}$$

e) Interprétation géométrique

Deux vecteurs $\vec{V}$ et $\vec{W}$ sont orthogonaux si et seulement si : $\vec{V} \cdot \vec{W} = 0$.

De manière plus générale, on a :

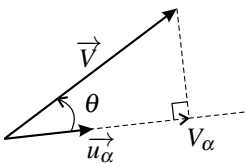
$$\vec{V} \cdot \vec{W} = \|\vec{V}\| \times \|\vec{W}\| \times \cos(\widehat{\vec{V}, \vec{W}}), \tag{B.16}$$

où $\widehat{(\vec{V}, \vec{W})}$ est l'angle entre les vecteurs $\vec{V}$ et $\vec{W}$. La formule est applicable quelle que soit l'orientation de l'angle, puisque le cosinus est une fonction paire.

Dans le cas où $\vec{W} = \vec{u}_\alpha$, vecteur unitaire d'une base, cette relation s'écrit :

$$V_\alpha = \|\vec{V}\| \cos \theta,$$

où θ est l'angle entre le vecteur unitaire $\vec{u}_\alpha$ et le vecteur $\vec{V}$ que l'on projette sur $\vec{u}_\alpha$ (voir figure).



Exemple

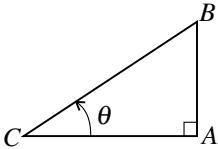
ABC étant un triangle rectangle en A quelconque (voir figure), on peut écrire, en utilisant la relation de Chasles et la bilinéarité du produit scalaire :

$$\vec{CA} \cdot \vec{CB} = \vec{CA} \cdot (\vec{CA} + \vec{AB}) = \vec{CA} \cdot \vec{CA} + \vec{CA} \cdot \vec{AB} = CA^2 + 0 = CA^2.$$

On a aussi, d'après la formule ci-dessus :

$$\vec{CA} \cdot \vec{CB} = \|\vec{CA}\| \times \|\vec{CB}\| \times \cos(\widehat{\vec{CA}, \vec{CB}}) = CA \times CB \times \cos \theta.$$

La comparaison des deux formules donne : $\cos \theta = \frac{CA}{CB}$, relation bien connue.



4.2 Produit vectoriel

a) Orientation de l'espace

L'espace physique est habituellement orienté avec la **règle de la main droite** : une base orthonormée $(\vec{u}_\alpha, \vec{u}_\beta, \vec{u}_\gamma)$ est **directe** si l'on peut placer la main droite de telle manière que le pouce, l'index et le majeur soient respectivement dans la direction et le sens des trois vecteurs $\vec{u}_\alpha, \vec{u}_\beta$ et $\vec{u}_\gamma$ (voir figure B.16). Dans le cas contraire elle est dite **indirecte**.

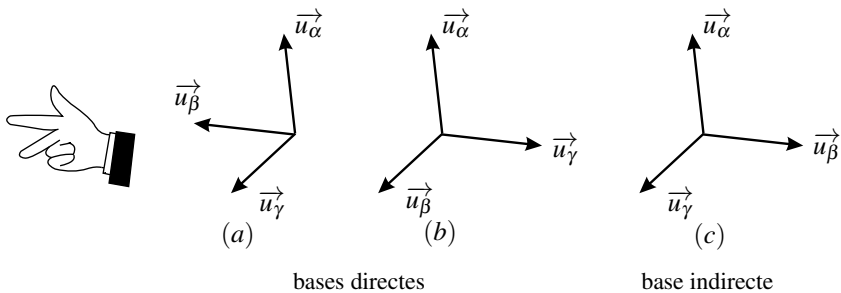


Figure B.16 – Convention d'orientation de l'espace.

Les bases (a) et (b) sur la figure ci-dessus sont directes, la base (c) est indirecte. En inversant le sens d'un des vecteurs, on transforme une base directe en une base indirecte (par exemple on transforme la base (a) en (c) ci-dessus en inversant le sens de $\vec{u}_\beta$). En permutant deux vecteurs, on transforme une base directe en une base indirecte (par exemple on transforme la base (b) en (c) en intervertissant les vecteurs $\vec{u}_\beta$ et $\vec{u}_\gamma$).



Les formules données dans la suite de ce paragraphe ne sont valables que dans le cas où la base $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ est directe.

CHAPITRE B – OUTILS MATHÉMATIQUES

b) Produit vectoriel de deux vecteurs

Le **produit vectoriel** des vecteurs $\vec{V}$ et $\vec{W}$ est un vecteur noté $\vec{V} \wedge \vec{W}$ et défini par :

$$\vec{V} \wedge \vec{W} = \begin{pmatrix} V_x \\ V_y \\ V_z \end{pmatrix} \wedge \begin{pmatrix} W_x \\ W_y \\ W_z \end{pmatrix} = \begin{pmatrix} V_y W_z - V_z W_y \\ V_z W_x - V_x W_z \\ V_x W_y - V_y W_x \end{pmatrix}. \quad (\text{B.17})$$



Dans la pratique, il faut faire attention à exprimer les vecteurs $\vec{V}$ et $\vec{W}$ sur la même base.

Le produit vectoriel est antisymétrique :

$$\vec{V} \wedge \vec{W} = -\vec{W} \wedge \vec{V}.$$

Par conséquent, le produit vectoriel de tout vecteur par lui-même est nul : $\vec{V} \wedge \vec{V} = \vec{0}$.

Le produit vectoriel est bilinéaire :

$$\begin{aligned} \vec{V} \wedge (\vec{W} + \vec{W}') &= \vec{V} \wedge \vec{W} + \vec{V} \wedge \vec{W}' \quad \text{et} \quad \vec{V} \wedge (\lambda \vec{W}) = \lambda (\vec{V} \wedge \vec{W}), \\ (\vec{V} + \vec{V}') \wedge \vec{W} &= \vec{V} \wedge \vec{W} + \vec{V}' \wedge \vec{W} \quad \text{et} \quad (\lambda \vec{V}) \wedge \vec{W} = \lambda (\vec{V} \wedge \vec{W}). \end{aligned}$$

c) Produits vectoriels des vecteurs de la base

Les produits scalaires des vecteurs d'une base orthonormée directe sont :

$$\vec{u}_x \wedge \vec{u}_x = \vec{u}_y \wedge \vec{u}_y = \vec{u}_z \wedge \vec{u}_z = \vec{0}$$

et

$$\vec{u}_x \wedge \vec{u}_y = \vec{u}_z, \quad \vec{u}_y \wedge \vec{u}_z = \vec{u}_x \quad \text{et} \quad \vec{u}_z \wedge \vec{u}_x = \vec{u}_y. \quad (\text{B.18})$$



Ces règles de calcul, ainsi que la bilinéarité du produit vectoriel, sont utiles pour calculer le produit vectoriel de deux vecteurs. Par exemple :

$$\vec{u}_x \wedge (2\vec{u}_x + 10\vec{u}_y + 3\vec{u}_z) = 2\vec{u}_x \wedge \vec{u}_x + 10\vec{u}_x \wedge \vec{u}_y + 3\vec{u}_x \wedge \vec{u}_z = 10\vec{u}_z - 3\vec{u}_y.$$

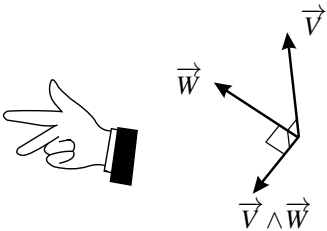
De manière générale, si $\vec{u}_\alpha$ et $\vec{u}_\beta$ sont deux vecteurs unitaires perpendiculaires, le vecteur $\vec{u}_\gamma = \vec{u}_\alpha \wedge \vec{u}_\beta$ est le vecteur unitaire tel que $(\vec{u}_\alpha, \vec{u}_\beta, \vec{u}_\gamma)$ est une base orthonormée directe.

d) Interprétation géométrique

Deux vecteurs $\vec{V}$ et $\vec{W}$ sont colinéaires si et seulement si : $\vec{V} \wedge \vec{W} = \vec{0}$.

Dans le cas où $\vec{V}$ et $\vec{W}$ ne sont pas colinéaires, $\vec{V} \wedge \vec{W}$ est un vecteur perpendiculaire à la fois au vecteur $\vec{V}$ et au vecteur $\vec{W}$.

On trouve le sens du produit vectoriel de deux vecteurs non colinéaires en appliquant la règle de la main droite : si l'on place le pouce selon le vecteur $\vec{V}$ et l'index selon le vecteur $\vec{W}$, alors le majeur donne le sens du produit vectoriel $\vec{V} \wedge \vec{W}$ (voir figure).



e) Norme du produit vectoriel

On admet la formule générale :

$$\|\vec{V} \wedge \vec{W}\| = \|\vec{V}\| \times \|\vec{W}\| \times |\sin(\widehat{(\vec{V}, \vec{W})})|,$$

(B.19)

où $\widehat{(\vec{V}, \vec{W})}$ est l'angle entre les vecteurs $\vec{V}$ et $\vec{W}$.

Les deux cas particuliers suivants sont très importants :

- si $\vec{V}$ et $\vec{W}$ sont colinéaires : $\widehat{(\vec{V}, \vec{W})} = 0$ ou π , donc le sinus de cet angle est nul ; on retrouve le fait que le produit vectoriel de vecteurs colinéaires est nul ;
- si $\vec{V}$ et $\vec{W}$ sont perpendiculaires, $\widehat{(\vec{V}, \vec{W})} = \pm \frac{\pi}{2}$ donc $|\sin(\widehat{(\vec{V}, \vec{W})})| = 1$, et dans ce cas :

$$\|\vec{V} \wedge \vec{W}\| = \|\vec{V}\| \times \|\vec{W}\|.$$

4.3 Transformations géométriques

a) Translations

Définition La **translation** $\mathcal{T}_{\vec{T}}$ de vecteur $\vec{T}$ est la transformation de l'espace qui au point M fait correspondre le point $M' = \mathcal{T}_{\vec{T}}(M)$ tel que :

$$\overrightarrow{MM'} = \vec{T}.$$

Expression dans un repère Dans un repère cartésien $(Oxyz)$, si $\vec{T} = T_x \vec{u}_x + T_y \vec{u}_y + T_z \vec{u}_z$, les coordonnées (x', y', z') de M' se déduisent des coordonnées (x, y, z) de M par les formules :

$$\begin{cases} x' = x + T_x \\ y' = y + T_y \\ z' = z + T_z \end{cases}$$

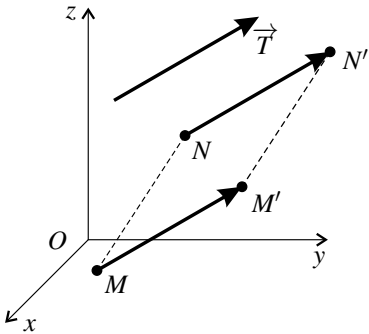


Figure B.17 – Translation de vecteur $\vec{T}$.

Propriétés La translation $\mathcal{T}_{\vec{T}}$ laisse les vecteurs invariants.

La translation $\mathcal{T}_{\vec{T}}$ conserve les distances.

La translation $\mathcal{T}_{\vec{T}}$ transforme un repère orthonormé direct de l'espace en un repère orthonormé direct.

CHAPITRE B – OUTILS MATHÉMATIQUES

b) Rotation autour d'un axe orienté

Définition La rotation $\mathcal{R}_{(O, \vec{u}), \alpha}$ d'angle α autour de l'axe orienté $(O, \vec{u})$ (axe passant par O et dirigé par le vecteur unitaire $\vec{u}$) est la transformation de l'espace qui au point M fait correspondre le point $M' = \mathcal{R}_{(O, \vec{u}), \alpha}(M)$ tel que :

$$\vec{OM'} = (\vec{u} \cdot \vec{OM}) \vec{u} + \cos \alpha \left(\vec{OM} - (\vec{u} \cdot \vec{OM}) \vec{u} \right) + \sin \alpha \vec{u} \wedge \vec{OM}. \tag{B.20}$$

Cette formule mathématique n'est pas à connaître et on retiendra plutôt l'interprétation géométrique.

Interprétation géométrique Soit H , le projeté orthogonal de M sur l'axe $(O, \vec{u})$, M' se trouve sur le cercle de centre H et d'axe $(O, \vec{u})$ (cercle contenu dans un plan perpendiculaire $(O, \vec{u})$). L'angle algébrique, orienté par $\vec{u}$, entre les vecteurs $\vec{HM}$ et $\vec{HM'}$ est égal à α (voir figure B.18).

La relation (B.20) peut s'écrire :

$$\vec{HM'} = \cos \alpha \vec{HM} + \sin \alpha \vec{u} \wedge \vec{HM}.$$

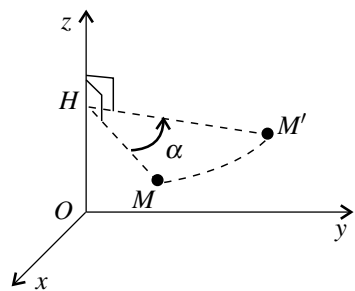


Figure B.18 – Rotation $\mathcal{R}_{(Oz), \alpha}$ d'angle α autour de l'axe (Oz) .

Expression dans un repère Dans le cas où $\vec{u} = \vec{u}_z$, troisième vecteur d'un repère orthonormé direct $(Oxyz)$, les formules de transformation des coordonnées sont :

$$\begin{cases} x' = x \cos \alpha - y \sin \alpha \\ y' = x \sin \alpha + y \cos \alpha \\ z' = z \end{cases}$$

La rotation d'angle α autour de (Oz) s'exprime en les coordonnées cylindriques (r, θ, z) d'axe (Oz) par les relations :

$$\begin{cases} r' = r \\ \theta' = \theta + \alpha \\ z' = z \end{cases}$$

La rotation d'angle α autour de (Oz) s'exprime en coordonnées sphériques (r, θ, φ) d'axe (Oz) par les relations :

$$\begin{cases} r' = r \\ \theta' = \theta \\ \varphi' = \varphi + \alpha \end{cases}$$

Transformation des vecteurs La rotation $\mathcal{R}_{(O,\vec{u}),\alpha}$ transforme les vecteurs de l'espace selon la formule (B.20). En coordonnées cartésiennes, dans le cas où $\vec{u} = \vec{u}_z$, le vecteur $\vec{V} = V_x\vec{u}_x + V_y\vec{u}_y + V_z\vec{u}_z$ est transformé en $\vec{V}' = V'_x\vec{u}_x + V'_y\vec{u}_y + V'_z\vec{u}_z$ selon les formules suivantes :

$$\begin{pmatrix} V_x \\ V_y \\ V_z \end{pmatrix} \mapsto \begin{pmatrix} V'_x \\ V'_y \\ V'_z \end{pmatrix} = \begin{pmatrix} V_x \cos \alpha - V_y \sin \alpha \\ V_x \sin \alpha + V_y \cos \alpha \\ V_z \end{pmatrix}.$$

Propriétés La rotation $\mathcal{R}_{(O,\vec{u}),\alpha}$ laisse invariants tous les points de l'axe $(O, \vec{u})$.
La rotation $\mathcal{R}_{(O,\vec{u}),\alpha}$ conserve les normes des vecteurs et les distances entre les points.
La rotation $\mathcal{R}_{(O,\vec{u}),\alpha}$ transforme un repère orthonormé direct de l'espace en un repère orthonormé direct.

c) Symétrie par rapport à un plan

Définition Soit (P) le plan passant par un point A et dirigé par deux vecteurs unitaires orthogonaux $\vec{u}_\alpha$ et $\vec{u}_\beta$. La **symétrie** $\mathcal{S}_{(P)}$ par rapport au plan (P) est la transformation d'espace qui au point M fait correspondre le point $M' = \mathcal{S}_{(P)}(M)$ défini par :

$$\overrightarrow{AM'} = (\vec{u}_\alpha \cdot \overrightarrow{AM})\vec{u}_\alpha + (\vec{u}_\beta \cdot \overrightarrow{AM})\vec{u}_\beta - (\vec{u}_\gamma \cdot \overrightarrow{AM})\vec{u}_\gamma, \tag{B.21}$$

où $\vec{u}_\gamma = \vec{u}_\alpha \wedge \vec{u}_\beta$.

Cette formule mathématique n'est pas à connaître et on retiendra plutôt l'interprétation géométrique.

Interprétation géométrique M' se trouve sur la droite passant par M et orthogonale au plan (P) , à même distance que M de ce plan (voir figure B.19).

Expression dans un repère Dans le cas où $\vec{u}_\alpha = \vec{u}_x$ et $\vec{u}_\beta = \vec{u}_y$ sont les deux premiers vecteurs de la base d'un repère orthonormé $(Oxyz)$, le plan (P) peut être noté (Axy) et les formules de transformations des coordonnées sont :

$$\begin{cases} x' = x \\ y' = y \\ z' = 2z_A - z \end{cases}$$

où z_A est la troisième coordonnée du point A .

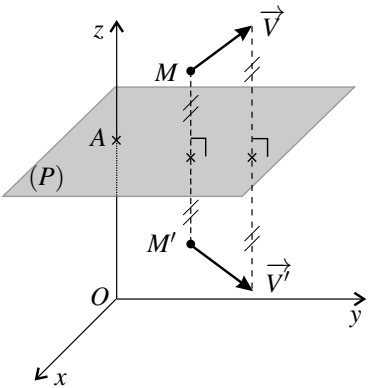


Figure B.19 – Symétrie $\mathcal{S}_{(P)}$ par rapport au plan $(P) = (Axy)$.

Transformation des vecteurs La symétrie $\mathcal{S}_{(Axy)}$ transforme les vecteurs de l'espace selon les formules suivantes :

$$\begin{pmatrix} V_x \\ V_y \\ V_z \end{pmatrix} \mapsto \begin{pmatrix} V'_x \\ V'_y \\ V'_z \end{pmatrix} = \begin{pmatrix} V_x \\ V_y \\ -V_z \end{pmatrix}.$$

CHAPITRE B – OUTILS MATHÉMATIQUES

De manière plus générale, si $(\vec{u}_\alpha, \vec{u}_\beta, \vec{u}_\gamma)$ est une base orthonormée telle que $\vec{u}_\alpha$ et $\vec{u}_\beta$ sont parallèles au plan (P) de la symétrie $\mathcal{S}_{(P)}$, la loi de transformation des vecteurs est :

$$\begin{pmatrix} V_\alpha \\ V_\beta \\ V_\gamma \end{pmatrix} \mapsto \begin{pmatrix} V'_\alpha \\ V'_\beta \\ V'_\gamma \end{pmatrix} = \begin{pmatrix} V_\alpha \\ V_\beta \\ -V_\gamma \end{pmatrix}.$$

Propriétés La symétrie $\mathcal{S}_{(P)}$ laisse tous les points du plan (P) invariants.

La symétrie $\mathcal{S}_{(P)}$ conserve les normes des vecteurs et les distances entre les points.

La symétrie $\mathcal{S}_{(P)}$ transforme une base repère orthonormé direct en un repère orthonormé *indirect*. Par exemple, la base directe $(\vec{u}_\alpha, \vec{u}_\beta, \vec{u}_\gamma)$ est transformée en $(\vec{u}_\alpha, \vec{u}_\beta, -\vec{u}_\gamma)$ qui est indirecte.

4.4 Courbes planes

Ce paragraphe indique comment reconnaître différents types de courbes par leur équation cartésienne. On se restreint à ces courbes planes, contenus dans le plan (Oxy) et le point M est repéré uniquement par deux coordonnées (x, y) . Le repère $(O, \vec{u}_x, \vec{u}_y)$ est orthonormé.

a) Droites

L'équation cartésienne d'une droite (D) est de la forme :

$$ax + by + c = 0, \tag{B.22}$$

où a , b et c sont des constantes et $(a, b) \neq (0, 0)$. Tout ensemble défini par une équation cartésienne de cette forme est une droite.

Un vecteur directeur de la droite est le vecteur $\vec{V} = b\vec{u}_x - a\vec{u}_y$. La droite est orthogonale au vecteur $a\vec{u}_x + b\vec{u}_y$.

b) Cercles

L'équation cartésienne d'un cercle est de la forme :

$$x^2 + y^2 + ax + by + c = 0 \tag{B.23}$$

où a , b et c sont des constantes telles que $c < \frac{1}{4}(a^2 + b^2)$. Tout ensemble défini par une équation cartésienne de cette forme est un cercle.

Où se trouve le centre du cercle et quel est son rayon ? Pour répondre à ces questions il suffit de réécrire l'équation sous la forme :

$$\left(x - \frac{a}{2}\right)^2 + \left(y - \frac{b}{2}\right)^2 = \frac{1}{4}(a^2 + b^2) - c.$$

Ainsi, si C est le point de coordonnées $\left(\frac{a}{2}, \frac{b}{2}\right)$ et si $R = \sqrt{\frac{1}{4}(a^2 + b^2) - c}$ (cette définition à

un sens du fait de l'inégalité ci-dessus), l'équation cartésienne se réécrit :

$$CM^2 = R^2.$$

Donc la courbe d'équation (B.23) est le cercle de centre C et de rayon R .

c) Coniques

L'équation cartésienne d'une conique est de la forme :

$$ax^2 + 2bxy + cy^2 + dx + ey + k = 0 \tag{B.24}$$

où a, b, c, d, e et k sont des constantes. On suppose dans ce qui suit ces constantes sont telles que la courbe n'est pas l'ensemble vide.

Classification des coniques : On appelle discriminant de la conique $\Delta = ac - b^2$.

- si $\Delta > 0$, la courbe est une conique de type **ellipse**,
- si $\Delta = 0$, la courbe est une conique de type **parabole**,
- si $\Delta < 0$, la courbe est une conique de type **hyperbole**.

Remarque

Un cercle est une conique de type ellipse.

Équation réduite d'une ellipse :

La courbe d'équation :

$$\frac{x^2}{A^2} + \frac{y^2}{B^2} = 1,$$

où A et B sont des constantes positives, est une ellipse. Elle admet l'origine O pour centre de symétrie et les axes (Ox) et (Oy) pour axes de symétrie (voir figure B.20). Ses dimensions sont $2A$ suivant (Ox) et $2B$ suivant (Oy) . Le **demi-grand axe** de l'ellipse est $\max\{A, B\}$ et le **demi-petit axe** est $\min\{A, B\}$.

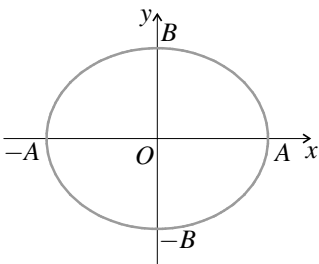


Figure B.20 – L'ellipse d'équation $\frac{x^2}{A^2} + \frac{y^2}{B^2} = 1$.

Équation réduite d'une parabole : La courbe d'équation $y = Ax^2$, où A est une constante, est une parabole admettant l'axe (Oy) pour axe de symétrie. La courbe d'équation $x = By^2$, où B est une constante, est une parabole admettant l'axe (Ox) pour axe de symétrie (voir figure B.21).

CHAPITRE B – OUTILS MATHÉMATIQUES

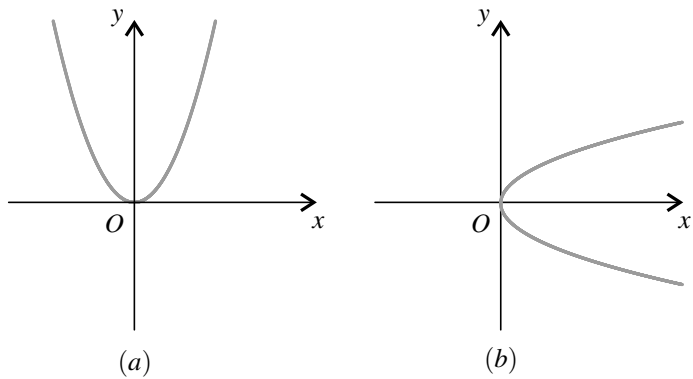


Figure B.21 – Les paraboles d'équation $y = Ax^2$ avec $A > 0$ (courbe (a)) et $x = By^2$ avec $B > 0$ (courbe (b)).

Équation réduite d'une hyperbole : La courbe d'équation $\frac{x^2}{A^2} - \frac{y^2}{B^2} = 1$, où A et B sont des constantes positives, est une hyperbole admettant l'origine O pour centre de symétrie et les axes (Ox) et (Oy) pour axes de symétrie.

Il en est de même pour la courbe d'équation $-\frac{x^2}{A^2} + \frac{y^2}{B^2} = 1$.

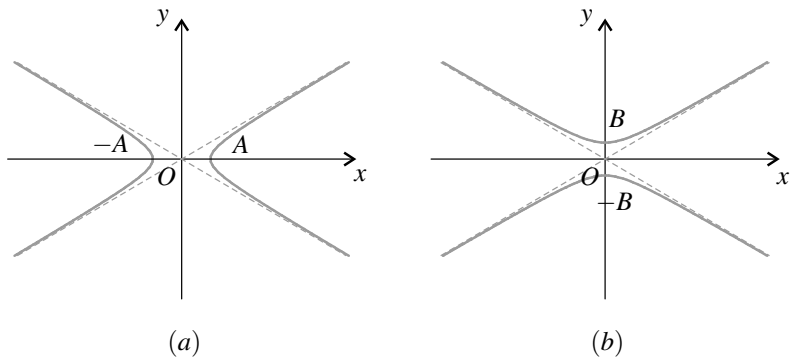


Figure B.22 – Les hyperboles d'équation $\frac{x^2}{A^2} - \frac{y^2}{B^2} = 1$ (courbe (a)) et $-\frac{x^2}{A^2} + \frac{y^2}{B^2} = 1$ (courbe (b)). Les droites en pointillé sont les asymptotes.

Pour des points éloignés de l'origine, tels que $|x| \gg A$ et $|y| \gg B$, les équations de ces deux hyperboles se ramènent à : $y^2 \simeq \frac{B^2}{A^2}x^2$. Elles sont donc quasiment confondues avec les droites d'équations $y = \frac{B}{A}x$ et $y = -\frac{B}{A}x$. Ces droites sont les **asymptotes** de l'hyperbole.

d) Courbe définie par une équation polaire

Définition Une courbe définie par une équation polaire (r, θ) est une courbe d'équation :

$$r = f(\theta),$$

où f est une fonction quelconque à valeurs positives. Elle est formée des points M tels que :

$$\overrightarrow{OM} = f(\theta)\overrightarrow{u_r} = f(\theta)\cos\theta\overrightarrow{u_x} + f(\theta)\sin\theta\overrightarrow{u_y},$$

pour θ variable dans $\mathbb{R}$ ou sur un intervalle de $\mathbb{R}$.

Méthode de tracé Les calculettes sont dotées de grapheurs permettant de visualiser les courbes définies par une équation polaire. Il est attendu du candidat aux concours qu'il sache utiliser cette fonction.

Propriétés Si $f(\theta)$ est périodique de période $2n\pi$, où n est un entier, la courbe est fermée et fait n fois le tour de l'origine.

Un vecteur tangent à la courbe définie par l'équation polaire $r = f(\theta)$ est :

$$\overrightarrow{T} = \frac{d\overrightarrow{OM}}{d\theta} = f'(\theta)\overrightarrow{u_r} + f(\theta)\overrightarrow{u_\theta}.$$

Exemple

1. La courbe définie par

$$r = \frac{p}{1 + e \cos \theta},$$

où $\theta \in [0, 2\pi]$, avec p constante positive et e constante telle que $0 < e < 1$, est une ellipse non centrée sur l'origine.

La distance $r(\theta)$ entre O et M varie entre $r_{\min} = \frac{p}{1+e}$ (pour $\theta = 0$) et $r_{\max} = \frac{p}{1-e}$ (pour $\theta = \pi$, voir figure B.23).

2. La courbe définie par

$$r = \frac{p}{1 + \cos \theta}$$

où $\theta \in]-\pi, \pi[$, avec p constante positive, est une parabole.

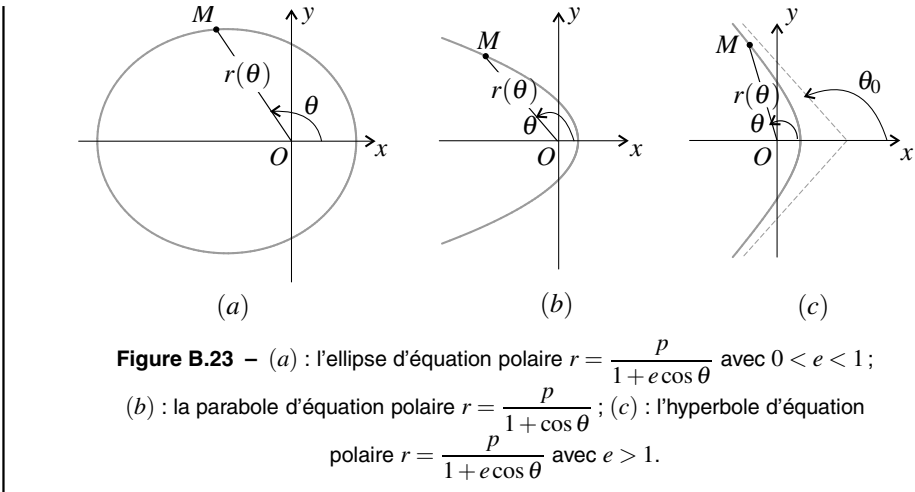
3. La courbe définie par

$$r = \frac{p}{1 + e \cos \theta}$$

où $\theta \in]-\theta_0, \theta_0[$, où p est une constante positive, e une constante supérieure à 1 et $\theta_0 = \arccos\left(-\frac{1}{e}\right)$, est une branche d'hyperbole.

Le dénominateur dans l'expression de r s'annule pour $\theta = \pm\theta_0$ ce qui correspond aux deux asymptotes (voir figure B.23).

CHAPITRE B – OUTILS MATHÉMATIQUES



4.5 Courbes paramétrées

Définition Une courbe paramétrée du plan (Oxy) est une courbe décrite au cours du temps par un point M de coordonnées :

$$\begin{cases} x(t) = f(t) \\ y(t) = g(t) \end{cases} \quad (\text{B.25})$$

où $f(t)$ et $g(t)$ sont deux fonctions continues et dérivables et t un paramètre variant dans $\mathbb{R}$ ou un intervalle.

Méthode de tracé Les calculettes sont dotées de graphes permettant de visualiser les courbes paramétrées. Il est attendu du candidat aux concours qu'il sache utiliser cette fonction.

Propriétés Un vecteur tangent à la courbe définie par (B.25) est :

$$\vec{T} = \frac{d\vec{OM}}{dt} = f'(t)\vec{u}_x + g'(t)\vec{u}_y.$$

Si t est le temps, $\vec{T}$ est la vitesse $\vec{v}$ du point mobile M .

Exemple

Un exemple à connaître est le cas où $f(t)$ et $g(t)$ sont des fonctions sinusoïdales de même pulsation ω , soit :

$$\begin{cases} x(t) = A \cos(\omega t) \\ y(t) = B \cos(\omega t - \varphi) \end{cases} \quad (\text{B.26})$$

où A et B sont des constantes positives et φ une constante appartenant à $[0, 2\pi[$. On a supposé la phase initiale de la fonction $f(t)$ nulle sans perte de généralité puisqu'on peut toujours se ramener à ce cas par un choix adapté de l'origine des temps.

Expérimentalement, on visualise ce type de courbe sur un oscilloscope en mode XY ou avec un logiciel d'acquisition de signaux lorsqu'on porte en abscisse et en ordonnée deux signaux sinusoïdaux de même pulsation.

Quelle est la nature géométrique de cette courbe paramétrée par la formule (B.25) suivant la valeur de φ ?

Cas où $\varphi \neq 0$ et $\varphi \neq \pi$: D'après la formule du cosinus d'une différence (voir paragraphe 5) :

$$\begin{cases} x(t) = A \cos(\omega t) \\ y(t) = B \cos(\omega t) \cos \varphi + B \sin(\omega t) \sin \varphi \end{cases} \quad \text{d'où} \quad \begin{cases} \cos(\omega t) = \frac{x(t)}{A} \\ \sin(\omega t) = -\frac{x(t) \cos \varphi}{A \sin \varphi} + \frac{y(t)}{B \sin \varphi} \end{cases}$$

puisque $\sin \varphi \neq 0$. La relation $\cos^2(\omega t) + \sin^2(\omega t) = 1$ donne alors :

$$\left(\frac{x(t)}{A}\right)^2 + \left(\frac{y(t)}{B \sin \varphi} - \frac{x(t) \cos \varphi}{A \sin \varphi}\right)^2 = 1,$$

soit :

$$\frac{x(t)^2}{A^2} + \frac{y(t)^2}{B^2} - \frac{2 \cos \varphi}{AB} x(t)y(t) = \sin^2 \varphi.$$

On trouve l'équation d'une conique (voir ci-dessus) dont le discriminant est :

$$\Delta = \frac{1}{A^2} \frac{1}{B^2} - \left(\frac{\cos \varphi}{AB}\right)^2 = \frac{\sin^2 \varphi}{A^2 B^2}.$$

Δ étant positif, la courbe est une ellipse.

Dans le cas où $\cos \varphi = 0$, soit si $\varphi = \frac{\pi}{2}$ ou $\varphi = \frac{3\pi}{2}$, l'équation devient :

$$\frac{x(t)^2}{A^2} + \frac{y(t)^2}{B^2} = 1.$$

C'est une équation réduite : l'ellipse est symétrique par rapport aux axes (Ox) et (Oy) . Le paramétrage de la courbe prend dans ce cas la forme particulière :

$$\begin{cases} x(t) = A \cos(\omega t) \\ y(t) = \pm B \sin(\omega t) \end{cases}.$$

Pour une valeur quelconque de φ les axes de symétrie de l'ellipse sont inclinés par rapport aux axes du repère. Visuellement, les cas où $\varphi = \frac{\pi}{2}$ ou $\varphi = \frac{3\pi}{2}$ se reconnaissent donc très facilement (voir figure B.24).

CHAPITRE B – OUTILS MATHÉMATIQUES

Cas où $\varphi = 0$: On a alors : $y(t) = B \cos(\omega t) = \frac{B}{A} x(t)$. La courbe paramétrée est un segment de droite de pente positive (voir figure B.24). Visuellement, ce cas se reconnaît immédiatement.

Cas où $\varphi = \pi$: On a alors : $y(t) = -B \cos(\omega t) = -\frac{B}{A} x(t)$. La courbe paramétrée est un segment de droite de pente négative dans ce cas qui est très facilement reconnaissable (voir figure B.24).

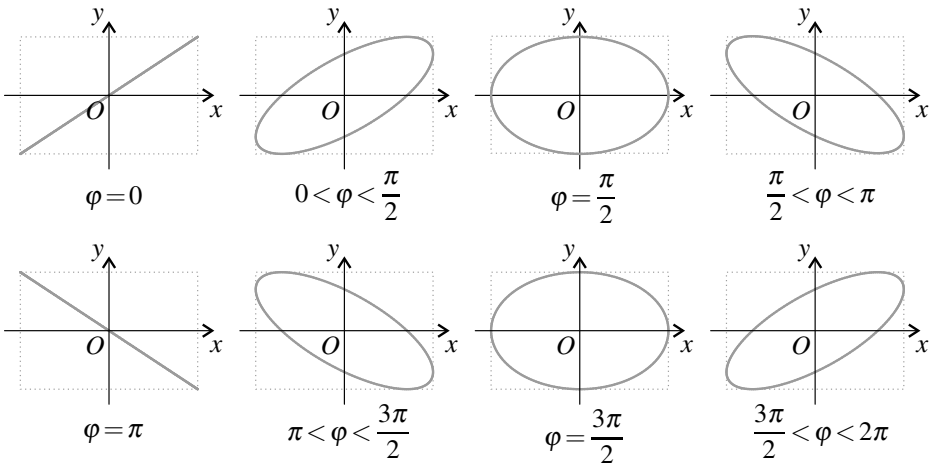


Figure B.24 – Courbe paramétrée $x(t) = A \cos(\omega t)$, $y(t) = B \cos(\omega t - \varphi)$ des valeurs de φ croissantes. En pointillé, le rectangle de cotés $2A$ suivant (Ox) et $2B$ suivant (Oy) dans lequel la courbe est inscrite.

4.6 Longueurs, aires et volumes classiques

Les expressions de périmètres, surfaces et volumes données dans ce paragraphe sont à connaître par cœur.

Cercle de rayon R Le périmètre P et la surface intérieure S sont :

$$P = 2\pi R \quad \text{et} \quad S = \pi R^2 \tag{B.27}$$

Sphère de rayon R La surface S et le volume intérieur V sont :

$$S = 4\pi R^2 \quad \text{et} \quad V = \frac{4}{3}\pi R^3 \tag{B.28}$$

Cylindre de rayon R et hauteur h La surface latérale S et le volume intérieur V sont :

$$S = 2\pi R h \quad \text{et} \quad V = \pi R^2 h \tag{B.29}$$

4.7 Barycentre d'un système de points

a) Cas de deux points

On définit le centre de gravité G des points M_1 et M_2 affectés des masses m_1 et m_2 par la relation :

$$(m_1 + m_2) \overrightarrow{AG} = m_1 \overrightarrow{AM_1} + m_2 \overrightarrow{AM_2}, \tag{B.30}$$

où A est un point quelconque.

En écrivant la relation (B.30) pour $A = O$, origine du repère cartésien $(Oxyz)$, on montre que les coordonnées du centre de gravité G vérifient :

$$\begin{cases} (m_1 + m_2) x_G = m_1 x_{M_1} + m_2 x_{M_2} \\ (m_1 + m_2) y_G = m_1 y_{M_1} + m_2 y_{M_2} \\ (m_1 + m_2) z_G = m_1 z_{M_1} + m_2 z_{M_2} \end{cases} \tag{B.31}$$

La relation (B.30), écrite pour $A = M_1$ devient

$$(m_1 + m_2) \overrightarrow{M_1 G} = m_2 \overrightarrow{M_1 M_2} \implies \overrightarrow{M_1 G} = \frac{m_2}{m_1 + m_2} \overrightarrow{M_1 M_2}.$$

Comme $\frac{m_2}{m_1 + m_2} \in]0; 1[$, le centre de gravité G est situé sur le segment $[M_1 M_2]$ et sa position dépend des masses m_1 et m_2 :

- lorsque $m_1 \gg m_2$, G est quasiment confondu avec M_1 ;
- lorsque $m_2 \gg m_1$, G est quasiment confondu avec M_2 ;
- lorsque $m_1 = m_2$, G est au milieu du segment $[M_1 M_2]$.

Lorsque les masses m_1 et m_2 sont identiques, le milieu du segment $[M_1 M_2]$ est le centre de symétrie du système. Le centre de gravité est donc confondu avec le centre de symétrie.

Remarque

En mathématiques, les masses m_1 et m_2 sont appelées poids, ils peuvent être négatifs et on parle de barycentre plutôt que de centre de gravité.

b) Généralisation

On peut généraliser cette relation pour un ensemble de N points M_i de masse m_i :

$$\left(\sum_{i=1}^N m_i \right) \overrightarrow{AG} = \sum_{i=1}^N m_i \overrightarrow{AM_i}.$$

Pour un solide situé dans le volume $\mathcal{V}$, cette somme devient une intégrale :

$$\left(\iiint_{M \in \mathcal{V}} \rho(M) dV(M) \right) \overrightarrow{AG} = \iiint_{M \in \mathcal{V}} \overrightarrow{AM} \rho(M) dV(M) \tag{B.32}$$

où $\rho(M) dV(M)$ est la masse infinitésimale contenue dans le volume infinitésimal $dV(M)$ entourant le point M .

CHAPITRE B – OUTILS MATHÉMATIQUES

c) Cas d'un solide

Pour un solide quelconque situé dans le volume $\mathcal{V}$, l'intégrale (B.32) conduit aux résultats suivants :

- pour un disque homogène, G est confondu avec le centre du disque ;
- pour un cylindre homogène, G est situé au centre de symétrie du cylindre (à mi-hauteur sur son axe) ;
- pour une boule homogène, G est confondu avec le centre de la boule.

D'une manière générale :

Lorsque le solide est homogène et qu'il présente un centre de symétrie, G est en ce point.

Lorsque le solide est homogène et qu'il possède un plan de symétrie, G appartient à ce plan de symétrie. Lorsqu'il en possède plusieurs, G se situe sur leur intersection.

d) Associativité du barycentre

Si l'on connaît le barycentre G_1 d'un solide $\mathcal{S}_1$ de masse m_1 et le barycentre G_2 d'un solide $\mathcal{S}_2$ de masse m_2 , le barycentre G du solide $\mathcal{S} = \mathcal{S}_1 \cup \mathcal{S}_2$ constitué de la réunion des deux solides précédents est égal au barycentre des points G_1 et G_2 respectivement affectés des masses m_1 et m_2 , ce qui s'écrit :

$$(m_1 + m_2)\overrightarrow{AG} = m_1\overrightarrow{AG_1} + m_2\overrightarrow{AG_2}. \tag{B.33}$$

5 Trigonométrie

5.1 Angle orienté

a) Orientation du plan

Orienter le plan de figure consiste à choisir un sens positif qui peut être le sens trigonométrique (ou antihoraire) ou bien le sens horaire. Ce choix est équivalent au choix d'un vecteur unitaire $\vec{n}$ perpendiculaire à la figure : $\vec{n}$ dirigé vers l'avant de la figure pour le sens trigonométrique, et $\vec{n}$ dirigé vers l'arrière de la figure pour le sens horaire (voir figure B.25).

Plus généralement, on peut orienter un plan quelconque de l'espace en choisissant un vecteur $\vec{n}$ orthogonal à ce plan et en appliquant la règle de la main droite : le pouce de la main droite pointant dans la direction et le sens de $\vec{n}$, les autres doigts repliés donne le sens positif correspondant à $\vec{n}$ qui est appelé *sens positif autour du vecteur $\vec{n}$* (voir figure B.25).

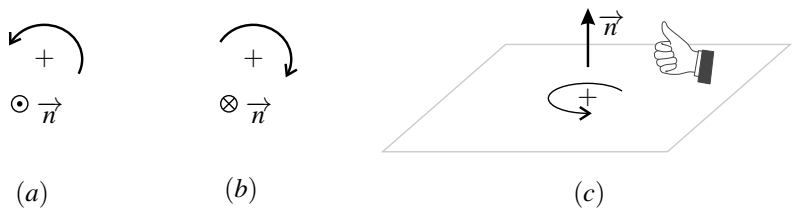


Figure B.25 – (a) : plan de figure orienté dans le sens trigonométrique, $\vec{n}$ pointe vers l'avant de la figure ; (b) : plan de figure orienté dans le sens horaire, $\vec{n}$ pointe vers l'arrière de la figure ; (c) : règle de la main droite donnant le sens positif autour du vecteur $\vec{n}$.



Cette convention est à bien connaître. Elle est liée au choix d'un trièdre direct et à la définition du produit vectoriel. Le produit vectoriel de deux vecteurs $\vec{V}$ et $\vec{W}$ appartenant au plan de figure, orienté par $\vec{n}$, est :

$$\vec{V} \wedge \vec{W} = \|\vec{V}\| \|\vec{W}\| \sin \theta \vec{n}, \tag{B.34}$$

où $\theta = (\widehat{\vec{V}, \vec{W}})$ est l'angle orienté, de $\vec{V}$ vers $\vec{W}$ dans cet ordre.

b) Utilisation des angles orientés

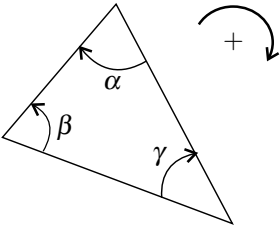
L'utilisation des **angles orientés** (on dit aussi angles algébriques) suppose d'avoir précisé clairement un choix d'orientation. Les angles doivent être représentés sur la figure par une flèche courbe. Ils sont positifs quand cette flèche est dans le sens positif et négatif quand elle est dans l'autre sens. Un angle représenté par une courbe sans flèche est par convention un angle arithmétique, toujours positif.

L'angle orienté allant d'un vecteur $\vec{V}$ à un vecteur $\vec{W}$ est noté $(\widehat{\vec{V}, \vec{W}})$.

CHAPITRE B – OUTILS MATHÉMATIQUES



Certaines lois géométriques sont vérifiées uniquement par les angles arithmétiques. Pour les appliquer avec les angles algébriques, il faut s'appuyer sur une figure claire. Par exemple, la propriété selon laquelle la somme des angles d'un triangle est égale à π s'écrit dans le cas de la figure ci-contre :



$$|\alpha| + |\beta| + |\gamma| = \pi \text{ soit } \alpha - \beta + \gamma = \pi,$$

puisque $\alpha > 0$, $\beta < 0$ et $\gamma > 0$.

5.2 Fonctions trigonométriques

a) Définitions, cercle trigonométrique

Le cercle trigonométrique, cercle de rayon 1 centré à l'origine d'un repère (Oxy) permet de visualiser les fonctions trigonométriques sinus et cosinus d'un angle θ .

M étant le point du cercle tel que l'angle entre $\vec{u}_x$ et $\vec{OM}$ dans cet ordre est $\widehat{(\vec{u}_x, \vec{OM})} = \theta$, on pose par définition (voir figure B.26) :

$$x_M = \cos \theta \text{ et } y_M = \sin \theta. \tag{B.35}$$

La fonction tangente se définit à partir des fonctions cosinus et sinus par :

$$\tan \theta = \frac{\sin \theta}{\cos \theta}. \tag{B.36}$$

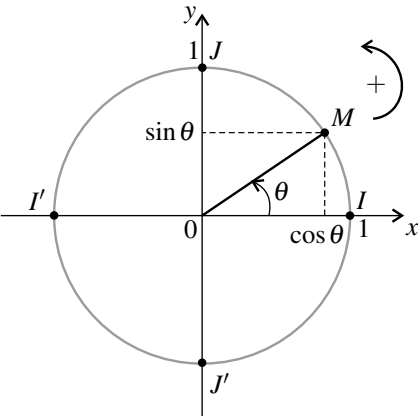


Figure B.26 – Cercle trigonométrique, définition des fonctions sinus et cosinus.

La représentation du cercle trigonométrique permet de retrouver très facilement les valeurs remarquables suivantes :

CHAPITRE B – OUTILS MATHÉMATIQUES

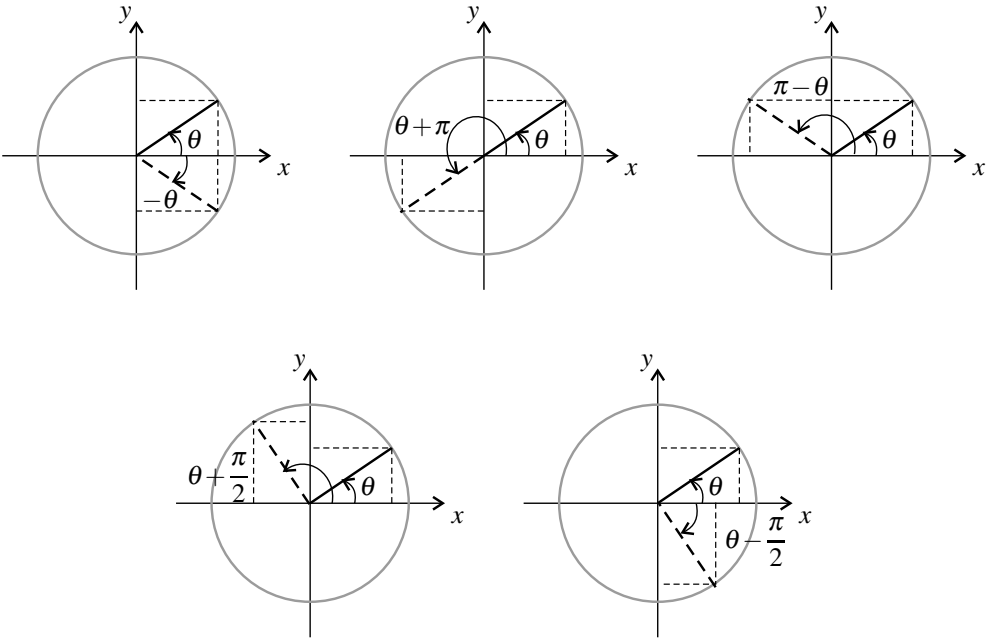


Figure B.27 – Schémas pour retrouver les relations entre les lignes trigonométriques de θ et les lignes trigonométriques de $-\theta$, $\theta + \pi$, $\pi - \theta$, $\theta + \frac{\pi}{2}$ ou $\theta - \frac{\pi}{2}$.

$$\begin{aligned} \cos p + \cos q &= 2 \cos \frac{p+q}{2} \cos \frac{p-q}{2} \\ \cos p - \cos q &= -2 \sin \frac{p+q}{2} \sin \frac{p-q}{2} \\ \sin p + \sin q &= 2 \sin \frac{p+q}{2} \cos \frac{p-q}{2} \\ \sin p - \sin q &= 2 \cos \frac{p+q}{2} \sin \frac{p-q}{2} \end{aligned} \quad , \quad \text{(B.41)}$$

ainsi que les les formules suivantes qui transforment un produit de deux cosinus ou sinus en une somme de deux cosinus ou sinus :

$$\begin{aligned} \cos a \cos b &= \frac{1}{2} (\cos(a+b) + \cos(a-b)) \\ \sin a \cos b &= \frac{1}{2} (\sin(a+b) + \sin(a-b)) \\ \sin a \sin b &= \frac{1}{2} (\cos(a+b) - \cos(a-b)) \end{aligned} \quad . \quad \text{(B.42)}$$

5.3 Nombres complexes

Les **nombres complexes** sont les nombres de la forme : $\underline{z} = x + iy$ où x et y sont réels et où i est une racine carré de -1 : $i^2 = -1$.

x est la **partie réelle** de $\underline{z}$ qui est notée $\text{Re}(\underline{z})$; y est la **partie imaginaire** de $\underline{z}$ qui est notée $\text{Im}(\underline{z})$.

Remarque

Il est de tradition en physique de noter les nombres complexes par des lettres soulignées, ce qui permet d'utiliser la même notation pour un réel s et un complexe $\underline{s}$ correspondant à la même grandeur physique.

Représentation des nombres complexes dans le plan Au nombre complexe $\underline{z} = x + iy$ on peut associer le point du plan (Oxy) de coordonnées cartésiennes (x, y) . x est l'abscisse de M et la partie réelle de $\underline{z}$; y est l'ordonnée de M et la partie imaginaire de $\underline{z}$ (voir figure B.28).

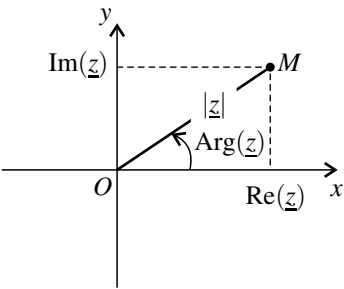


Figure B.28 – Interprétation géométrique de la partie réelle, de la partie imaginaire, du module et de l'argument d'un nombre complexe $\underline{z}$.

Module et argument On peut aussi repérer le point M par ses coordonnées polaires (r, θ) où $r = OM$ et $\theta = (\vec{u}_x, \vec{OM})$, angle orienté dans le sens trigonométrique).

Le **module** du nombre complexe $\underline{z}$, noté $|\underline{z}|$, est :

$$|\underline{z}| = r = \sqrt{x^2 + y^2}. \tag{B.43}$$

L'**argument** du nombre complexe $\underline{z}$, noté $\text{Arg}(\underline{z})$, est une détermination de l'angle entre $\vec{u}_x$ et $\vec{OM}$ soit :

$$\text{Arg}(\underline{z}) = \theta = \begin{cases} \arccos\left(\frac{\text{Re}(\underline{z})}{|\underline{z}|}\right) & \text{si } \text{Im}(\underline{z}) > 0 \\ -\arccos\left(\frac{\text{Re}(\underline{z})}{|\underline{z}|}\right) & \text{si } \text{Im}(\underline{z}) < 0 \end{cases} \tag{B.44}$$

où $\arccos(x)$ est la fonction réciproque de la fonction cosinus. Cette formule donne une valeur

CHAPITRE B – OUTILS MATHÉMATIQUES

comprise entre $-\pi$ et π . L'argument est défini modulo 2π : toute valeur différant de celle-ci d'un multiple entier de 2π est une valeur possible pour l'argument de $\underline{z}$.

Addition de deux nombres complexes On additionne les nombres complexes en additionnant leurs parties réelles et leurs parties imaginaires : $(x + iy) + (x' + iy') = (x + x') + i(y + y')$, soit :

$$\begin{aligned}\operatorname{Re}(\underline{z} + \underline{z}') &= \operatorname{Re}(\underline{z}) + \operatorname{Re}(\underline{z}') \\ \operatorname{Im}(\underline{z} + \underline{z}') &= \operatorname{Im}(\underline{z}) + \operatorname{Im}(\underline{z}')\end{aligned}\tag{B.45}$$



Le module d'une somme n'est pas égale à la somme des modules : $|\underline{z} + \underline{z}'| \neq |\underline{z}| + |\underline{z}'|$!

Produit de deux nombres complexes La multiplication est une opération bilinéaire :

$$(x + iy)(x' + iy') = xx' + ixy' + iyx' + i^2yy' = (xx' - yy') + i(xy' + yx'),$$

soit :

$$\begin{aligned}\operatorname{Re}(\underline{z} \underline{z}') &= \operatorname{Re}(\underline{z}) \operatorname{Re}(\underline{z}') - \operatorname{Im}(\underline{z}) \operatorname{Im}(\underline{z}') \\ \operatorname{Im}(\underline{z} \underline{z}') &= \operatorname{Re}(\underline{z}) \operatorname{Im}(\underline{z}') + \operatorname{Im}(\underline{z}) \operatorname{Re}(\underline{z}')\end{aligned}\tag{B.46}$$

De plus :

$$|\underline{z} \underline{z}'| = |\underline{z}| |\underline{z}'|\tag{B.47}$$

et

$$\operatorname{Arg}(\underline{z} \underline{z}') = \operatorname{Arg}(\underline{z}) + \operatorname{Arg}(\underline{z}')\tag{B.48}$$

6 Gradient d'un champ scalaire

6.1 Champ scalaire f et surfaces iso- f

Dans ce paragraphe on considère un **champ scalaire**, c'est-à-dire une fonction f qui associe un scalaire $f(M)$ à tout point M de l'espace.

On appelle **surface iso- f** une surface définie comme l'ensemble de points pour lesquels :

$$f(M) = C,$$

où C est une constante. Il existe une infinité de surfaces iso- f , une pour chaque valeur de C prise par f .

6.2 Dérivées partielles et différentielle

Le champ scalaire $f(M)$ s'exprime en fonction des coordonnées de M dans un système choisi. Par exemple en coordonnée cartésiennes f s'exprime en fonction des coordonnées x , y et z de M . C'est une fonction de trois variables $f(x, y, z)$.

On appelle **dérivée partielle** de f par rapport à x , notée $\frac{\partial f}{\partial x}$, la dérivée de la fonction $f(x, y, z)$ par rapport à sa variable x que l'on calcule en faisant comme si y et z étaient des constantes.

On définit de même la dérivée partielle par rapport à y notée $\frac{\partial f}{\partial y}$ et la dérivée partielle par rapport à z notée $\frac{\partial f}{\partial z}$.

Exemple

Pour la fonction $f(x, y, z) = x^2y^3 + z$:

$$\frac{\partial f}{\partial x} = 2xy^3, \quad \frac{\partial f}{\partial y} = 3x^2y^2 \quad \text{et} \quad \frac{\partial f}{\partial z} = 1$$

La **différentielle** de f est la variation élémentaire df de f associée à un déplacement élémentaire du point M . Pour l'exprimer, il faut se placer dans un système de coordonnées. Par exemple si M est représenté par des coordonnées cartésiennes (x, y, z) on a :

$$df = \frac{\partial f}{\partial x}dx + \frac{\partial f}{\partial y}dy + \frac{\partial f}{\partial z}dz, \tag{B.49}$$

expression dans laquelle apparaissent les dérivées partielles de f par rapport à chacune des trois coordonnées.

CHAPITRE B – OUTILS MATHÉMATIQUES

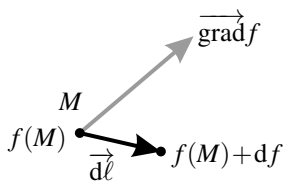
6.3 Vecteur gradient

a) Définition

Le **gradient** du champ scalaire f , noté $\overrightarrow{\text{grad}}f$, est le vecteur tel que la variation df de f que l'on observe si on se déplace de $\overrightarrow{d\ell}$ à partir de M (voir figure) est :

$$df = \overrightarrow{\text{grad}}f \cdot \overrightarrow{d\ell}.$$

(B.50)



Ce vecteur dépend du point M .

D'après l'interprétation géométrique du produit scalaire, df est maximale, pour un déplacement de longueur $\|\overrightarrow{d\ell}\|$ donnée, si le déplacement est dans la direction et le sens du vecteur $\overrightarrow{\text{grad}}f$. Si le déplacement est orthogonal au vecteur $\overrightarrow{\text{grad}}f$, df est nul donc f ne varie pas. Ainsi :

Le vecteur $\overrightarrow{\text{grad}}f$ est orthogonal aux surfaces iso- f .
Le vecteur $\overrightarrow{\text{grad}}f$ est orienté dans le sens où f croît le plus rapidement.

Exemple

Dans une modélisation très simple on peut considérer que la température en un point M situé à l'intérieur de la Terre dépend uniquement de la distance entre M et le centre O de la Terre. Les surfaces isothermes sont alors les sphères de centre O (en pointillé sur la figure).

La température augmente quand on se rapproche du centre de la Terre. Le gradient de température en tout point M est dirigé vers O . Il est orthogonal à la surface isotherme passant par M (voir figure B.29).

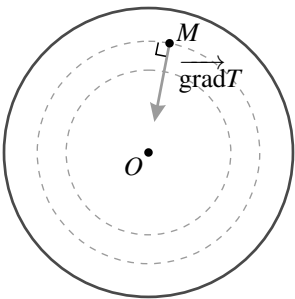


Figure B.29 –
Champ de
température à
l'intérieur de la Terre.

b) Expression du gradient en coordonnées cartésiennes

En coordonnées cartésiennes le déplacement élémentaire s'écrit :

$$\overrightarrow{d\ell} = dx\overrightarrow{u_x} + dy\overrightarrow{u_y} + dz\overrightarrow{u_z}.$$

Ainsi, la relation (B.50) définissant le gradient s'écrit :

$$df = (\overrightarrow{\text{grad}}f)_x dx + (\overrightarrow{\text{grad}}f)_y dy + (\overrightarrow{\text{grad}}f)_z dz,$$

en notant $(\overrightarrow{\text{grad}f})_x$, $(\overrightarrow{\text{grad}f})_y$ et $(\overrightarrow{\text{grad}f})_z$ les composantes du vecteur gradient de f sur les trois axes. En comparant cette expression à la relation (B.49) on trouve que :

$$(\overrightarrow{\text{grad}f})_x = \frac{\partial f}{\partial x}, \quad (\overrightarrow{\text{grad}f})_y = \frac{\partial f}{\partial y}, \quad \text{et} \quad (\overrightarrow{\text{grad}f})_z = \frac{\partial f}{\partial z},$$

soit :

$$\overrightarrow{\text{grad}f} = \frac{\partial f}{\partial x} \vec{u}_x + \frac{\partial f}{\partial y} \vec{u}_y + \frac{\partial f}{\partial z} \vec{u}_z.$$

 (B.51)

Cette expression à connaître par cœur.

c) Expressions du gradient en coordonnées cylindrique ou sphériques

Dans les systèmes de coordonnées cylindriques ou sphériques, le vecteur gradient a des expressions plus compliquées qui ne sont pas à connaître mais seront fournies lors de épreuves des concours si elles sont nécessaires.

En coordonnées cylindriques (r, θ, z) :

$$\overrightarrow{\text{grad}f} = \frac{\partial f}{\partial r} \vec{u}_r + \frac{1}{r} \frac{\partial f}{\partial \theta} \vec{u}_\theta + \frac{\partial f}{\partial z} \vec{u}_z. \tag{B.52}$$

En coordonnées sphériques (r, θ, φ) :

$$\overrightarrow{\text{grad}f} = \frac{\partial f}{\partial r} \vec{u}_r + \frac{1}{r} \frac{\partial f}{\partial \theta} \vec{u}_\theta + \frac{1}{r \sin \theta} \frac{\partial f}{\partial \varphi} \vec{u}_\varphi. \tag{B.53}$$

Les dérivées partielles qui apparaissent dans ces expressions se calculent en dérivant l'expression de f par rapport à l'une des variables, en faisant comme si les deux autres variables étaient des constantes.

Index

- +20 dB/décade, 385
- +40 dB/décade, 395
- −20 dB/décade, 381
- −40 dB/décade, 390
- échelle mésoscopique, 757
- échelle macroscopique, 757
- échelle microscopique, 757
- échelles logarithmiques, 1088
- échelon de tension, 286
- éclairage, 124
- éclairage spectral, 124
- écoulement permanent, 898
- électron-volt, 659
- électrons de conduction, 245
- électrons libres, 245
- énergie, 265, 819
- énergie cinétique, 27, 581
- énergie interne, 767
- énergie mécanique, 27, 589
- énergie magnétique, 976
- énergie potentielle, 27
- énergie potentielle de la force, 585
- énergie potentielle effective, 725
- équation d'état, 764
- équation d'état du gaz parfait, 764
- équation différentielle, 24
- équations horaires, 468
- équilibre thermodynamique local, 762
- état de diffusion, 593, 727
- état lié, 593, 727
- état libre, 593
- états, 755
- évaluation de type A, 1053
- évaluation de type B, 1053
- accélééré, 488
- accélération angulaire, 485
- additive, 244
- adiabatique, 809
- agitation thermique, 756
- alternateur, 1005
- ampèremètre, 247
- amplitude, 30, 62
- amplitude complexe, 341
- amplitude crête à crête, 338
- analyse spectrale, 49
- angles orientés, 1107
- apériodique, 319
- aplanétisme, 156
- approximation des régimes quasi-stationnaire, 255
- argument, 1111
- argument cosinus hyperbolique, 1082
- argument sinus hyperbolique, 1082
- ARQS, 255
- aspect corpusculaire, 207
- aspect ondulatoire, 207
- axe optique, 155
- bande passante, 398, 427
- base de projection fixe, 460
- base de projection mobile, 461
- base locale, 461
- base mobile, 461
- bras de levier, 694
- célérité, 58
- célérité de la lumière dans le vide, 120
- calorimètre, 835
- capacité, 257

- capacité thermique à pression constante, 826
- capacité thermique à volume constant, 768
- centre optique, 159
- champ, 929
- champ magnétique, 930
- champ scalaire, 1113
- changement d'état, 772
- charge électrique, 244
- climatiseur, 887
- coefficient d'auto-inductance, 971
- coefficient de corrélation, 1058
- coefficient de frottement fluide, 538
- composante continue, 50
- composante orthoradiale, 462
- composante radiale, 462
- composante sinusoïdale, 49
- composantes du vecteur, 1091
- condition initiale, 288, 289
- conditions aux limites, 94
- conditions de Gauss, 156
- conditions initiales, 25
- conduction thermique, 808
- conjugués, 156
- conservative, 244, 585
- constante de gravitation universelle, 531
- constante de temps, 287
- constante des gaz parfaits, 764
- convection thermique, 808
- convention générateur, 255
- convention récepteur, 255
- coordonnées cartésiennes, 459
- coordonnées polaires, 461
- coordonnées sphériques, 465
- cosinus, 1081
- cosinus hyperbolique, 1082
- coupe-bande, 399, 401
- couple de freinage, 699
- couple magnétique, 944, 945
- couple moteur, 699
- courant électrique, 245
- courants de Foucault, 983, 1004
- courbe d'ébullition, 779
- courbe de rosée, 780
- courbe de saturation, 779
- critique, 320
- cycle de Carnot, 891
- cycle moteur, 806
- cycle récepteur, 806
- décéléré, 488
- décade, 380
- décibel, 377, 378
- décrément logarithmique, 626
- déphasage, 34, 340
- déplacement élémentaire, 468
- dérivée partielle, 1113
- dérivateur, 386, 393
- dérive d'un potentiel, 585
- développement en série de Fourier, 50
- développements limités, 1085
- demi-grand axe, 1099
- demi-petit axe, 1099
- deuxième loi de Joule, 828
- deuxième principe, 852
- deuxième vitesse cosmique, 734
- diagramme d'Amagat, 765
- diagramme de Bode, 377
- diagramme de Clapeyron, 766
- diagramme des frigoristes, 900
- diagrammes de phases, 772
- diaphragme, 177
- différentielle de f , 1113
- diffusion Compton, 209
- dioptre, 151
- dioptrie, 161
- dipôle actif, 258
- dipôle passif, 256
- directe, 1093
- dispersion, 120
- distance focale image, 160
- distance focale objet, 160
- diviseur de courant, 262
- diviseur de tension, 261
- domaine audible, 55
- domaine visible, 121
- double périodicité spatio-temporelle, 62
- droite d'action d'une force, 694
- dualité onde-particule, 207
- durée d'exposition, 177
- effet de bord, 935

INDEX

- effet Joule, 266
- effet photoélectrique, 208
- efficacité, 889, 890
- ellipse, 1099
- en opposition phase, 35
- en phase, 34
- en quadrature de phase, 35
- enthalpie, 826
- enthalpie massique de changement d'état, 831
- enthalpie molaire de changement d'état, 831
- entrefer, 1016
- entropie, 852
- entropie massique de changement d'état, 858
- entropie molaire de changement d'état, 858
- erreur de mesure, 1048
- erreurs aléatoires, 1049
- erreurs systématiques, 1049
- exponentielle, 1080
- extérieur, 758

- facteur d'amortissement, 317
- facteur de qualité, 317, 349, 394
- faisceau laser, 126
- farad, 257
- ferromagnétique, 980
- feuille, 983
- filtre, 397
- fluide supercritique, 775
- flux magnétique, 958
- flux propre, 971
- focale, 177
- fonction d'état, 767
- fonction d'onde, 223
- fonction de transfert, 355
- fonctions de transfert asymptotiques, 380
- fondamental, 50
- force, 524
- force électromotrice induite, 962
- force centrale, 696
- force centrale conservative, 717
- force de freinage, 1004
- force de frottement fluide, 1002, 1009
- force de Lorentz, 651
- forces à distances, 529
- forces de contact, 529
- forces dissipatives, 589
- forces newtoniennes, 719
- forme d'onde, 37
- formule de Taylor, 1084
- formule des interférences, 81
- formules de conjugaison, 167
- foyer image principal, 158
- foyer objet principal, 158
- foyers image secondaires, 158
- foyers objet secondaires, 158
- fréquence, 31
- fréquence atténuée, 398
- fréquence passante, 398
- fréquences propres, 95
- franges rectilignes, 214

- générateur, 258
- générateur de tension réel, 259
- générateur idéal de tension, 258
- gain en décibel, 377
- gain maximum, 398
- gain minimum, 398
- gaz parfait, 764
- gaz parfait monoatomique, 768
- gaz réels, 766
- gradient, 1114
- grandeur d'influence, 1050
- grandeur massique, 761
- grandeur molaire, 761
- grandissement, 155
- grandissement linéaire, 157

- harmonique de rang n , 50
- haut-parleur, 1011
- henry, 258, 972, 977
- hyperbole, 1099
- hypermétrope, 177

- impédance, 344
- impédance cinétique, 1013
- inégalité de Clausius, 887
- incertitude élargie, 1052
- incertitude élargie relative, 1056
- incertitude de mesure, 1051
- incertitude-type, 1053

- incertitude-type de mesure., 1051
- incertitude-type relative, 1051
- indice optique, 120
- indirecte, 1093
- inductance, 258
- inductance mutuelle, 977
- inductance propre, 971
- induction électromagnétique, 961
- inertie, 521
- instable, 594
- intégrale première, 1078
- intégrale première du mouvement, 554, 590
- intégrateur, 382, 387, 393
- intensité, 246
- interaction coulombienne, 532
- interactions fondamentales, 529
- interférence constructive, 85
- interférence destructive, 85
- intervalle de confiance, 1051
- invariants dans le temps, 410
- isentropique, 853
- isobare, 796
- isochore, 796
- isochronisme, 29, 618
- isochronisme des petites oscillations, 553
- isolé, 523
- isotherme, 798
- la force de portance, 536
- la force de traînée, 536
- la poussée d'Archimède, 536
- lame semi-réfléchissante, 212
- lampe spectrale, 125
- lentille, 159
 - sphérique, 159
- lentille mince, 159
- lentilles convergentes, 159
- lentilles divergentes, 159
- liaison pivot, 700
- lignes de champ, 930
- liquide saturant seul, 780
- logarithme décimal, 1080
- logarithme népérien, 1080
- loi d'Ohm, 257
- loi de Faraday, 962
- loi de la propagation rectiligne de la lumière, 127
- loi de Laplace, 855
- loi de Lenz, 961, 1002
- Loi des aires, 720
- loi des mailles, 252
- loi des nœuds, 248
- Loi des orbites, 720
- Loi des périodes, 720
- lois de Descartes, 150
- lois de Descartes pour la réflexion, 150
- lois de Descartes pour la réfraction, 151
- longitudinale, 47
- longueur d'onde, 62, 63
- longueur d'onde de de Broglie, 215
- LTI systems, 410
- lumière blanche, 125
- lunette
 - astronomique, 174
 - de Galilée, 173
- mécaniquement réversible, 804
- méthode des moindres carrés, 1058
- métrologie, 1047
- machine frigorifique, 887
- maille, 252
- masse, 251
- masse inerte, 521
- masse inertielle, 521
- masse pesante, 521
- mesurage, 1048
- mesurande, 1048
- modes de transfert thermique, 808
- modes propres, 94, 95
- module, 1111
- Moment cinétique par rapport à un axe, 687
- Moment cinétique par rapport à un point, 686
- moment d'inertie, 689, 691
- moment magnétique, 937
- moteur thermique, 887
- mouvement circulaire et uniforme, 661
- mouvement rectiligne et uniformément accéléré, 481
- moyenneur, 421
- nœuds de vibration, 89, 90

INDEX

- Neumann, 977
- neutre, 253
- niveaux d'énergie, 227
- nombre d'onde, 62
- nombres complexes, 1111
- nœud, 248
- œil, 175
- ohm, 256
- onde, 47
- onde électromagnétique, 119
- onde de matière, 215
- onde harmonique, 62
- onde progressive, 60, 61
- onde progressive à une dimension, 58
- onde progressive sinusoïdale, 62
- onde sinusoïdale, 62
- onde stationnaire, 88
- ondes, 47
- ondes élastiques, 47
- ondes électromagnétiques, 48
- ondes acoustiques, 48
- ondes gravitationnelles, 48
- ondes monochromatiques, 120
- ondes sinusoïdales, 47
- opalescence critique, 782
- oscillateur harmonique, 23, 25
- période, 31
- périodique, 31
- parabole, 1099
- parabole de sûreté, 548
- paraxial, 156
- paroi diathermane, 811
- partie imaginaire, 1111
- partie réelle, 1111
- passe-bande, 399, 401
- passe-bas, 398, 401
- passe-haut, 398, 401
- passe-tout, 388
- pendule pesant, 702
- permittivité du vide, 533
- phénomène d'interférences, 83
- phénomène de Gibbs, 416
- phase, 253
- phase condensée indilatable et incompressible, 766
- phase initiale, 30
- phase initiale à l'origine, 62
- phase instantanée, 30
- phases, 755
- phases condensées, 755
- photon, 210
- photons, 207
- plan d'incidence, 150
- plan de phase, 596
- plan focal image, 158
- plan focal objet, 158
- poids, 531
- point critique, 774, 781
- point de fonctionnement, 264
- point image réel, 155
- point image virtuel, 154
- point objet réel, 153
- point objet virtuel, 154
- point triple, 774
- pompe à chaleur, 887
- portrait de phase, 291, 596
- potentiel électrique, 250, 656
- pouvoir séparateur, 176
- première loi de Joule, 769
- première vitesse cosmique, 734
- premier principe de la thermodynamique, 820
- pression, 759
- pression de vapeur saturante, 777
- primaire, 980
- primitive, 1085
- principe d'indétermination de Heisenberg, 225
- principe de complémentarité, 223
- produit scalaire, 1091
- produit vectoriel, 1094
- propagation, 47, 58
- propriété caractéristique, 90
- pseudo-isolés, 525
- pseudo-période, 626
- pseudopériode, 326
- pseudopériodique, 320
- puissance, 265, 578
- puits de potentiel, 597, 615
- puits infini à une dimension, 226

- pulsation, 30
- pulsation caractéristique, 317
- pulsation cyclotron, 661
- pulsation de résonance, 349
- pulsation propre, 25
- Punctum Proximum, 176
- Punctum Remotum, 176
- quantifiée, 226
- quantification, 226
- quantité de matière, 760
- quantité de mouvement, 522
- quantum d'énergie, 208
- quasistatique, 851
- référence de l'énergie potentielle, 585
- référentiel, 466
- référentiel du laboratoire, 524
- référentiel géocentrique, 524
- référentiel héliocentrique, 524
- référentiel terrestre, 524
- référentiels galiléens, 523, 524
- réflexion, 150
- réflexion totale, 152
- réfraction, 151
- régime établi, 286, 294, 314, 337
- régime aperiodique, 624, 627
- régime critique, 624
- régime libre, 292, 624
- régime permanent, 286, 294, 314, 337
- régime pseudo-périodique, 624, 626
- régime sinusoïdal forcé, 338
- régime transitoire, 286, 294, 314, 338
- régression linéaire, 1058
- réjecteur de bande, 399
- réponse indicielle, 286
- résistance, 256
- résistance d'entrée, 262, 263
- résistance de sortie, 263
- résonance, 91, 349, 351, 390, 395, 399
- réversible, 851
- règle de la main droite, 931, 933, 937, 958, 1093
- raideur du ressort, 535
- raies spectrales, 126
- rapport de transformation, 981
- rapport des capacités thermiques, 829
- rayonnement thermique, 808
- rayons lumineux, 127
- réflexion, 150
- relation de de Broglie, 215
- relation de Mayer, 829
- relation(s)
 - de Descartes, 167
- rendement, 887
- rendement de Carnot, 888
- RMS, 417
- rotation, 1096
- rotor, 699, 1015
- rupture de contact, 536
- série de Fourier, 412
- secondaire, 980
- signal complexe, 341
- signal physique, 23, 47
- signal sinusoïdal, 23, 30
- signaux périodiques, 23
- sinus, 1081
- sinus hyperbolique, 1082
- solénoïde, 934
- solide, 509
- solution homogène, 288, 301, 319
- solution particulière, 288, 301, 319
- source ponctuelle, 128
- source ponctuelle et monochromatique, 128
- sources de photons uniques, 211
- spectre, 47
- spectre continu, 53
- spectre d'énergie, 49
- spectre d'amplitude, 49
- spectre de phase, 49
- spectre du signal, 49
- stable, 594
- stator, 699, 1015
- stigmatisme, 156
- stroboscope, 92
- superposition linéaire, 96
- surface de contrôle, 758
- surface iso- f , 1113
- symétrie, 1097

INDEX

- synthèse de Fourier, 50
- système à un degré de liberté, 596
- système afocal, 159
- système de coordonnées cylindriques, 463
- système fermé, 759
- Système International d'unités, 1048
- système isolé, 820
- système optique centré, 155
- système ouvert, 759
- système thermodynamique, 758
- système thermodynamique à l'équilibre, 762
- systèmes diphasés, 772
- systèmes linéaires, 409

- tangente, 1081
- température, 760
- température d'ébullition, 778
- temps de réponse, 290, 294, 301
- temps de réponse à 5%, 322
- tension, 249, 534
- terre, 254
- Tesla, 931
- théorème de Carnot, 888–890
- thermostat, 810
- titres massiques, 775
- titres molaires, 776
- trajectoire, 468
- trajectoire de phase, 596
- transfert thermique, 807
- transformée de Fourier, 53
- transformateur, 980
- transformateur d'isolement, 983
- transformation irréversible, 849
- transformation thermodynamique, 795
- translation, 510, 1095
- translation circulaire, 512
- translation rectiligne, 511
- transmittance, 355
- transversale, 47
- travail, 799
- travail élémentaire, 579

- uniformément accéléré, 488
- uniformément décéléré, 488
- uniforme, 488

- valeur efficace, 417, 419
- valeur en eau du calorimètre, 835
- valeur mesurée, 1048
- valeur vraie, 1048
- vapeur saturante seule, 780
- variable
 - d'état, 759
- variable extensive, 760
- variable intensive, 761
- variables d'état, 759
- vecteur d'onde, 62
- vecteur déplacement, 468
- vecteur de Fresnel, 33
- vecteur surface, 937, 958
- vecteurs de Fresnel, 340
- ventres de vibration, 89, 90
- vergence, 161
- vitesse angulaire, 485
- vitesse aréolaire, 723
- vitesse de libération, 734
- vitesse de propagation, 58
- vitesse instantanée, 469
- vitesse moyenne, 469
- vitesse quadratique moyenne, 770
- voltmètre, 250

- weber, 959