

Séverine Bagard | Nicolas Simon

# VISA POUR LA PRÉPA 2021-2022

## PHYSIQUE-CHIMIE

MPSI - PCSI - MP2I - PTSI - TSI - BCPST

l'intelligence

DUNOD

Conception graphique de la couverture : Hokus Pokus Créations

© Dunod, 2021  
11 rue Paul Bert, 92240 Malakoff  
[www.dunod.com](http://www.dunod.com)  
ISBN 978-2-10-083034-3

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

|                                                                                                           |           |
|-----------------------------------------------------------------------------------------------------------|-----------|
| <b>1. Structure de la matière</b>                                                                         | <b>1</b>  |
| 1.1 Les atomes                                                                                            | 1         |
| 1.2 La liaison de covalence                                                                               | 9         |
| 1.3 Les associations d'atomes                                                                             | 10        |
| 1.4 Les molécules de la chimie organique                                                                  | 16        |
| 1.5 Spectres d'émission et d'absorption des atomes                                                        | 26        |
| 1.6 Analyse spectrale en chimie                                                                           | 29        |
| Exercices                                                                                                 | 32        |
| Solutions des exercices                                                                                   | 33        |
| <b>2. Méthodes physiques de détermination d'une quantité de matière</b>                                   | <b>39</b> |
| 2.1 Mesure de l'abondance d'un échantillon de matière                                                     | 39        |
| 2.2 Les méthodes physiques d'analyse d'une solution                                                       | 42        |
| Exercices                                                                                                 | 54        |
| Solutions des exercices                                                                                   | 56        |
| <b>3. Transformations chimiques en solutions aqueuses</b>                                                 | <b>61</b> |
| 3.1 Modélisation d'une transformation chimique                                                            | 61        |
| 3.2 Interprétation des équilibres chimiques et détermination de la composition d'un système à l'équilibre | 68        |

|            |                                |    |
|------------|--------------------------------|----|
| <b>3.3</b> | Équilibres acido-basiques      | 74 |
| <b>3.4</b> | Équilibres d'oxydoréduction    | 84 |
|            | <b>Exercices</b>               | 95 |
|            | <b>Solutions des exercices</b> | 97 |

## 4. Méthodes chimiques de détermination d'une quantité de matière 103

|            |                                |     |
|------------|--------------------------------|-----|
| <b>4.1</b> | Principe d'un titrage          | 103 |
| <b>4.2</b> | Méthodes de suivi d'un titrage | 108 |
|            | <b>Exercices</b>               | 119 |
|            | <b>Solutions des exercices</b> | 121 |

## 5. Cinétique d'une transformation chimique 127

|            |                                               |     |
|------------|-----------------------------------------------|-----|
| <b>5.1</b> | Facteurs cinétiques                           | 127 |
| <b>5.2</b> | Cinétique chimique : loi de vitesse d'ordre 1 | 129 |
|            | <b>Exercices</b>                              | 134 |
|            | <b>Solutions des exercices</b>                | 136 |

## 6. Les lois de Newton et leurs applications 141

|            |                                                      |     |
|------------|------------------------------------------------------|-----|
| <b>6.1</b> | La cinématique                                       | 141 |
| <b>6.2</b> | La dynamique et les lois de Newton                   | 146 |
| <b>6.3</b> | Les interactions                                     | 148 |
| <b>6.4</b> | Mouvement d'un point matériel dans un champ uniforme | 150 |
| <b>6.5</b> | Mécanique céleste                                    | 151 |
|            | <b>Exercices</b>                                     | 155 |
|            | <b>Solutions des exercices</b>                       | 156 |

## 7. Caractérisation des phénomènes ondulatoires 161

|            |                                                         |     |
|------------|---------------------------------------------------------|-----|
| <b>7.1</b> | Généralités sur les ondes mécaniques progressives (OMP) | 161 |
| <b>7.2</b> | Les phénomènes liés aux ondes                           | 168 |
|            | <b>Exercices</b>                                        | 180 |
|            | <b>Solutions des exercices</b>                          | 182 |

## 8. Introduction à la mécanique quantique 189

|            |                                                         |     |
|------------|---------------------------------------------------------|-----|
| <b>8.1</b> | La physique des quantas                                 | 189 |
| <b>8.2</b> | La structure de l'atome et l'interaction noyau-électron | 194 |
| <b>8.3</b> | La dualité onde-corpuscule                              | 196 |
|            | <b>Exercices</b>                                        | 198 |
|            | <b>Solutions des exercices</b>                          | 200 |

## 9. Optique géométrique 205

|            |                                |     |
|------------|--------------------------------|-----|
| <b>9.1</b> | Généralités sur la lumière     | 205 |
| <b>9.2</b> | Les lentilles minces           | 208 |
| <b>9.3</b> | Les instruments d'optique      | 217 |
|            | <b>Exercices</b>               | 222 |
|            | <b>Solutions des exercices</b> | 223 |

## 10. Électrocinétique 227

|             |                                                             |     |
|-------------|-------------------------------------------------------------|-----|
| <b>10.1</b> | Phénomènes, grandeurs et lois de base de l'électrocinétique | 227 |
| <b>10.2</b> | Les dipôles et leurs caractéristiques                       | 233 |
| <b>10.3</b> | Exemple d'étude d'un régime transitoire : cas du circuit RC | 239 |
|             | <b>Exercices</b>                                            | 245 |
|             | <b>Solutions des exercices</b>                              | 247 |

|                                                                  |            |
|------------------------------------------------------------------|------------|
| <b>11. Aspects énergétiques</b>                                  | <b>255</b> |
| <b>11.1</b> Énergie et puissance                                 | 255        |
| <b>11.2</b> Aspects énergétiques des phénomènes mécaniques       | 258        |
| <b>11.3</b> Aspects énergétiques des phénomènes électriques      | 264        |
| <b>11.4</b> Aspects énergétiques des phénomènes thermodynamiques | 269        |
| <b>11.5</b> Un mot concernant les gaz                            | 276        |
| <b>Exercices</b>                                                 | 279        |
| <b>Solutions des exercices</b>                                   | 281        |
| <b>Index</b>                                                     | <b>289</b> |

### 1.1 Les atomes

#### — Qu'est-ce qu'un élément chimique ?

Un **élément chimique** est l'ensemble des atomes possédant le même nombre de protons au noyau. Chaque élément sera donc caractérisé par ce nombre, appelé **numéro atomique** et noté **Z**.

**Remarque :** une transformation chimique repose sur un réarrangement des liaisons de covalence, qui n'engagent que les électrons. Les éléments chimiques sont donc conservés dans ce cas. Si les noyaux eux-mêmes sont affectés, on quitte le domaine de la chimie pour entrer dans celui de la physique nucléaire.

Par exemple, l'élément fer est contenu dans un morceau de fer solide (également dit *métallique*) ainsi que dans les ions fer (II)  $\text{Fe}^{2+}$  et fer (III)  $\text{Fe}^{3+}$ . Ainsi, quand on dit qu'il y a du fer dans les épinards, on entend bien sûr par là la présence, dans le légume, des ions et non de la forme métallique.

Chaque élément possède un nom et un symbole commençant par une lettre majuscule, éventuellement suivi d'une lettre minuscule. La lettre majuscule correspond le plus souvent à la première lettre du de l'élément (en Français souvent, mais également en Anglais, Danois, Latin, etc.). Par exemple, l'élément carbone a pour symbole C, l'élément chlore pour symbole Cl, et l'élément sodium Na (du latin *natrium*).

**Remarque :** les données simultanées du symbole, et du numéro atomique Z d'un élément sont redondantes.

#### — Qu'est-ce qu'un isotope ?

Des **isotopes** sont des atomes dont les noyaux possèdent des nombres identiques de protons, mais des nombres de neutrons différents. Ils appartiennent donc au même élément, mais possèdent des nombres de nucléons différents. Le nombre de nucléons d'un noyau est une autre donnée fréquemment utilisée ; il est appelé **nombre de masse** du noyau, noté **A**, et figure souvent en haut à gauche du noyau considéré. Par exemple, l'hydrogène  $1\text{ }^1_1\text{H}$ , le deutérium  $2\text{ }^2_1\text{H}$  et le tritium  $3\text{ }^3_1\text{H}$  sont trois isotopes de l'élément hydrogène.

**Remarque :** la dénomination *isotope* provient du Grec *ἴσος* (« le même ») et *τόπος* (« lieu »). En effet, deux isotopes appartenant à un même élément occuperont forcément la même case dans la **classification périodique des éléments**.

Deux isotopes diffèrent donc par leurs **nombres de masse A**. En conséquence, deux isotopes d'un même élément ont des masses molaires différentes, ce qui pose problème pour la définition de la masse molaire d'un élément chimique. Aussi cette dernière sera-t-elle déterminée en tenant compte des proportions respectives des différents isotopes de cet élément dans la nature (ceci explique le fait que la masse molaire d'un élément donné ne soit jamais exactement un nombre entier).

## — Comment peut-on caractériser l'état d'un électron dans un atome ?

À l'échelle atomique, la mécanique newtonienne ne suffit plus pour décrire convenablement le comportement dynamique des particules. En particulier, une analyse fine révèle que l'énergie des électrons peut adopter uniquement certaines valeurs bien particulières, et non varier continument comme c'est le cas dans le modèle newtonien : on dit que leur énergie est **quantifiée**. La mécanique quantique est un modèle qui rend compte de cette quantification, ainsi que de celle d'autres grandeurs (évoquées ci-dessous à titre d'information et dont vous ferez la connaissance en CPGE).

À chacune de ces grandeurs quantifiées est alors associé un nombre quantique :

| Grandeur quantifiée   | Nombre quantique (NQ) associé | Notation | Valeurs                               | Structure associée |
|-----------------------|-------------------------------|----------|---------------------------------------|--------------------|
| Énergie               | NQ principal                  | $n$      | $n = 0, 1, 2, \dots$                  | Niveau             |
| Moment cinétique (MC) | NQ secondaire                 | $l$      | $l = 0, 1, 2, \dots, (n - 1)$         | Orbitale           |
| Projection du MC      | NQ magnétique                 | $m$      | $m = -l, \dots, 0, \dots, (l - 1), l$ | Case quantique     |
| MC intrinsèque        | NQ de spin                    | $s$      | $s = \pm 1 / 2$                       | État quantique     |

Ainsi un électron dans un atome ne peut-il exister que dans certains états, chacun défini par un jeu de valeurs de ces nombres. Outre les valeurs effectives des grandeurs physiques qui leur sont associées (énergie, moment cinétique, etc.), ces nombres conditionnent la fonction décrivant la probabilité de présence de l'électron, autrement dit la forme du fameux **nuage électronique**.

Une règle, appelée **principe d'exclusion de Pauli**, affirme alors qu'au sein d'un même atome, deux électrons ne peuvent exister dans le même état quantique, c'est-à-dire qu'il est impossible pour deux électrons d'un même atome, de posséder exactement le même jeu de nombres quantiques.

## — Comment le cortège électronique d'un atome est-il structuré ?

Le principe d'exclusion de Pauli fait que les électrons d'un atome vont tous occuper un état différent. Pour un jeu de valeurs  $(n, l, m)$  donné, il existe donc 2 valeurs du nombre de spin  $s$ . On appelle ainsi **case quantique** la donnée d'un triplet  $(n, l, m)$ , et un atome peut donc contenir **2 électrons par case quantique**. Cette association explique notamment la formation des **doublets d'électrons**.

Par ailleurs, pour un doublet de valeurs  $(n, l)$ , on sait qu'il existe  $(2l + 1)$  valeurs possibles du nombre  $m$  (toutes les valeurs entières de  $-l$  à  $+l$ ). On appelle alors **orbitale atomique** la donnée d'un doublet  $(n, l)$ , et un atome va donc présenter  **$(2l + 1)$  cases quantiques par orbitale**, soit un total de  $(4l + 2)$  électrons pour une orbitale de nombre quantique  $l$ .

À une valeur donnée de  $l$  peut être associée une forme donnée de nuage électronique, indépendamment de la valeur de  $n$ . On définit ainsi différents types d'orbitales, désignés par des lettres de l'alphabet :

| Valeur de $l$                                    | 0   | 1   | 2   | 3   |
|--------------------------------------------------|-----|-----|-----|-----|
| Type d'orbitale                                  | $s$ | $p$ | $d$ | $f$ |
| Nombre de cases quantiques                       | 1   | 3   | 5   | 7   |
| Peut accueillir un nombre total de ... électrons | 2   | 6   | 10  | 14  |

**Remarque : attention** à ne pas confondre les nombres quantiques ( $n, l, m, s$ ) avec les noms des orbitales atomiques ( $s, p, d, f$ , etc.). Les deux concernent la mécanique quantique mais décrivent des concepts très différents (nombres quantiques pour les premiers, type d'orbitale pour les seconds).

Enfin pour une valeur donnée de  $n$ , le nombre  $l$  peut prendre  $n$  valeurs (valeurs entières de 0 à  $(n - 1)$ ). On appelle **niveau d'énergie** de l'atome la donnée d'une valeur de  $n$ , et un atome va donc disposer de  **$n$  orbitales atomiques au sein du niveau d'énergie  $n$** .

| $l$<br>$n$ | 0<br>( $2 e^-$<br>par<br>orbitale $s$ ) | 1<br>( $6 e^-$<br>par<br>orbitale $p$ ) | 2<br>( $10 e^-$<br>par<br>orbitale $d$ ) | 3<br>( $14 e^-$<br>par<br>orbitale $f$ ) | $l$<br>( $(4l + 2) e^-$<br>par<br>orbitale) | $N_{e^-, \max}$<br>sur le<br>niveau $n$ |
|------------|-----------------------------------------|-----------------------------------------|------------------------------------------|------------------------------------------|---------------------------------------------|-----------------------------------------|
| $n$        | $ns$                                    | $np$                                    | $nd$                                     | $nf$                                     | ...                                         | $2n^2$                                  |
| $\vdots$   | $\vdots$                                | $\vdots$                                | $\vdots$                                 | $\vdots$                                 | ...                                         | $\vdots$                                |
| 4          | $4s$                                    | $4p$                                    | $4d$                                     | $4f$                                     |                                             | 32                                      |
| 3          | $3s$                                    | $3p$                                    | $3d$                                     |                                          |                                             | 18                                      |
| 2          | $2s$                                    | $2p$                                    |                                          |                                          |                                             | 8                                       |
| 1          | $1s$                                    |                                         |                                          |                                          |                                             | 2                                       |

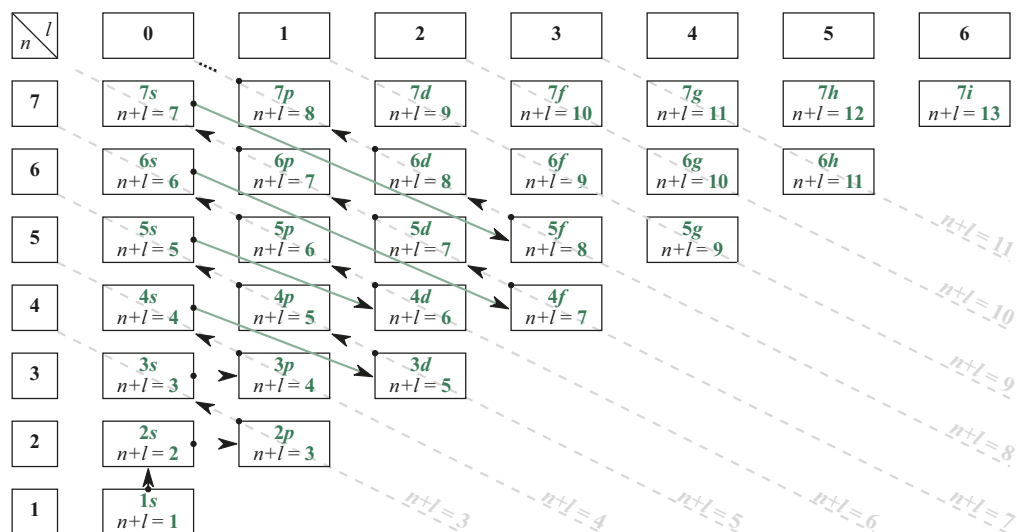
**Remarque :** le nombre de niveaux et d'orbitales est théoriquement infini. Cependant à mesure qu'il croît, le noyau gagne en instabilité : tous les atomes dont le numéro atomique excède 83 sont radioactifs, avec globalement une espérance de vie d'autant plus courte que leur numéro atomique est plus grand. C'est la raison pour laquelle le tableau de Mendeleïev se limite à environ 120 éléments, limite au-delà de laquelle les noyaux se désintègrent spontanément en une fraction de seconde. Il existe donc des orbitales situées au-delà des orbitales  $f$  et nommées en déroulant l'alphabet (orbitales  $g, h$ , etc.) cependant vous n'aurez guère d'occasion de les rencontrer avant quelques années, et encore seulement dans certains domaines très spécialisés.

## — Dans quel ordre les électrons d'un atome se répartissent-ils sur ses orbitales ?

Lorsqu'un atome se trouve dans son état fondamental (aucun électron excité), ses électrons vont occuper les différents états possibles en suivant la **règle de Klechkowski** :

- les électrons remplissent les orbitales atomiques selon les valeurs croissantes de  $(n + l)$  ;
- si deux orbitales présentent la même valeur de  $(n + l)$ , c'est celle de nombre  $n$  le plus bas qui se remplit en premier.

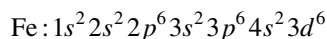
On obtient ainsi :



## — Qu'est-ce que la structure électronique d'un atome ?

On appelle ainsi la donnée de la répartition des électrons du cortège électronique sur les différentes orbitales. Elle s'écrit en énumérant lesdites orbitales dans l'ordre décrit par la règle de Klechkowski, et en précisant pour chacune le nombre d'électrons qu'elle contient, indiqué en exposant.

Par exemple, si l'on considère l'atome neutre de fer ( $Z = 26$ , soit 26 électrons pour un atome neutre), on écrira :



Notons que la structure électronique d'un atome ou d'un ion dépend uniquement du nombre d'électrons que comporte son cortège électronique. Ainsi par exemple un ion  $\text{Ni}^{2+}$  ( $Z = 28$ , soit 26 électrons pour un cation chargé deux fois) a-t-il la même structure électronique que l'atome neutre de fer, même si les énergies associées à leurs orbitales respectives diffèrent.

## — Qu'appelle-t-on électrons de valence, ou électrons externes ?

On appelle ainsi les électrons situés sur la **couche de plus haut nombre quantique principal  $n$**  qui en contienne encore. Attention : si pour les 18 premiers électrons cette couche est toujours la dernière figurant dans l'écriture de la structure électronique (jusqu'à  $1s^2 2s^2 2p^6 3s^2 3p^6$ , donc), les choses sont un peu moins évidentes au-delà. Par exemple, entre 21 et 30 électrons, les derniers placés par la règle de Klechkowski appartiennent aux orbitales 3d (sur la couche  $n = 3$ ). Cependant l'orbitale 4s ( $n = 4$ ) ayant déjà accueilli des électrons, ce sera bien elle la couche externe, malgré le fait qu'elle n'ait pas été la dernière remplie.

Ces électrons de valence présentent une importance toute particulière puisque c'est sur eux que reposent les liaisons de covalence (d'où leur nom) ; leur nombre définit donc la **valence d'un élément**.

## — Qu'est-ce que la valence d'un élément ?

On appelle ainsi le nombre de doublets liants qu'un atome de cet élément est capable d'établir avec d'autres atomes.

On constate alors que la valence d'un atome neutre est directement liée à son numéro atomique, puisque ce dernier conditionne le nombre d'électrons de cet atome, et que la répartition de ceux-ci sur les orbitales, imposée par la nature, définit à son tour le nombre d'électrons externes.

Ainsi :

- la valence de l'hydrogène est limitée à 1 puisque sa couche externe ( $n = 1$ ) ne peut accueillir que 2 électrons (soit un doublet) ;
- pour les atomes neutres jusqu'à  $Z = 18$  (dont font partie C, N et O), la couche externe ( $n = 2$ ) peut accueillir au plus 8 électrons, soit 4 doublets, d'où une valence maximale égale à 4 ;

| Atome                       | H      | C                | N                | O                | F                |
|-----------------------------|--------|------------------|------------------|------------------|------------------|
| Numéro atomique             | 1      | 6                | 7                | 8                | 9                |
| Structure électronique      | $1s^1$ | $1s^2 2s^2 2p^2$ | $1s^2 2s^2 2p^3$ | $1s^2 2s^2 2p^4$ | $1s^2 2s^2 2p^5$ |
| Nombre d'électrons externes | 1      | 4                | 5                | 6                | 7                |
| Valence                     | 1      | 4                | 3                | 2                | 1                |

- au-delà, la structure à 8 électrons externes reste une structure favorisée par la nature (règle de l'octet, *cf.* plus loin), mais la valence peut monter plus haut et vous découvrirez en CPGE les merveilles de l'hypervalence (possibilité pour un atome de se lier plus de 4 fois).

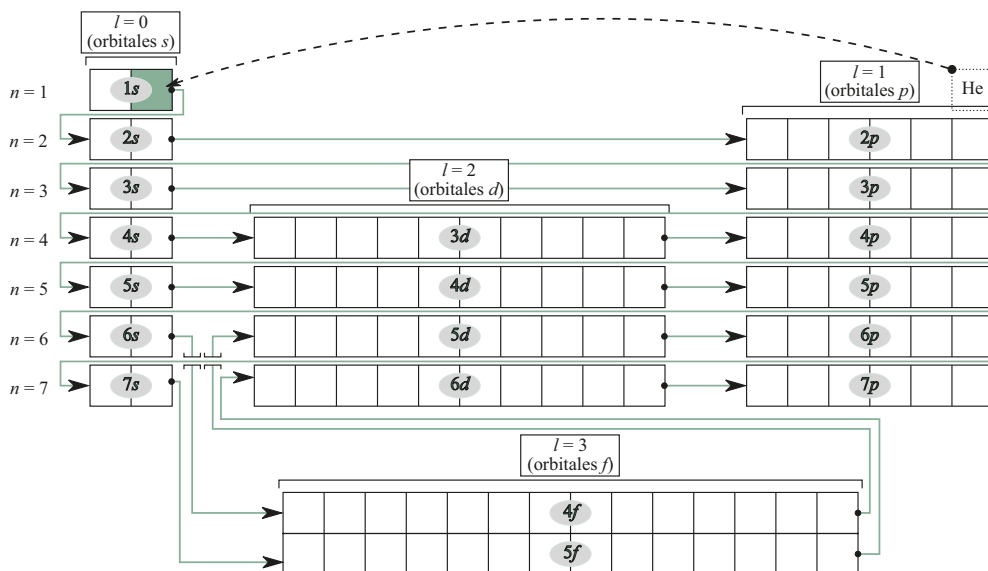
## — Existe-t-il un lien entre la règle de Klechkowski et la structure du tableau de Mendeleïev ?

Et pas qu'un peu... Mendeleïev a établi son tableau sur une base purement expérimentale, admettant avec une humilité rare ses diverses insuffisances :

- exceptions à la règle de remplissage selon les masses atomiques croissantes (puisque c'est en réalité le numéro atomique qui doit être pris en considération, mais le concept de proton est encore lointain en 1869) ;
- cases vides (correspondant à des éléments inconnus à l'époque, et dont il prédira l'existence ainsi que les propriétés chimiques par analogie avec les éléments de la même famille) ;
- impossibilité d'expliquer cette forme singulière, dans laquelle s'insèrent régulièrement de nouvelles rangées d'éléments à mesure que l'on descend dans les lignes.

Or ce troisième point doit justement tout à la règle de remplissage des orbitales atomiques, et la classification de Mendeleïev constituera l'un des guide-lignes qui permettront d'élaborer la mécanique quantique.

Une fois connue la structure du cortège électronique, il devient assez évident que la structure du tableau répond en fait au remplissage des orbitales pour des éléments dont les atomes neutres possèdent de plus en plus d'électrons :



Notons ici que l'hélium ( $Z = 2$ ) correspond bien au remplissage de l'orbitale  $1s$ . Cependant ses propriétés (stabilité chimique exceptionnelle entre autres), dues à la satisfaction de la règle du duet (analogue à la règle de l'octet, mais lorsque la couche externe est  $n = 1$ ), font de lui un gaz noble. C'est à ce titre que Ménédeleïev, fondant son tableau sur des critères expérimentaux, le plaça dans la 18<sup>e</sup> colonne et non dans la 2<sup>e</sup> comme l'exigerait le point de vue des orbitales.

## — Comment la classification périodique est-elle organisée ?

La classification se présente sous la forme d'un tableau à **sept lignes** (également appelées **périodes**) et **dix-huit colonnes** (cf. page suivante). À chaque colonne correspond une **famille** au sein de laquelle les éléments présentent tous un même type de réactivité. En effet, tous les éléments d'une même famille ont la même structure électronique externe. Les lignes, ou périodes, possèdent un nombre d'éléments croissants au fur et à mesure que l'on descend dans le tableau : il s'agit donc d'une classification périodique de période variable.

**Remarque :** deux lignes de 14 éléments chacune figurent sous la partie principale du tableau. Elles correspondent à des éléments appelés respectivement les **lanthanides** et les **actinides**, des noms de leurs premiers éléments respectifs. À noter que le lanthane et l'actinium font partie du bloc f mais sont en général représentés dans le corps principal du tableau (donc dans le bloc d), tandis que le cérium et le thorium (qui font partie du bloc d), premiers éléments de leurs blocs d respectifs, sont représentés dans ces deux lignes affichées en marge du tableau.

## — Quelles sont les principales familles de la classification à connaître ?

Les éléments de la **première colonne** (hormis l'hydrogène dont la réactivité est très particulière) constituent la famille des **alcalins**. Ces éléments sont des métaux ductiles

et très réducteurs. Leurs atomes s'oxydent très facilement en un cation monoatomique chargé une fois.

Les éléments de la **deuxième colonne** constituent la famille des **alcalino-terreux**. Ce sont également des métaux, bon réducteurs, qui vont facilement donner un cation monoatomique chargé deux fois.

Les éléments de la **dix-septième colonne** constituent la famille des **halogènes**. Les corps purs simples associés à ces éléments sont des molécules diatomiques constituant de bons oxydants. Les anions monoatomiques associés à ces éléments, chargés une fois négativement, donnent par réaction avec les ions argent (I) des solides peu solubles.

Les éléments de la **dix-huitième colonne** constituent la famille des **gaz nobles**. Dans les conditions usuelles de température et de pression ce sont des gaz monoatomiques avec une grande inertie chimique.

**Remarque :** plus anciennement, les éléments de la dix-huitième colonne étaient également appelés gaz rares, ou encore gaz inertes.

## — Quelles grandes tendances peut-on dégager de la classification ?

Plusieurs grandes **tendances** peuvent être dégagées de la lecture de la classification périodique des éléments :

- Les trois quarts des éléments sont des **métaux**, les **non-métaux** se situent dans la partie la plus à droite de la classification (cases les plus foncées).
- L'**électronégativité** est une grandeur sans dimension traduisant l'aptitude d'un atome à attirer à lui les électrons de toute liaison covalente dans laquelle il se trouve engagé. Elle croît de la gauche vers la droite et du bas vers le haut de la classification périodique. L'élément le plus électronégatif est le fluor (premier membre de la famille des halogènes), et le deuxième, l'oxygène. Notons que l'électronégativité nécessitant l'établissement d'une liaison de covalence pour être établie, elle n'est pas définie pour les gaz nobles.

## — Doit-on connaître par cœur toute la classification périodique des éléments ?

La réponse est évidemment non. Néanmoins, en CPGE, on attendra de vous une connaissance partielle de la classification. En particulier, vous devez être capable de retrouver :

- les trois premières lignes ;
- les trois premiers éléments de la famille des alcalins et de celle des alcalino-terreux ;
- les quatre premiers éléments de la famille des halogènes et de celle des gaz nobles.

## — En quoi les règles de l'octet et du duet consistent-elles ?

Les **règles du duet et de l'octet** permettent de prévoir comment des atomes isolés vont se transformer en ions monoatomiques, ou bien s'associer en molécules ou ions polyatomiques. Ces règles énoncent que les atomes ont tendance à acquérir une structure électronique externe identique à celle du gaz noble dont ils sont le plus proche dans la classification périodique.

|                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                  |                   |                  |                  |                 |                 |                  |                  |
|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|------------------|-------------------|------------------|------------------|-----------------|-----------------|------------------|------------------|
| 1,0<br>H<br>1     |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                  | 4,0<br>He<br>2    |                  |                  |                 |                 |                  |                  |
| 6,9<br>Li<br>3    | 9,0<br>Be<br>4    |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                  |                   | 12,0<br>C<br>6   | 14,0<br>N<br>7   | 16,0<br>O<br>8  | 19,0<br>F<br>9  | 20,2<br>Ne<br>10 |                  |
| 23,0<br>Na<br>11  | 24,3<br>Mg<br>12  |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                  |                   | 27,0<br>Al<br>13 | 28,1<br>Si<br>14 | 31,0<br>P<br>15 | 32,1<br>S<br>16 | 35,5<br>Cl<br>17 | 39,9<br>Ar<br>18 |
| 39,10<br>K<br>19  | 40,1<br>Ca<br>20  | 45,0<br>Sc<br>21  | 47,9<br>Ti<br>22  | 50,9<br>V<br>23   | 52,00<br>Cr<br>24 | 54,9<br>Mn<br>25  | 55,8<br>Fe<br>26  | 58,9<br>Co<br>27  | 58,7<br>Ni<br>28  | 63,5<br>Cu<br>29  | 65,4<br>Zn<br>30  | 69,7<br>Ga<br>31  | 72,6<br>Ge<br>32  | 74,9<br>As<br>33  | 79,0<br>Se<br>34  | 79,9<br>Br<br>35 | 83,8<br>Kr<br>36  |                  |                  |                 |                 |                  |                  |
| 85,5<br>Rb<br>37  | 87,6<br>Sr<br>38  | 88,9<br>Y<br>39   | 91,2<br>Zr<br>40  | 92,9<br>Nb<br>41  | 95,9<br>Mo<br>42  | 98,9<br>Tc<br>43  | 101,1<br>Ru<br>44 | 102,9<br>Rh<br>45 | 106,4<br>Pd<br>46 | 107,9<br>Ag<br>47 | 112,4<br>Cd<br>48 | 114,8<br>In<br>49 | 118,7<br>Sn<br>50 | 121,8<br>Sb<br>51 | 127,6<br>Te<br>52 | 129,9<br>I<br>53 | 131,3<br>Xe<br>54 |                  |                  |                 |                 |                  |                  |
| 132,9<br>Cs<br>55 | 137,3<br>Ba<br>56 | 57-71             | 178,5<br>Hf<br>72 | 180,9<br>Ta<br>73 | 183,8<br>W<br>74  | 186,2<br>Re<br>75 | 190,2<br>Os<br>76 | 192,2<br>Ir<br>77 | 195,1<br>Pt<br>78 | 197,0<br>Au<br>79 | 200,6<br>Hg<br>80 | 204,4<br>Tl<br>81 | 207,2<br>Pb<br>82 | 209,0<br>Bi<br>83 | 209<br>Po<br>84   | 210<br>At<br>85  | 222<br>Rn<br>86   |                  |                  |                 |                 |                  |                  |
| 223<br>Fr<br>87   | 226<br>Ra<br>88   | 89-103            | 261<br>Rf<br>104  | 262<br>Db<br>105  | 263<br>Sg<br>106  |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                  |                   |                  |                  |                 |                 |                  |                  |
| Lanthanides:      |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                  |                   |                  |                  |                 |                 |                  |                  |
| 138,9<br>La<br>57 | 140,1<br>Ce<br>58 | 140,1<br>Pr<br>59 | 144,2<br>Nd<br>60 | 145<br>Pm<br>61   | 150,4<br>Sm<br>62 | 152,0<br>Eu<br>63 | 157,3<br>Gd<br>64 | 158,9<br>Tb<br>65 | 162,5<br>Dy<br>66 | 164,9<br>Ho<br>67 | 167,3<br>Er<br>68 | 168,9<br>Tm<br>69 | 173,0<br>Yb<br>70 | 175,0<br>Lu<br>71 |                   |                  |                   |                  |                  |                 |                 |                  |                  |
| Actinides:        |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                   |                  |                   |                  |                  |                 |                 |                  |                  |
| 227<br>Ac<br>89   | 232,0<br>Th<br>90 | 231,0<br>Pa<br>91 | 238,0<br>U<br>92  | 237<br>Np<br>93   | 244<br>Pu<br>94   | 243<br>Am<br>95   | 247<br>Cm<br>96   | 247<br>Bk<br>97   | 251<br>Cf<br>98   | 252<br>Es<br>99   | 257<br>Fm<br>100  | 258<br>Md<br>101  | 259<br>No<br>102  | 262<br>Lr<br>103  |                   |                  |                   |                  |                  |                 |                 |                  |                  |

Actinides:

Les atomes dont le numéro atomique est voisin de celui de l'hélium ( $Z = 2$  ; en pratique les atomes de numéro atomique inférieur ou égal à 4) vont avoir tendance à acquérir une structure électronique externe en  $1s^2$  (donc à 2 électrons externes), d'où le nom de règle du duet. Au-delà, ils vont avoir tendance à acquérir une structure électronique à 8 électrons externes (d'où le nom de *règle de l'octet*), soit donc en  $ns^2np^6$ .

## — Comment prévoir la nature de l'ion monoatomique que peut éventuellement former un élément ?

Les atomes d'un élément participent à des édifices ioniques ou moléculaires, ou encore donnent des ions monoatomiques afin d'acquérir la **structure électronique externe** du gaz noble le plus proche d'eux dans la classification périodique. Une manière de réaliser cela est donc la formation d'ions monoatomiques en acquérant ou en perdant un, deux ou trois électrons de valence.

Les atomes des éléments des colonnes I, II et III perdent leurs électrons externes pour mener à des cations chargés respectivement une, deux ou trois fois.

Les atomes des éléments des colonnes XV, XVI et XVII gagnent respectivement trois, deux ou un électron(s) supplémentaire(s) pour mener à des anions chargés respectivement trois, deux ou une fois.

## 1.2 La liaison de covalence

### — Qu'est-ce qu'une liaison de covalence ?

Une **liaison de covalence** consiste en la mise en commun, par deux atomes, d'un électron externe chacun, formant ainsi un **doublet liant d'électrons externes**.

Ces deux électrons mis en commun peuvent provenir :

- soit de chacun des deux atomes ;
- soit d'un seul des deux atomes engagés dans la liaison.

**Remarque :** la provenance exacte des électrons formant une liaison de covalence sera détaillée en chimie organique, lors de l'étude des mécanismes réactionnels.

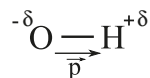
Les liaisons de covalence établies entre deux atomes peuvent être **multiples** ; en pratique, on rencontre des **double** et des **triple** liaisons correspondant respectivement à la mise en commun par les deux atomes de deux et trois doublets d'électrons externes.

**Remarque :** dans une liaison multiple, les doublets mis en commun ne sont pas tous de même nature. Le premier doublet liant mène à une liaison dite « sigma », alors que le (ou les) suivant(s) mène(nt) à une (des) liaison(s) « pi ». Une liaison sigma est plus difficile à rompre qu'une liaison pi. Cette différenciation sera elle aussi détaillée en CPGE.

### — Qu'est-ce qu'une liaison polarisée ?

Une liaison est dite **polarisée** lorsque l'un des atomes entre lesquels elle est engagée attire à lui le(s) doublet(s) d'électrons de ladite liaison, plus que l'autre. On sera donc en présence d'une liaison polarisée dès lors qu'il existe une différence significative d'électronégativité entre les atomes qui la forment.

L'atome le plus électronégatif acquiert alors une **charge partielle négative** notée  $-\delta$ , l'atome le moins électronégatif une **charge partielle positive** notée  $+\delta$ . Par exemple, une liaison O-H est polarisée ce qui pourra se représenter comme suit :



**Remarque :** la liaison restant tout de même de nature covalente, même si un atome semble accaparer le doublet liant, les charges partielles qui apparaissent restent forcément inférieures en valeur absolue à la charge électrique élémentaire  $e$ .

On caractérise une liaison polarisée par un vecteur appelé **moment dipolaire électrique**. Ce vecteur généralement noté  $\vec{p}$  présente les caractéristiques suivantes :

- il est dirigé parallèlement à la liaison ;
- il est orienté de l'atome le plus électronégatif (porteur de la charge partielle négative) vers le moins électronégatif (porteur de la charge partielle positive) ;
- il a pour valeur le produit de la valeur absolue de la charge partielle apparaissant sur chacun des deux atomes, par la longueur  $l$  de la liaison engagée par ceux-ci.

$$p = \delta \times l$$

|          |        |
|----------|--------|
| $p$      | en C.m |
| $\delta$ | en C   |
| $l$      | en m   |

## 1.3 Les associations d'atomes

### — Comment détermine-t-on le nombre de doublets d'électrons externes présents dans une molécule ou un ion polyatomique ?

Nous l'avons dit, les liaisons de covalence sont l'apanage des électrons externes des atomes. Il est donc intéressant, pour évaluer comment les atomes peuvent s'organiser en une molécule, d'anticiper le nombre total de ces doublets. Pour ce faire, on applique la méthode suivante :

- on détermine le **nombre d'électrons externes de chaque atome** en écrivant sa structure électronique, et on en déduit le nombre total d'électrons externes que peut fournir l'ensemble des atomes de l'entité (molécule ou ion polyatomique) considérée ;
- on soustrait à ce dernier un nombre d'électrons égal à sa charge globale algébrique (traduisant une diminution de ce nombre pour un cation, aucun changement pour une entité neutre, et une augmentation pour un anion) ;
- si le nombre total  $N$  d'électrons ainsi calculé est pair, on a alors  $\frac{N}{2}$  doublets d'électrons externes à répartir dans l'entité ;
- si  $N$  est impair, on a  $\frac{N-1}{2}$  doublets, plus un électron célibataire à répartir dans l'entité.

**Remarque :** lors de vos études secondaires, les entités polyatomiques que vous avez étudiées avaient systématiquement un nombre pair d'électrons externes. Ce ne sera plus forcément le cas en CPGE, où vous aboutirez ainsi parfois à ce que l'on appelle un « radical », très réactif en raison de la présence de cet électron célibataire. Un radical contenant un nombre impair d'électrons, il est donc bien entendu qu'il comportera (au moins) un atome ne satisfaisant pas à la règle de l'octet.

## — Comment établit-on la représentation de Lewis d'une molécule ou d'un ion polyatomique ?

Il n'y a pas de méthode universelle, la technique vient essentiellement en s'entraînant. On peut néanmoins dégager une méthodologie pouvant vous guider dans l'écriture de la **structure de Lewis** d'entités simples :

- on détermine la **covalence**, c'est-à-dire le nombre de liaisons de covalence que va naturellement établir chaque atome participant à l'édifice. Cette covalence est égale à 8 moins le nombre d'électrons externes de l'atome. Pour les atomes de la 2<sup>e</sup> ligne du tableau de Mendeleïev (typiquement *C, N, O*), cette covalence est égale à 8 moins le nombre d'électrons externes. Ainsi le carbone, possédant 4 électrons externes, est tétravalent, l'oxygène, possédant 6 électrons externes, est bivalent... Cette règle vaut également dans les lignes situées plus bas, tout en sachant qu'à partir de la 3<sup>e</sup> ligne, la disponibilité d'orbitales au-delà de l'orbitale *p* rend possible des *promotions de valence* pouvant aboutir à des valences supérieures à 4, comme vous le verrez en CPGE ;
- on lie les atomes par des doublets liants en respectant leur covalence (ce qui n'exclut pas de réaliser des liaisons multiples) ;
- on affecte les doublets restants en tant que doublets non liants afin de satisfaire à la règle de l'octet (pour les éléments des deux premières lignes de la classification au moins).

**Remarque :** dans le décompte des électrons pour chaque atome, il importe de bien distinguer les électrons externes dont l'atome est entouré (chaque doublet au contact d'un atome « entoure » celui-ci de ses 2 électrons), du nombre d'électrons externes qu'il possède (un doublet non liant est possédé en totalité, un doublet liant ne lui appartient que pour moitié). Le premier nombre intervient dans la satisfaction de la règle de l'octet, le second dans l'électroneutralité de l'atome.

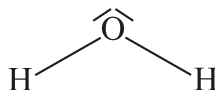
Se pose alors la question de la localisation de la (ou des) charge(s) d'un ion polyatomique. Pour ce faire, on doit introduire la notion de **charge formelle**. On la calcule pour chaque atome impliqué dans l'édifice en soustrayant le nombre d'électrons externes qu'il possède effectivement, de celui que possède un atome neutre du même élément.

Lorsqu'une charge formelle non nulle apparaît, on l'écrit à côté de l'atome concerné et on l'entoure (afin de ne pas la confondre avec la charge globale de l'entité).

**Remarque :** des charges formelles peuvent apparaître également dans des molécules malgré leur neutralité globale. On vérifiera à ce moment que les différentes charges formelles apparues se compensent bien.

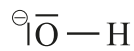
Illustrons cette méthodologie par deux exemples.

- La molécule d'eau, de formule brute  $\text{H}_2\text{O}$ , comporte 8 électrons externes en tout (1 pour chaque atome d'hydrogène et 6 pour celui d'oxygène) ; on doit donc répartir en tout  $8/2 = 4$  doublets pour construire la molécule. L'hydrogène étant monovalent, il paraît alors légitime d'établir deux liaisons simples entre l'atome d'oxygène et chacun des deux atomes d'hydrogène présents. L'hydrogène satisfait alors d'ores et déjà à la règle du duet, il ne reste plus qu'à placer les deux derniers doublets en doublets non liants sur l'atome d'oxygène qui satisfait alors à la règle de l'octet.



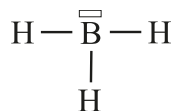
- L'ion hydroxyde  $\text{HO}^-$  possède également 8 électrons externes en tout (1 venant de l'atome d'hydrogène, 6 venant de l'atome d'oxygène et un dernier en raison de la charge globale

de l'anion). On répartit ces électrons en 4 doublets, dont un liant entre les deux atomes. On place alors les trois doublets restants en doublets non liants autour de l'atome d'oxygène. On détermine enfin les charges formelles. L'atome d'hydrogène, qui possède normalement un électron de valence, est entouré d'un doublet liant et possède donc 1 électron ; sa charge formelle est donc nulle. L'atome neutre d'oxygène possédant 6 électrons de valence, celui considéré dans le cas présent en possède pour sa part 7 (6 provenant des doublets non liants et 1 de sa liaison simple avec H) ; sa charge formelle est donc égale à  $6 - 7 = -1$ .



**Remarque :** vous remarquerez, en résolvant les exercices, qu'il est parfois possible d'écrire plusieurs répartitions de charge pour une même structure de Lewis. Ceci s'expliquera en CPGE par la théorie de la **mésomérie** qui traduit des déplacements d'électrons entre sites dits conjugués. Toutes les formes de Lewis correspondantes sont correctes mais diffèrent par leur probabilité de réalisation effective, la forme réelle de la molécule consistant en une moyenne de ces différentes formes, chacune pondérée par la probabilité en question.

La règle de l'octet a également ses limites, dans le sens où il n'y a parfois pas assez d'électrons pour la satisfaire ! Illustrons ce propos par un autre exemple. Le bore B appartient à la 13<sup>e</sup> colonne de la classification ; il possède donc 3 électrons externes sur sa couche de valence (2 de la couche 2s et 1 de la couche 2p). Associé à trois atomes d'hydrogène (monovalents), la molécule de borane  $\text{BH}_3$  va donc afficher 6 électrons externes soit trois doublets en tout. La règle de l'octet ne saurait donc être atteinte pour l'atome de bore, faute d'électrons ! On traduit cet état de fait sur le schéma de Lewis de la molécule en faisant explicitement apparaître ce manque d'électrons par une **lacune électronique**, schématisée par un rectangle vide :



## — Qu'est-ce que la méthode VSEPR ?

La **méthode VSEPR** (*Valence Shell Electron Pair Repulsion*), encore appelée **méthode de Gillespie**, explique la géométrie des molécules en se fondant sur la répulsion qu'exercent les uns sur les autres les différents doublets d'électrons externes présents dans un même édifice polyatomique (tous portent une charge électrique de même signe). La géométrie qui sera finalement adoptée par l'édifice sera alors celle qui permet d'écarter ces doublets au maximum les uns des autres (afin de minimiser leur répulsion électrostatique).

## — Quelles sont les principales géométries à savoir reconnaître ?

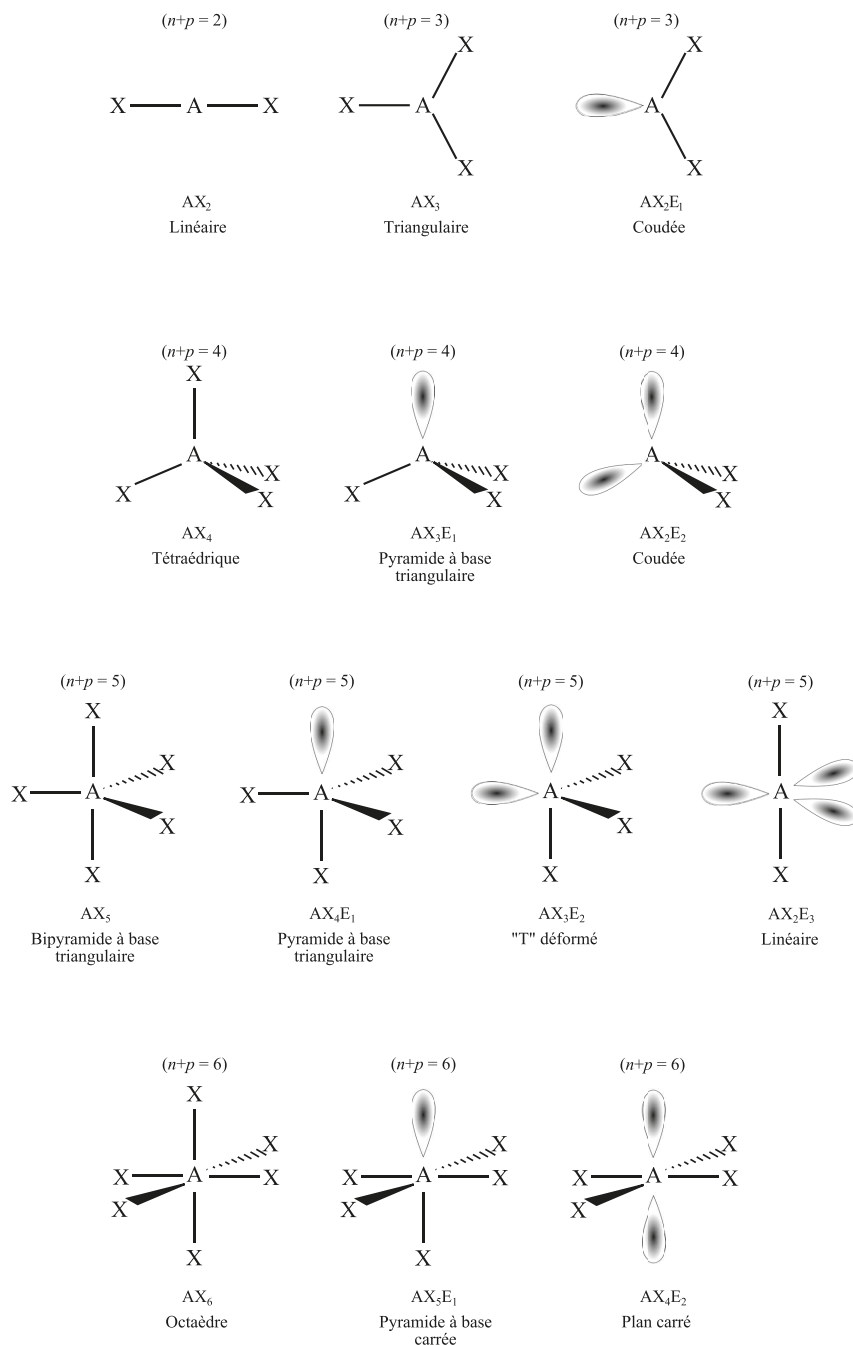
La **géométrie d'un édifice polyatomique** se décrit localement, autour d'un atome particulier de l'édifice, traditionnellement noté **atome central A** (rien à voir avec le nombre de masse). Pour ce faire :

- On détermine le nombre  $n$  (rien à voir avec le nombre quantique principal) d'atomes auxquels est lié l'atome A (notés X). Dans ce décompte, la multiplicité des liaisons n'intervient pas ; autrement dit, une liaison, même double ou triple, ne multiplie pas le nombre effectif d'atomes voisins.

- On détermine le nombre  $p$  (rien à voir avec les orbitales  $p$ ) de doublets non liants (souvent notés E) entourant l'atome A.
- On se reporte au tableau ci-après, que vous devez connaître, pour déterminer la géométrie de la molécule autour de l'atome A.

**Remarque :** afin de retenir plus facilement le tableau des géométries, on peut remarquer que les édifices de même valeur de  $n+p$  ont une géométrie dérivant directement de celle de  $AX_{n+p}$ .

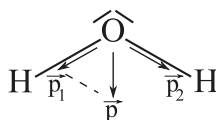
**Tableau des géométries**



## — À quelle(s) condition(s) une molécule est-elle polaire ?

Une molécule est dite **polaire** si elle possède un moment dipolaire électrique global (égal à la somme des moments dipolaires électriques de ses différentes liaisons) significatif. Pour déterminer la polarité d'une molécule, il est donc nécessaire de repérer d'éventuelles liaisons polarisées, mais également d'étudier sa géométrie à l'aide de la méthode VSEPR, afin de voir si les moments dipolaires électriques de ces liaisons ne se compensent pas entre eux.

Un exemple classique de molécule polaire est la molécule d'eau. En reprenant son schéma de Lewis et sa géométrie coudée, on peut représenter le moment dipolaire total de la molécule et vérifier qu'il n'est pas nul, ce qui confère à l'eau son caractère polaire.



## — À quelles échelles et par quelles interactions la cohésion de la matière est-elle assurée ?

Nous savons que les nucléons sont liés au sein du noyau par l'interaction nucléaire forte et que les électrons sont liés aux noyaux par l'interaction électromagnétique, assurant ainsi les cohésions de la matière respectivement aux échelles nucléaire et atomique.

Les molécules, elles, ne sont jamais qu'un genre de gros atome dont le noyau serait délocalisé sur plusieurs centres. La cohésion de l'ensemble est assurée par des liaisons de covalence, qui peuvent être vues comme résultant d'un équilibre entre la répulsion qu'exercent les noyaux des atomes les uns sur les autres, et l'attraction qu'ils exercent chacun sur les électrons de la molécule. Ainsi la cohésion des atomes au sein des molécules est-elle également d'origine électromagnétique.

Reste à expliquer la cohésion des corps à l'échelle macroscopique. En effet, si les solides ioniques et les métaux peuvent être vus comme de gigantesques molécules (aux liaisons très polarisées dans le cas des solides ioniques), l'existence des solides moléculaires, et plus encore des liquides, est moins évidente. En effet les molécules de ces corps restent solidaires les unes des autres alors qu'elles ne sont liées entre elles par aucune liaison de covalence, ni attraction coulombienne globale puisqu'elles sont électriquement neutres.

En réalité c'est à nouveau l'interaction coulombienne qui est à l'œuvre, mais dans une version plus subtile :

- si les molécules du corps considéré sont polaires, elles se comportent comme des **dipôles électriques permanents**, dotés d'un pôle positif et d'un pôle négatif. Le pôle positif d'une molécule pourra ainsi attirer le pôle négatif d'une autre (et réciproquement), entraînant donc dans une population de molécules un ensemble d'interactions attractives qui vont assurer une certaine cohésion à l'ensemble ;
- même si une molécule est apolaire, cette apolarité résulte d'un moment dipolaire électrique **moyen** nul. Cependant les électrons étant en mouvement autour des noyaux, une molécule possède toujours un moment dipolaire électrique **instantané** non nul. Et deux molécules, même apolaires, vont ainsi s'attirer en vertu d'une interaction analogue à la précédente, quoique significativement moins forte du fait du manque de constance des moments dipolaires électriques sur lesquels elle repose ;

- dans le cas où l'on met en contact des molécules polaires avec des molécules apolaires, le champ électrique généré par les premières va accentuer la polarisation des secondes, créant ainsi des moments dipolaires électriques **induits**.

Ces trois interactions, respectivement dites de Keesom, London et Debye, sont rassemblées sous l'appellation d'**interactions de Van der Waals**. Fondamentalement elles sont de même nature que les liaisons hydrogène (ou *liaisons H*) dont vous avez déjà entendu parler. Elles sont cependant moins fortes, ces dernières reposant sur des liaisons extrêmement polarisées (O – H, N – H, etc.). Cette différence trouve son illustration notamment dans les températures de changement d'état très élevées des corps purs au sein desquels elles s'exercent, en tête desquels l'eau.

On peut retenir, à l'échelle globale, les ordres de grandeur suivants (ici donnés pour plus de commodité en kilojoules d'énergie par mole de liaison, sachant que  $1 \text{ kJ.mol}^{-1}$  équivaut à environ 0,01 eV par molécule) :

| Cohésion             | Du noyau        | De l'atome        | De la molécule | Intermoléculaire |               |
|----------------------|-----------------|-------------------|----------------|------------------|---------------|
| Liaison              | Nucléaire forte | Électromagnétique |                | Liaison H        | Van der Waals |
| $\text{kJ.mol}^{-1}$ | $10^8$          | $10^3$            |                | $10^1$           | $10^0$        |

Cette disparité explique la difficulté croissante à dissocier des corps à mesure que l'on pénètre plus profondément dans la matière :

- la séparation des molécules dans un corps dense pour en faire un gaz engage peu d'énergie (un peu de chaleur et/ou de la diffusion) ;
- une réaction altérant les nuages électroniques (réaction chimique, donc, avec dissociation de molécules et/ou ionisation d'atomes) engage généralement des énergies plus importantes (réactions exothermiques notamment) ;
- quant à celles altérant le noyau lui-même, elles définissent le domaine de la radioactivité, dont nous connaissons les quantités vertigineuses d'énergies qu'elles sont susceptibles de libérer, pour le meilleur comme pour le pire.

Il est intéressant de noter que les différents niveaux auxquels peuvent s'exercer les interactions de Van der Waals (plus fortes entre les molécules polaires qu'entre les molécules apolaires) sont à l'origine des différences de solubilité ou de miscibilité des corps purs entre eux. En effet, les molécules s'associeront toujours préférentiellement avec les partenaires qui les attirent le plus (logique matrimoniale de base). Donc :

- si l'on mélange deux corps purs dont les molécules présentent une grande différence de polarité, les molécules polaires resteront entre elles plutôt que de s'associer avec les molécules apolaires, et tendront à former une phase à part qu'il sera difficile, voire impossible de dissoudre dans l'autre ;
- idem si l'on introduit un corps pur apolaire dans un solvant polaire : les molécules de solvant se serreront les coudes et empêcheront le corps pur apolaire de se dissoudre ;
- si en revanche on mélange entre eux deux corps purs tous deux constitués de molécules polaires, les molécules de l'un pourront s'associer avec celles de l'autre, et ce d'autant plus facilement que leurs niveaux de polarité seront proches ;
- enfin, si l'on mélange deux corps purs tous deux constitués de molécules apolaires, la faiblesse des attractions que chaque molécule exercera envers chacun deux types de partenaires disponibles fait que les deux corps pourront se mélanger dans la plus parfaite indifférence.

## 1.4 Les molécules de la chimie organique

### — Pourquoi la chimie organique constitue-t-elle une branche à part de la chimie ?

Historiquement, la chimie organique s'intéressait aux molécules produites exclusivement par les êtres vivants. Cette définition est devenue obsolète lorsque l'on s'est rendu compte qu'il était possible de produire certaines de ces molécules en laboratoire (emblématique : l'urée, découverte en 1773 par Hilaire Rouelle, est finalement synthétisée en 1828 par Friedrich Wöhler).

La chimie organique se redéfinit alors comme une chimie des molécules carbonées. Typiquement, une molécule organique comporte deux parties :

- un **squelette carboné**, constitué essentiellement d'atomes de **carbone** liés entre eux par des liaisons simples, et dont les liaisons vacantes sont essentiellement partagées avec des atomes d'**hydrogène** ;
- un ou plusieurs **groupes fonctionnels** : des sites constitués d'**hétéroatomes** (c'est-à-dire qui ne sont ni du carbone, ni de l'hydrogène) ou d'atomes de carbone particulièrement réactifs (liaisons multiples notamment).

Cette structure présente plusieurs intérêts :

- en jouant astucieusement avec les sites réactifs, il est possible de modéliser ces molécules à volonté : **transformer certains groupes fonctionnels** en d'autres, voire **altérer le squelette carboné**, confectionnant ainsi des molécules sur mesure pour répondre à tel ou tel besoin ;
- si l'origine biologique évoquée plus haut n'est plus au centre de la définition de la chimie organique, la structure des molécules organiques les rend malgré tout particulièrement aptes à réagir avec le vivant, et offre donc de **multiples applications en biologie**, dans le domaine pharmacologique en particulier ;
- ceci sans compter les multiples applications hors du domaine du vivant, notamment en **pétrochimie** (polymères, carburants, etc.).

La complexité et la diversité des molécules étudiées en chimie organique (environ 10 millions de molécules connues à ce jour), la multiplicité des réactions auxquelles celles-ci peuvent donner lieu, ainsi que celle des synthèses que ces réactions permettent d'élaborer, les techniques nombreuses et spécifiques qu'elle met en jeu pour y parvenir, lui confèrent donc un statut particulier.

### — Comment caractériser le squelette carboné d'une molécule organique ?

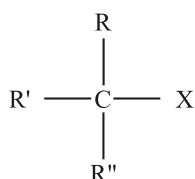
Les atomes de carbone à la base d'une molécule organique sont liés entre eux par des liaisons covalentes et forment une chaîne constituant le **squelette carboné** de la molécule. Ce squelette carboné peut prendre différentes formes ; il peut être :

- **linéaire**, c'est-à-dire constitué d'une succession d'atomes de carbone liés les uns à la suite des autres, sans qu'aucun n'en touche plus de 2 autres ;
- **ramifié**, c'est-à-dire constitué d'une chaîne principale (elle-même constituée de l'enchaînement le plus long d'atomes de carbone) et de chaînes secondaires liées à la principale ;
- **cyclique**, c'est-à-dire tel que la chaîne carbonée principale se referme sur elle-même ;
- **saturé**, c'est-à-dire sans cycle ni liaison ;
- **insaturé**, c'est-à-dire comportant au moins un cycle et/ou une liaison multiple.

## Quels sont les principaux groupes fonctionnels ?

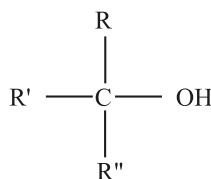
Les hétéroatomes présents dans les molécules organiques vont s'agencer de manière précise pour former des groupements appelés **groupes fonctionnels** et conférant à la molécule qui les possède une réactivité particulière. Il faut être capable de reconnaître dans une molécule les groupes fonctionnels suivants (sauf indication contraire, R, R' et R'' désignent des chaînes hydrogéné-carbonées quelconques) :

Halogénure d'alkyle

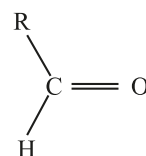


avec X = Cl, Br ou I

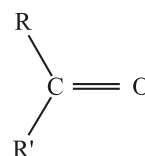
Alcool



Aldéhyde

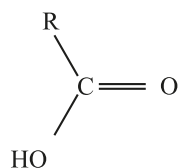


Cétone

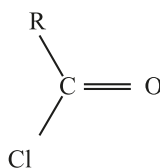


R et R' ≠ H

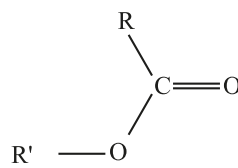
Acide carboxylique



Chlorure d'acyle

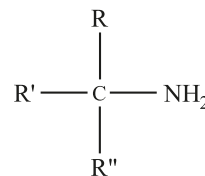


Ester



R' ≠ H

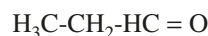
Amine



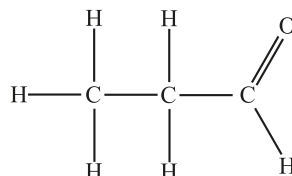
## Comment peut-on représenter une molécule organique ?

Plusieurs **formules**, de la plus concise à la plus détaillée, permettent de représenter une molécule organique :

- La **formule brute** indique simplement la constitution de la molécule en nombre d'atomes de chaque élément. Par exemple, on pourra considérer la molécule de propanal de formule brute  $\text{C}_3\text{H}_6\text{O}$ . Notons qu'à une même formule brute peuvent souvent être associées plusieurs molécules de structures distinctes appelées **isomères de constitution**.
- La **formule semi-développée** explicite les liaisons carbone-carbone et carbone-hétéroatome, mais sous-entend les liaisons carbone-hydrogène. Elle permet une bonne visualisation du squelette carboné de la molécule. La molécule de propanal sera ainsi représentée:

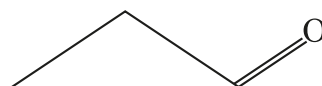


- La **formule développée** explicite en plus les liaisons carbone-hydrogène. Lourde à manier, elle ne sera plus guère utilisée en CPGE. Le propanal serait alors représenté :



**Remarque :** la géométrie est souvent à peu près respectée dans cette représentation ; par contre il ne faut pas confondre formule semi-développée plane et schéma de Lewis de la molécule. À ce titre, il n'y a pas à faire figurer de doublets non liants sur une formule développée plane...

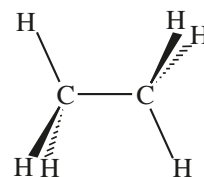
- La **formule topologique** est une représentation épurée qui ne fait plus apparaître que les hétéroatomes, ainsi que les atomes d'hydrogène directement liés à ceux-ci. La molécule est dessinée comme un polygone aux sommets duquel se trouvent les atomes de carbone. Les atomes d'hydrogène liés à ceux-ci ne sont pas représentés, non plus que leurs liaisons. On peut ainsi savoir combien d'atomes d'hydrogène sont liés à un atome de carbone donné, comme la différence entre le nombre total de doublets liants engagés par un atome de carbone (4) et le nombre de doublets représentés au départ de cet atome. Par exemple dans la représentation ci-dessous, un seul doublet est représenté au départ du carbone de gauche : il est donc lié à 3 hydrogène. 2 doublets sont représentés au départ du carbone médian : celui-ci est donc lié à 2 hydrogène. Enfin celui de droite, au départ duquel on trouve 3 doublets liants, est donc lié à un unique atome H. Cette représentation, très légère, est particulièrement adaptée aux molécules organiques (souvent copieuses). C'est essentiellement celle-ci que vous utiliserez en CPGE.



## — Comment représente-t-on de façon plane des structures spatiales ?

On utilise pour cela la **représentation de Cram**, qui repose sur les conventions de représentation suivantes :

- une liaison en **trait plein** se situe dans le plan de la figure ;
- une liaison en forme de **triangle allongé plein** représente une liaison orientée depuis le plan de représentation (situé à la pointe du triangle) vers l'observateur (situé à la base) ;
- une liaison en forme de **triangle allongé hachuré** représente une liaison orientée depuis l'observateur (situé à la pointe du triangle) vers le plan de représentation (situé à la base).



On obtient ainsi une vision spatiale de la molécule, comme nous pouvons le voir sur la molécule d'éthane représentée ci-contre :

## — Qu'appelle-t-on conformations d'une molécule ?

Les liaisons simples carbone-carbone permettent une libre rotation de la molécule autour des axes C-C correspondants. Une même molécule peut donc adopter différentes structures spatiales grâce à ce degré de liberté.

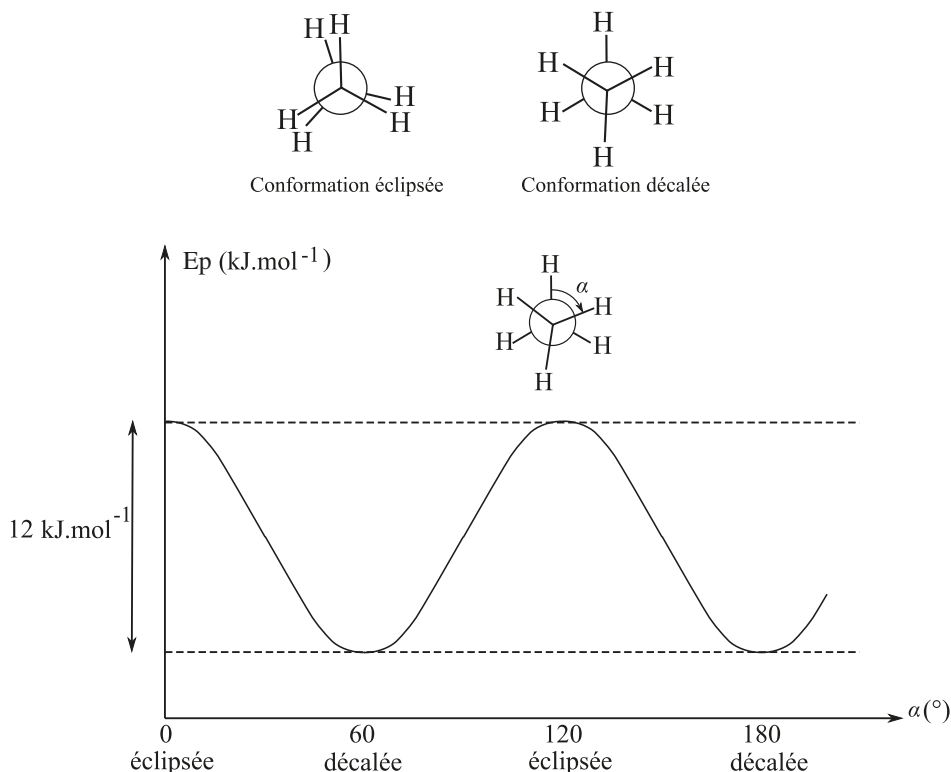
On appelle **conformations** d'une même molécule des structures qui ne diffèrent que par rotation autour de liaisons simples.

**Remarque :** il apparaît de manière évidente que plus la molécule possède de liaisons simples, plus elle possède de conformations différentes. Pour fixer les idées, nous envisagerons ici uniquement les rotations autour d'une unique liaison simple carbone-carbone.

Pour illustrer la notion de conformation, considérons le cas de la molécule d'éthane. Deux conformations remarquables apparaissent : la **conformation décalée** et la **conformation éclipsée**. Ces deux conformations ne présentent pas la même stabilité. En conformation éclipsée les atomes d'hydrogène sont plus proches les uns des autres qu'en conformation

décalée. Il en résulte une certaine déstabilisation du fait de la **gêne stérique** (c'est-à-dire liée à l'encombrement).

Comme un système est d'autant plus stable que son énergie est basse, on peut alors représenter sur un diagramme énergétique les différentes conformations de cette molécule, en prenant comme paramètre pertinent un angle de torsion de la molécule :



On appelle **conformère** une conformation particulière d'une molécule. Dans le cas de l'éthane, on constate que passer du conformère décalé au conformère éclipsé réclame  $12\text{ kJ}$  par mole de molécules affectées.

## — Comment nomme-t-on un hydrocarbure acyclique ?

La **nomenclature** de la chimie organique se construit en décrivant, à l'aide de règles très précises, la constitution exacte de la molécule. Nous allons nous limiter ici à rappeler les règles pour nommer deux types d'hydrocarbures acycliques : les **alcane**s et les **alcène**s.

Les **alcane**s désignent les hydrocarbures **saturés acycliques**. On peut aisément montrer qu'ils ont pour formule brute générale  $\text{C}_n\text{H}_{2n+2}$  où  $n$  est un entier naturel non nul. Les quatre premiers alcane linéaires de la famille portent des noms consacrés par l'usage :

- l'alcane linéaire de formule brute  $\text{CH}_4$  s'appelle le **méthane** ;
- l'alcane linéaire de formule brute  $\text{C}_2\text{H}_6$  s'appelle l'**éthane** ;
- l'alcane linéaire de formule brute  $\text{C}_3\text{H}_8$  s'appelle le **propane** ;
- l'alcane linéaire de formule brute  $\text{C}_4\text{H}_{10}$  s'appelle le **butane**.

**Remarque** : les **substituants alkyles** correspondants (de formule générique  $-\text{C}_n\text{H}_{2n+1}$ ), obtenus en ôtant un atome d'hydrogène, porteront respectivement comme noms : méthyle, éthyle, propyle et butyle.

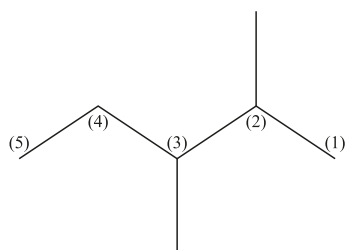
Pour déterminer le nom d'un alcane :

- On commence par identifier la chaîne carbonée la plus longue présente dans la molécule ; c'est elle qui détermine la base du nom (pour une chaîne en  $C_5$ , le nom dérivera ainsi du pentane).

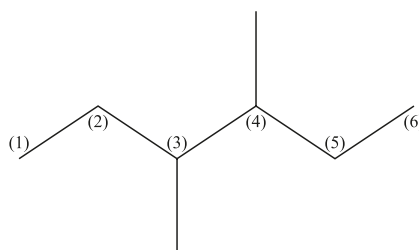
**Remarque :** s'il y a plusieurs chaînes candidates, on donne la priorité à celle portant le plus grand nombre de substituants.

- On numérote la chaîne de telle sorte que la somme des indices attribués aux substituants soit la plus faible possible.
- On place les noms des substituants alkyles en préfixe, dans l'ordre alphabétique et précédés de leurs indices de position devant le nom de base de l'alcane. En cas de répétition d'un même substituant, on utilise les préfixes « di », ou « tri » (qui n'interviennent alors pas dans le classement par ordre alphabétique).

Illustrons ceci par quelques exemples :



2,3-diméthylpentane



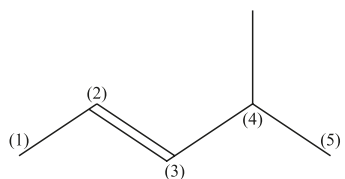
3,4-diméthylhexane

Les **alcènes** sont des hydrocarbures acycliques non saturés comportant une double liaison dans leur chaîne carbonée principale (notons que, pour établir le nom de la molécule, le fait de contenir le groupement alcène prime sur le fait que cette chaîne soit la plus longue ; double liaison dans la chaîne principale constitue donc, dans notre approche encore très simple de la nomenclature, une forme de pléonasme). Leur formule générale s'écrit  $C_nH_{2n}$ . Leur nom dérive de celui de l'alcane de même squelette carboné en y remplaçant la terminaison -ane par une terminaison -ène. La position de la double liaison dans la chaîne carbonée est indiquée par un indice avant la terminaison -ène, indice indiquant le numéro du premier carbone rencontré qui est engagé dans la double liaison. S'il existe une isomérisie *Z/E* (cf. les questions d'isométrie, plus loin), elle est précisée à la fin du nom.

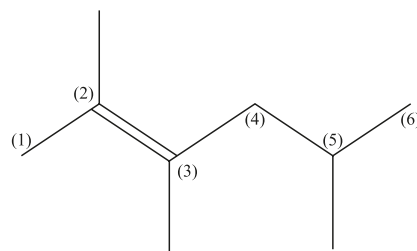


Le sens de la numérotation doit donner à la double liaison le plus petit indice possible. Ce critère a priorité sur la minimisation de la somme des indices des substituants alkyles.

Illustrons ceci par quelques exemples :



4-méthylpent-2-ène (E)



2,3,5-triméthylhex-2-ène

**Remarque :** l'alcène  $\text{H}_2\text{C}=\text{CH}_2$  porte comme nom d'usage l'**éthylène** en lieu et place de « éthène ».

**Remarque :** il existe d'autres hydrocarbures acycliques non saturés tels que les alcynes, les diènes...

## — Qu'est-ce qu'un alcool ?

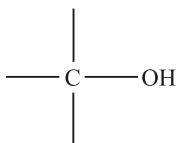
Un **alcool** est une espèce organique, dont la molécule possède le groupe **hydroxyle** -OH sur un atome de carbone tétraédrique. Cet atome de carbone, porteur du groupe hydroxyle, est appelé **carbone fonctionnel**.

On peut formellement considérer qu'un alcool dérive d'un alcane dont l'un des atomes d'hydrogène a été remplacé par le groupe hydroxyle. De ce fait, la formule brute générale d'un alcool saturé s'écrit  $\text{C}_n\text{H}_{2n+1}\text{OH}$ , soit  $\text{C}_n\text{H}_{2n+2}\text{O}$ .

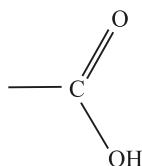


Le groupe -OH apparaît également dans d'autres groupements fonctionnels tels que celui des acides carboxyliques ou des énols. On reconnaît un alcool au fait que le carbone fonctionnel ne porte que des liaisons de covalence simples.

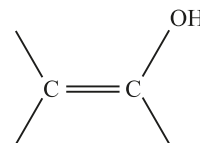
Alcool



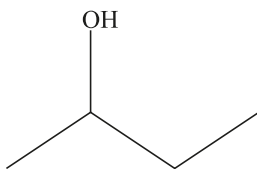
Acide carboxylique



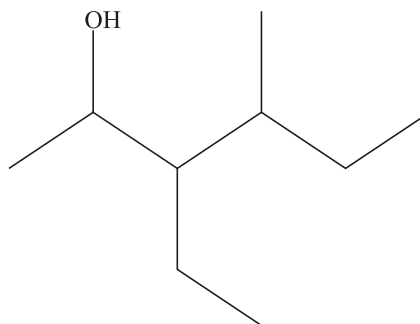
Enol



Enfin, en termes de nomenclature, celle des alcools se déduit de celle des alcanes de même chaîne carbonée en remplaçant le « e » final par la terminaison -ol, précédée de son indice de position dans la chaîne carbonée la plus longue la contenant. Là encore, la minimisation de l'indice de position du groupe hydroxyle l'emporte sur la minimisation des indices (ou de la somme des indices) des éventuels substituants alkyles. Les exemples suivants illustrent ce propos :



Butan-2-ol

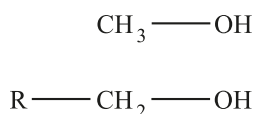


3,4 éthylméthylheptan-2-ol

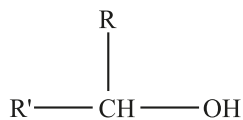
## — Qu'est-ce que la classe d'un alcool ?

La **classe d'un alcool** est déterminée par le nombre d'atomes de carbone auxquels est lié le carbone fonctionnel. Ce nombre peut varier de zéro (pour le méthanol  $\text{CH}_3\text{OH}$ ) à trois.

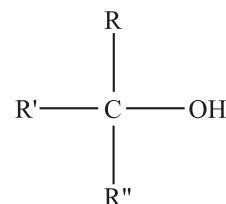
Alcool primaire



Alcool secondaire



Alcool tertiaire



R, R' et R'' ≠ H

On appelle ainsi :

- **alcool primaire** un alcool dont le carbone fonctionnel est lié au plus à un autre atome de carbone ;
- **alcool secondaire** un alcool dont le carbone fonctionnel est lié à deux autres atomes de carbone exactement ;
- **alcool tertiaire** un alcool dont le carbone fonctionnel est lié à trois autres atomes de carbone.

On peut schématiser les différentes classes d'alcool comme présentées ci-dessus.

**Remarque :** cette notion de classe est importante car elle influe sur la réactivité de l'alcool considéré.

## — Qu'est-ce qu'un composé carbonylé ?

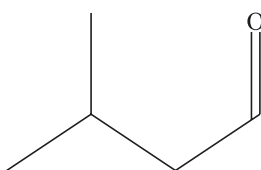
Un **composé carbonylé** est une molécule organique contenant le groupe caractéristique **carbonyle** C=O. L'atome de carbone de la chaîne carbonée portant le groupe carbonyle est appelé carbone fonctionnel ; on peut déduire de la théorie VSEPR qu'autour de ce carbone fonctionnel, la molécule présente une géométrie triangulaire plane.

Parmi les composés carbonylés, on distingue :

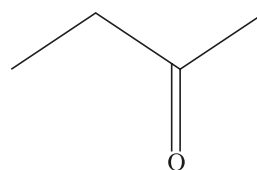
- les **aldéhydes**, pour lesquels le carbone fonctionnel est lié à **au moins un atome d'hydrogène** ;
- les **cétones**, pour lesquelles l'atome de carbone fonctionnel est lié uniquement à des atomes de carbone et à l'atome d'oxygène du groupe carbonyle.

La nomenclature de ce type de composé se construit de la même manière que celle des alcanes, alcènes et alcools. La terminaison dédiée aux aldéhydes est -al, celle dédiée aux cétones est -one. Notez au passage que le carbone fonctionnel d'un aldéhyde est forcément terminal (le groupe fonctionnel est dit *trivalent* : il mobilise une liaison pour l'hydrogène, 2 pour l'oxygène, n'en laissant qu'une seule vacante pour raccrocher le carbone fonctionnel au reste du squelette) ; il n'y a alors aucun indice de position à noter pour la fonction aldéhyde dont le carbone porteur est d'emblée le numéro 1 de la chaîne carbonée principale...

Illustrons ceci par des exemples :



3-méthylbutanal



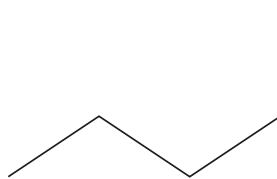
Butan-2-one

## — Qu'est-ce que l'isomérisie de constitution ?

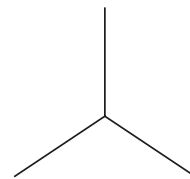
L'imprécision liée à la donnée de la formule brute d'une molécule organique mène à définir la notion d'**isomérisie de constitution**. On appelle **isomères de constitution** des molécules de formules brutes identiques mais différant par leurs formules semi-développées (ainsi que développées et topologiques). La différence pouvant se faire à plusieurs niveaux, on distingue généralement trois grands types d'isomérisie de constitution :

- L'**isomérisie de chaîne** lorsque la différence entre les deux molécules se situe au niveau de la chaîne carbonée.

Exemple :  $C_4H_{10}$



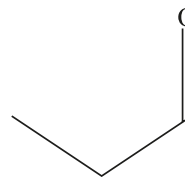
Butane



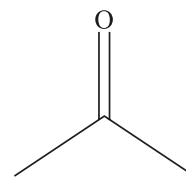
2-méthylpropane

- L'**isomérisie de fonction** lorsque la différence entre les deux molécules se situe au niveau des groupes fonctionnels présents.

Exemple :  $C_3H_6O$



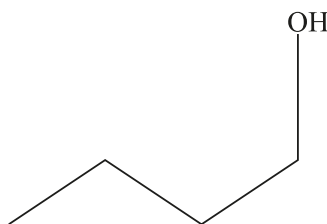
Propanal



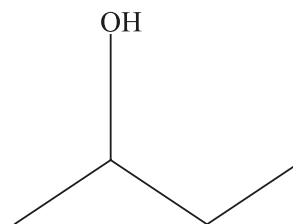
Propanone

- l'**isomérisie de position** lorsque les chaînes carbonées et les groupes fonctionnels sont identiques, mais la position de ces derniers différente.

Exemple :  $C_4H_{10}O$



Butan-1-ol



Butan-2-ol

## — Qu'est-ce que la stéréoisomérisie de configuration ?

La **stéréoisomérisie de configuration** est une isomérisie qui résulte uniquement de la position relative des atomes d'une molécule. Deux stéréoisomères de configuration ont donc en commun, outre la formule brute, leur formule semi-développée.

Contrairement aux conformations d'une molécule, dans le cas de stéréoisomères de configuration il est nécessaire de rompre une liaison pour passer d'un stéréoisomère à un autre.

On va distinguer ici deux types de stéréoisomérisie de configuration :

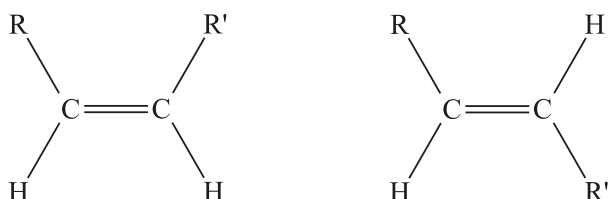
- la stéréoisomérisie de configuration Z/E ;
- la stéréoisomérisie « optique ».

## — Qu'est-ce que l'isomérisie de configuration Z/E ?

Un type particulier d'isomérisie de configuration est l'**isomérisie Z/E** qui se rencontre pour les **composés éthyléniques**, c'est-à-dire les composés comportant une double liaison carbone-carbone. Cette isomérisie concerne plus exactement les molécules du type  $R-CH=CH-R'$  où R et R' sont des substituants alkyles. Ces molécules sont planes ; on peut alors envisager de placer R et R' de deux manières différentes dans le plan de molécule :

- si R et R' sont situés du même côté de l'axe de la double liaison, on sera en présence de l'isomère Z (de l'allemand *zusammen*, « ensemble ») ;
- si R et R' sont situés de part et d'autre de l'axe de la double liaison, on sera en présence de l'isomère E (de l'allemand *entgegen*, « de part et d'autre »).

On peut alors schématiser ces deux stéréoisomères de configuration (qui constituent des **diastéréoisomères**, cf. plus loin) comme suit :



Isomère Z

Isomère E

Une autre manière de définir la notion d'isomérisie de configuration est de vérifier qu'il n'est pas possible de passer d'un stéréoisomère à l'autre sans casser de liaison. C'est bien le cas ici ; en effet, contrairement au cas d'une liaison simple autour de laquelle existe une possibilité de libre rotation, une double (ou triple) liaison est rigide.

**Remarque :** la différence entre deux stéréoisomères Z/E n'est pas simplement formelle ; elle se manifeste sur le plan expérimental, par des propriétés physiques et chimiques souvent dissemblables.

## — Qu'est-ce que la chiralité ?

Une molécule est dite **chirale** si elle n'est pas superposable à son image dans un miroir plan.

**Remarque :** un exemple typique d'objet chiral est la main (qui d'ailleurs est à l'origine du mot *chiral*). En effet, une main et son image par un miroir ne sont pas superposables puisque l'image d'une main droite par un miroir est une main gauche.

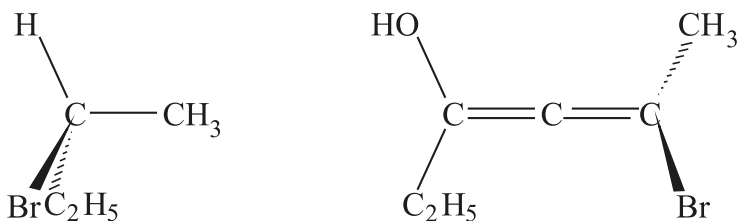
Un carbone tétraédrique est dit **asymétrique** lorsqu'il est lié à quatre atomes ou groupes d'atomes tous différents les uns des autres.

Il existe alors plusieurs façon d'obtenir une molécule chirale. De manière formelle, on peut affirmer (condition nécessaire et suffisante) qu'une molécule est chirale s'il n'existe pas d'entier  $n$  permettant d'obtenir une image par rotation d'angle  $\frac{2\pi}{n}$  autour d'un axe, suivie d'une

symétrie plane par rapport à un plan orthogonal à cet axe, qui soit identique à la molécule d'origine. Cette définition très théorique n'est pas pratique pour l'usage courant ! On retiendra alors plus simplement :

- Une condition nécessaire mais non suffisante de chiralité d'une molécule est de ne posséder ni plan, ni centre de symétrie. Si la molécule en possède un, on peut affirmer qu'elle n'est pas chirale. Cependant une molécule peut être chirale sans forcément posséder de carbone asymétrique.
- Une condition suffisante mais non nécessaire de chiralité est de posséder un seul carbone asymétrique. Notons aussi qu'une molécule avec plusieurs carbones asymétriques peut ne pas être chirale.

Voici deux exemples de molécules chirales (la première avec un carbone asymétrique, la seconde sans) :



## — Qu'est-ce que l'énantiométrie ?

Des **énantiomères** sont des molécules stéréoisomères, images l'une de l'autre par un miroir plan et non superposables.

**Remarque :** deux molécules énantiomères sont donc forcément chirales.

On appelle **mélange racémique** un mélange équimolaire de deux espèces énantiomères l'une de l'autre.

**Remarque :** on rencontre en biologie de nombreuses espèces chirales dont seul l'un des énantiomères présente un intérêt (l'autre étant sans action ou, pire, nocif).

Deux molécules énantiomères ont des propriétés physiques et chimiques identiques (sauf vis-à-vis de la lumière polarisée, comme vous le reverrez en CPGE), mais le plus souvent des propriétés biologiques différentes.

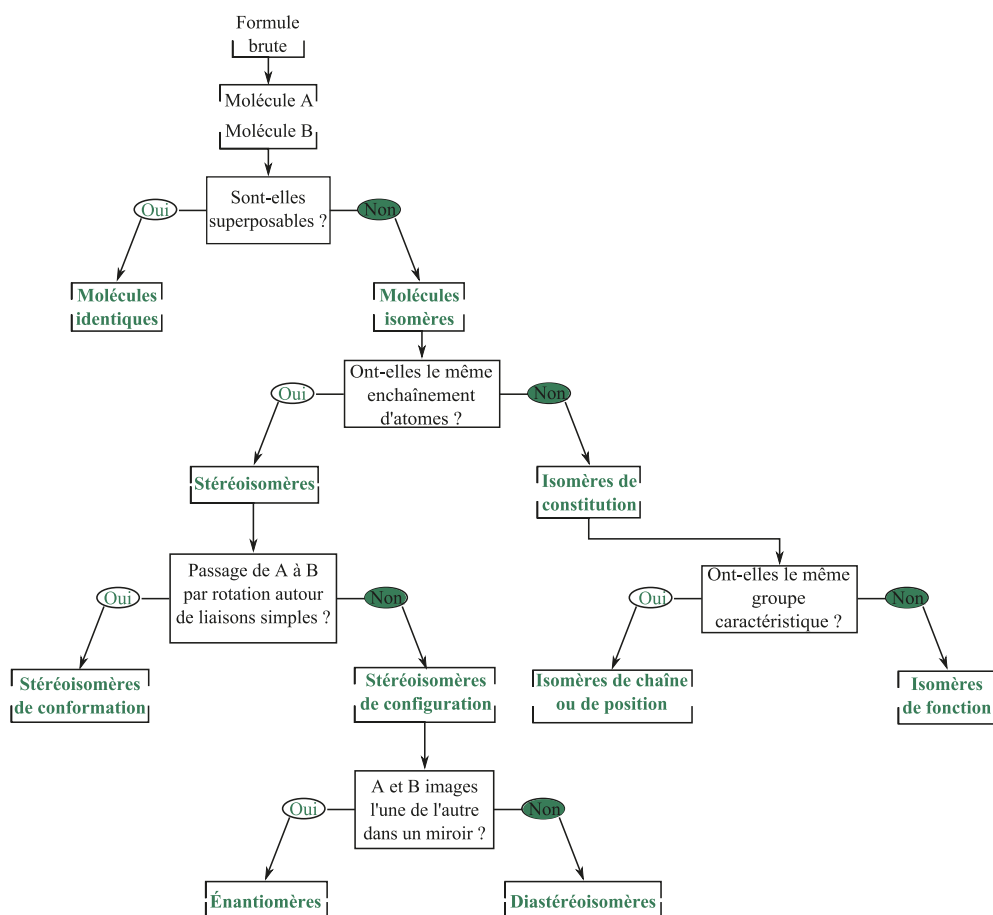
## — Qu'est-ce que la diastéréoisométrie ?

Des **diastéréoisomères** sont des molécules stéréoisomères qui ne sont pas des énantiomères. On peut donc rencontrer comme diastéréoisomères :

- les stéréoisomères de configuration Z et E d'un composé comportant une double liaison carbone-carbone ;
- certains composés comportant plus d'un carbone asymétrique.

Des diastéréoisomères ont des propriétés physiques et chimiques différentes. En particulier leurs températures d'ébullition et de fusion sont différentes (ce qui peut constituer un moyen d'identification).

Pour conclure, on peut résumer l'étude de la relation d'isomérisie par le tableau suivant :



## 1.5 Spectres d'émission et d'absorption des atomes

### Qu'est-ce que le phénomène de résonance ?

On appelle **oscillateur** tout système possédant au moins deux états physiques, et passant périodiquement de l'un à l'autre. On appelle par ailleurs **état d'équilibre** d'un oscillateur, tout état de cet oscillateur dans lequel il persévère, lorsqu'il s'y trouve abandonné à lui-même. Un exemple classique est le pendule simple, qui passe d'une position extrême à une autre avec un mouvement de balancier, et possède une position d'équilibre à la verticale de son point de fixation.

Lorsqu'un oscillateur simple (c'est-à-dire ne possédant qu'une seule articulation libre de bouger ou encore ne possédant qu'un seul degré de liberté) est abandonné à lui-même hors situation d'équilibre, on dit qu'il oscille **en régime libre**. On appelle alors **fréquence propre** de cet oscillateur la fréquence à laquelle il oscille spontanément en régime libre.

On peut également coupler un oscillateur à un excitateur (autre oscillateur, qui impose au premier sa propre fréquence d'oscillation). On dit alors de cet oscillateur qu'il oscille **en régime forcé** : après un petit temps d'adaptation (phase dite *transitoire*, au cours de laquelle le régime forcé se met en place tandis que le régime libre s'éteint sous l'effet des forces dissipatives), l'oscillateur finit par osciller à la fréquence de cet excitateur.

Cependant, les autres caractéristiques des oscillations obtenues (amplitude notamment) dépendent de la fréquence à laquelle l'oscillateur est forcé. Et l'on constate que si la plupart des fréquences n'aboutissent qu'à des réponses modérées de l'oscillateur (amplitude faible, retard de l'oscillateur par rapport à l'excitateur), son comportement devient en revanche spectaculaire lorsque le régime forcé lui impose une fréquence justement égale à sa fréquence propre. Le couplage est cette fois optimal (la structure de l'oscillateur le faisant spontanément osciller à cette fréquence, il reçoit chaque nouvelle séquence d'excitation au moment même où il se trouve en phase avec l'excitateur), et l'oscillateur accumule rapidement des quantités d'énergie qui vont aboutir à une croissance très rapide de l'amplitude de ses oscillations.

Ce phénomène d'amplification des oscillations d'un oscillateur lorsqu'il se trouve excité à sa fréquence propre est appelé **résonance** de l'oscillateur avec l'excitateur. Dans ces circonstances, l'oscillateur forcé est alors qualifié de **résonateur**.

Lorsqu'un oscillateur est plus complexe qu'un pendule simple (en particulier lorsque sa structure offre plus d'un unique degré de liberté : articulations, axes de rotation, liaisons élastiques, etc.), il peut posséder plusieurs fréquences propres (il en possède en fait autant que de degrés de liberté). Son fonctionnement en régime libre est alors une superposition de réponses sinusoïdales oscillant chacune à l'une de ces fréquences propres, et il est possible d'en privilégier une sur les autres, en choisissant les bonnes conditions initiales de fonctionnement.

Mais plus que le régime libre, c'est dans le cas présent le régime forcé qui va nous intéresser, lorsque le phénomène de résonance sera cette fois observé à chaque fois que la fréquence imposée par l'excitateur correspondra à l'une des fréquences propres de cet oscillateur complexe.

## — Comment interpréter le spectre d'émission d'un gaz sous basse pression ?

Lorsque l'on enferme un gaz dans une ampoule, dans des conditions telles que la pression y soit très faible, on constate que l'envoi, au moyen de deux électrodes, de décharges électriques à travers le gaz contenu dans cette ampoule, provoque l'émission d'une lumière particulière. En effet, si on la disperse (au moyen d'un prisme par exemple), on observe que cette lumière comporte quelques raies lumineuses seulement, dont les longueurs d'onde, et donc les fréquences, ont des valeurs bien déterminées.

Or nous savons que la fréquence  $\nu$  d'une onde électromagnétique est liée à l'énergie  $E$  qu'elle véhicule, par la relation de Planck :  $E = h\nu$ . Nous savons également que cette énergie n'est autre que l'écart énergétique entre deux des niveaux quantiques de l'atome concerné. La conclusion de tout ceci est que l'atome se comporte en fait comme un oscillateur, possédant plusieurs fréquences propres. Lorsqu'il est excité (par une décharge de courant électrique, par exemple), l'atome passe dans l'un des états d'énergie supérieure que lui autorise la mécanique quantique : certains de ses électrons changent d'orbitale (voire de niveau), puis retombent dans des états d'énergie inférieure en oscillant à une fréquence  $\nu$ , fréquence que l'on retrouve dans le rayonnement électromagnétique émis à cette occasion. Les fréquences des raies lumineuses observées ne sont alors ni plus, ni moins, que les fréquences propres d'oscillation de l'atome. Dans le gaz sous basse pression, les atomes sont éloignés les uns des autres et interagissent très peu. Ainsi, lorsqu'ils

sont excités, chacun émet tout ou partie de son spectre propre, et la superposition de ces spectres, tous identiques entre eux, donne un spectre global d'intensité supérieure, mais fidèle à l'identité de chacun des atomes qui a concouru à son élaboration. Ce spectre est appelé **spectre de raies d'émission**, est caractéristique des atomes (ou des molécules) constituant le gaz considéré.

## — Comment interpréter le spectre d'émission d'un corps dense chauffé ?

Lorsque l'on augmente la pression d'un gaz initialement sous basse pression, les atomes interagissent de plus en plus entre eux du fait de l'augmentation de la fréquence de leurs collisions. Les raies du spectre précédent s'élargissent progressivement, puis se recouvrent et donnent finalement un spectre continu. Toutes les fréquences y sont présentes, mais certaines contribuent plus que d'autres. La courbe globale présente cependant une certaine continuité et la propriété très particulière de ne plus dépendre que de la température du corps qui l'émet, à l'exclusion notamment de la nature de ce corps (tant que celui-ci est un corps dense, s'entend).

Ce spectre est appelé **spectre continu d'origine thermique**. La recherche de sa modélisation à partir d'oscillateurs a amené Max Planck à devoir formuler l'hypothèse que l'énergie des oscillateurs constitutifs de la matière (autrement dit les atomes ou les molécules) était quantifiée, ouvrant la voie à la mécanique quantique. Bien que ce spectre ne nous soit pas vraiment utile dans le cadre de l'analyse spectrale des molécules (puisque indépendant de leur nature), il eût été dommage de traiter le sujet sans lui rendre ce modeste hommage.

## — Comment interpréter un spectre de raies d'absorption ?

Après avoir décrit les atomes comme des oscillateurs, et leur comportement en régime libre, on peut se demander comment utiliser ces oscillateurs en régime forcé. Travailler fréquence par fréquence serait relativement difficile, puisque nous aurions alors besoin d'une lumière monochromatique de fréquence réglable à volonté pour observer la réponse d'une population d'atomes sous basse pression, à telle ou telle fréquence excitatrice. Le problème étant que si nous possédons des sources monochromatiques (les lasers notamment), nous n'avons pas le moyen de choisir n'importe quelle fréquence au gré de nos besoins.

On opte alors pour une solution un peu plus brutale : envoyer une lumière contenant toutes les fréquences (par exemple un spectre continu d'origine thermique) sur l'ampoule contenant le gaz sous basse pression, et déterminer les fréquences des radiations pour lesquelles ce gaz semble avoir de l'appétit. Concrètement, on observe le spectre de la lumière issue d'une lampe à filament (émettant un spectre continu d'origine thermique) après sa traversée du gaz sous basse pression.

On observe alors un spectre qui ressemble fortement au spectre continu d'origine thermique, à ceci près qu'il est parsemé de raies sombres : certaines radiations sont particulièrement absorbées lors de leur passage dans l'ampoule. Si l'on analyse les fréquences de ces radiations, on constate qu'il s'agit précisément des fréquences que le gaz sous basse pression émettait lorsqu'il était parcouru de décharges électriques. On retrouve ainsi le fait que les oscillateurs que constituent les atomes absorbent bien les radiations correspondant à leurs fréquences propres, c'est-à-dire précisément celles qu'ils sont aussi capables d'émettre. Ce spectre est appelé **spectre de raies d'absorption**.

## 1.6 Analyse spectrale en chimie

### — Que se passe-t-il pour une association de plusieurs atomes ?

Une molécule est un ensemble de particules dont la plupart sont électriquement chargées, liées entre elles : protons et neutrons dans le noyau, noyau et électrons dans l'atome, doublets d'électrons dans les liaisons de covalence. Ces particules peuvent se trouver dans différents états, caractérisés notamment par l'énergie dont elles sont dotées. Elles peuvent passer d'un état à l'autre, et se comportent donc elles aussi comme de petits oscillateurs.

L'atome n'est donc qu'un exemple parmi d'autres des oscillateurs microscopiques que l'on peut trouver dans la matière. Par conséquent une molécule elle-même peut être décrite comme un assemblage plus ou moins gros, d'oscillateurs se situant à des échelles diverses.

Un modèle moléculaire fidèle ne représenterait pas les doublets d'électrons assurant les liaisons entre atomes par des tiges rigides, mais par de petits ressorts. Très détaillé, il comporterait également des ressorts de rigidité accrue pour figurer les liaisons nucléaires fortes assurant la cohésion des nucléons au sein des noyaux.

Une onde électromagnétique est la propagation d'une perturbation du champ électromagnétique, oscillant à une certaine fréquence. On sait qu'un tel champ est capable de mettre en mouvement les particules électriquement chargées. Ainsi, lorsqu'un rayonnement électromagnétique atteint une molécule, il agit comme un excitateur sur les liaisons de celles-ci, qui vont se mettre à vibrer à la fréquence qu'il leur impose. Mais cette excitation sera efficace, autrement dit les liaisons absorberont une part significative de l'énergie véhiculée par cette onde, uniquement si la fréquence de celle-ci coïncide avec l'une de leurs fréquences propres.

### — Dans quels domaines une molécule possède-t-elle des fréquences propres ?

Le spectre d'absorption d'une molécule permet ainsi de déterminer les fréquences propres associées aux diverses liaisons qui assurent sa cohésion. Cependant, toutes les liaisons n'ont pas les mêmes fréquences propres, et n'absorbent donc pas aux mêmes fréquences.

Pour déterminer la nature et la structure d'une molécule, la chimie recourt à 3 types de spectres. Deux d'entre eux (UV-visible et IR, vus au lycée) seront traités ici. Le troisième (la RMN), très performant mais requérant un formalisme très spécifique, sera cité pour mémoire, dans l'attente d'un développement en CPGE.

- Dans le domaine UV-visible : on excite l'ensemble de la molécule. L'absorption dans ces domaines nous permettra d'identifier la nature d'une molécule (ainsi que sa concentration en solution) par identification au spectre d'une espèce connue, mais pas sa structure.
- Dans le domaine IR : on excite les liaisons de covalence de la molécule. L'absorption dans ce domaine nous permettra d'identifier la présence de liaisons particulières au sein d'une molécule.
- Dans le domaine radio, mais avec des rayonnements de très grande amplitude : on excite les noyaux des atomes d'hydrogène. L'absorption dans ce domaine nous permettra donc d'identifier les populations d'atomes H équivalents ainsi que leurs environnements respectifs au sein de la molécule. Il est alors possible d'en déduire l'exacte répartition de ces atomes au sein de la molécule. Conjugée à la spectroscopie IR, cette technique appelée RMN (**R**ésonance **M**agnétique **N**ucléaire) permet d'identifier atome par atome la structure de la molécule considérée. Elle ne sera pas détaillée dans le cadre de cet ouvrage.

## — Quelles informations peut-on tirer de l'analyse UV-visible ?

La spectroscopie dans les domaines visible et ultraviolet est celle que vous avez menée au lycée. Elle n'apporte malheureusement aucune information sur la structure même de la molécule : le spectre est grossier et permet au mieux de la distinguer d'une autre molécule dont on connaîtrait également le spectre. Mais quant à savoir quels groupements fonctionnels elle comporte, l'éventuelle présence de carbone tétravalent... Elle ne donne rien.

Son intérêt se situe surtout dans les suivis cinétiques, lorsqu'une transformation lente engage une espèce colorée, dont la loi de Beer-Lambert nous permet de suivre la concentration. L'instrument dédié à la mesure de l'absorbance est appelé spectrophotomètre ; il sera détaillé au chapitre suivant.

## — Quelles informations peut-on tirer de l'analyse IR ?

Le spectre infrarouge est quant à lui beaucoup plus intéressant pour identifier le détail de la molécule : les oscillateurs qui vont y laisser leur marque sont certaines liaisons particulières. Or qui dit liaison particulière dit groupe fonctionnel. L'analyse de ce spectre permet donc de déterminer l'éventuelle présence de liaisons moins évidentes que C-H, par exemple.

Traditionnellement, les spectres IR présentent la transmittance de la solution (contrairement à l'UV-visible, où l'on utilise plutôt l'absorbance). La résonance (qui se manifeste par une montée en flèche de l'absorption) ne se traduit donc pas ici par un pic de la grandeur observée, mais au contraire par une chute.

Autre différence : l'abscisse n'est pas graduée en longueur d'onde  $\lambda$  ou en fréquence  $\nu$ , mais en une autre grandeur caractéristique du rayonnement considéré : le nombre d'onde  $\sigma = \frac{1}{\lambda}$  (attention à ne pas confondre avec la conductivité) usuellement exprimé en  $\text{cm}^{-1}$ .

On caractérise alors l'absorption due à une liaison par l'intervalle de valeurs du nombre d'onde où se produit cette absorption. Celui-ci est souvent large, l'environnement chimique des atomes engagés dans cette liaison étant susceptible de modifier plus ou moins la fréquence de résonance et donc le nombre d'onde associé. On lui adjoint en général quelques détails qualitatifs concernant la chute de transmittance : la profondeur et la largeur le plus souvent.

| Liaison                                       | $\sigma (\text{cm}^{-1})$ | Chute de transmittance |
|-----------------------------------------------|---------------------------|------------------------|
| O – H <sub>libre</sub>                        | 3 580-3 650               | Forte et fine          |
| O – H <sub>lié</sub>                          | 3 200-3 400               | Forte et large         |
| N – H                                         | 3 100-3 500               | Moyenne                |
| C <sub>trigonal</sub> – H                     | 3 000-3 100               | Moyenne                |
| C <sub>trigonal</sub> – H <sub>arom</sub>     | 3 030-3 080               | Moyenne                |
| C <sub>tétragonal</sub> – H                   | 2 800-3 000               | Forte                  |
| C <sub>trigonal</sub> – H <sub>aldéhyde</sub> | 2 750-2 900               | Moyenne                |
| O – H <sub>acide carboxylique</sub>           | 2 500-3 200               | Forte et large         |

| Liaison                                             | $\sigma(\text{cm}^{-1})$ | Chute de transmittance |
|-----------------------------------------------------|--------------------------|------------------------|
| $\text{C} = \text{O}_{\text{ester}}$                | 1 700-1 740              | Forte                  |
| $\text{C} = \text{O}_{\text{aldéhyde/cétone}}$      | 1 650-1 730              | Forte                  |
| $\text{C} = \text{O}_{\text{acide carboxylique}}$   | 1 680-1 710              | Forte                  |
| $\text{C} = \text{C}$                               | 1 625-1 685              | Moyenne                |
| $\text{C} = \text{C}_{\text{aromatique}}$           | 1 450-1 600              | Moyenne                |
| $\text{C}_{\text{tétra}} - \text{H}$                | 1 415-1 470              | Forte                  |
| $\text{C}_{\text{tétra}} - \text{O}$                | 1 050-1 450              | Forte                  |
| $\text{C}_{\text{tétra}} - \text{C}_{\text{tétra}}$ | 1 000-1 250              | Forte                  |

## — Quelles sont les autres sources d'information à notre disposition ?

Les spectres IR et RMN utilisés conjointement permettent de déterminer complètement la structure d'une molécule : la connaissance des groupes d'atomes d'hydrogène équivalents, couplée à celle des liaisons particulières, laisse en général peu de place à l'ambiguïté.

Il peut malgré tout être intéressant, pour recouper les informations, lever le doute sur une chute de transmittance un peu confuse dans un spectre IR, ou tout simplement s'épargner des recherches dans des voies sans issue, de faire précéder l'analyse de ces spectres par l'étude des sources d'informations suivantes :

- **Analyse centésimale** : diverses techniques permettent de déterminer les éléments dont est constituée une molécule, et même leurs pourcentages massiques. Nous verrons en exercice que ceux-ci permettent alors d'obtenir la formule brute de la molécule, et ce faisant de cadrer concrètement l'analyse des spectres.
- **Le nombre d'insaturations** : celui-ci peut être déterminé à partir de la seule formule brute d'une molécule (*cf.* exercice « Vers la prépa » du présent chapitre).
- **Tests caractéristiques** : un test positif à la liqueur de Fehling, par exemple, révèle la présence d'une fonction aldéhyde beaucoup plus vite qu'un spectre IR, et pour un coût significativement moins élevé.
- **Pouvoir rotatoire** : vous apprendrez en CPGE que les molécules dotées d'un carbone asymétrique possèdent la propriété de modifier la direction de polarisation d'une lumière polarisée. Ceci concerne uniquement les molécules dotées d'un carbone asymétrique, et se heurte au fait que deux atomes de carbone asymétriques dans des configurations différentes ont des pouvoirs rotatoires opposés : un mélange racémique ne dipose ainsi d'aucun pouvoir rotatoire global. Malgré tout, face par exemple à une réaction d'addition stéréospécifique, cette propriété peut être vue comme un test caractéristique des atomes de carbone asymétriques.

## Halte aux idées reçues

Expliquer, pour chacune des affirmations suivantes, pourquoi elle est fausse :

1. Le nombre de charge d'un atome neutre est égal à son numéro atomique.
2. Les éléments dans le tableau de classification périodique sont classés par masse molaire atomique croissante.

3. La liaison entre anions et cations au sein d'un solide ionique est fondamentalement différente d'une liaison de covalence.

4. Une molécule comportant des liaisons polarisées est forcément polaire.

5. Le spectre UV-visible d'une espèce chimique permet d'en déterminer le(s) groupement(s) fonctionnel(s).

## Du Tac au Tac

(Exercices sans calculatrice)

1. Le chlore possède deux isotopes stables : le chlore 35, de masse molaire atomique  $M_{35\text{Cl}} = 35,0 \text{ g.mol}^{-1}$ , et le chlore 37, de masse molaire atomique  $M_{37\text{Cl}} = 37,0 \text{ g.mol}^{-1}$ .

Sachant que la masse molaire atomique de l'élément chlore a pour valeur  $M_{\text{Cl}} = 35,5 \text{ g.mol}^{-1}$ , déterminer les abondances relatives de ses isotopes stables dans la nature.

2. Le fer est l'élément caractérisé par le numéro atomique  $Z = 26$ . Expliquer pourquoi existent les ions fer (II)  $\text{Fe}^{2+}$  et fer (III)  $\text{Fe}^{3+}$  ?

3. Déterminer les représentations topologiques de toutes les molécules de formule brute  $\text{C}_3\text{H}_6\text{O}$ .

4. Écrire un schéma de Lewis pour chacune des molécules et ions suivants :

– le dioxygène  $\text{O}_2$  et l'ozone  $\text{O}_3$  (cette dernière n'étant pas cyclique) ;

– l'ion azoture  $\text{N}_3^-$  et l'ion nitronium  $\text{NO}_2^+$ .

5. Déterminer les signes des charges partielles portées respectivement par les atomes engagés dans chacune des liaisons suivantes, et les classer par polarisation croissante :

- |         |          |
|---------|----------|
| (a) O-H | (d) H-Br |
| (b) O-C | (e) C-H  |
| (c) O-O | (f) C-Cl |

On donne les électronégativités des éléments :  $\chi_{\text{H}} = 2,2$ ,  $\chi_{\text{C}} = 2,6$ ,  $\chi_{\text{O}} = 3,4$ ,  $\chi_{\text{Cl}} = 3,2$ ,  $\chi_{\text{Br}} = 3,0$ .

## Vers la prépa

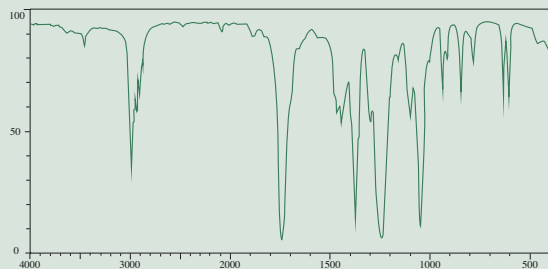
On considère une molécule A, au sujet de laquelle on a pu obtenir les informations suivantes :

Masse molaire :  $M_A = 88,0 \text{ g.mol}^{-1}$

Pourcentages massiques :

- Carbone :  $x_{\text{C}} = 54,5 \%$
- Hydrogène :  $x_{\text{H}} = 9,1 \%$
- Oxygène :  $x_{\text{O}} = 36,4 \%$

Spectre IR :



1. Déterminer la formule brute de la molécule A.

2. Montrer que le nombre  $N_i$  d'insaturations dans une molécule est donné par la relation :

$$N_i = 1 + N_4 + \frac{N_3}{2} - \frac{N_1}{2}$$

où  $N_1$ ,  $N_3$  et  $N_4$  représentent les nombres respectifs d'atomes monovalents, trivalents et tétravalents au sein de la molécule en question.

3. En déduire le nombre d'insaturations de A.

4. Déterminer les liaisons présentes dans cette molécule. En admettant qu'elle comporte un unique groupement fonctionnel, préciser la nature de celui-ci.

5. Donner les représentations topologiques et nommer tous les isomères envisageables à ce stade de l'analyse.

# Corrigés

## Halte aux idées reçues

**1.** Le nombre de charge  $Z$  d'un objet, quel qu'il soit, représente le nombre de fois où cet objet contient la charge électrique élémentaire  $e$ . Ainsi le nombre de charge d'un proton est-il  $+1$ , et celui d'un neutron,  $0$ . Mais on peut également définir le nombre de charge d'un électron :  $-1$ , ou même d'un objet ne se résumant pas à une particule élémentaire. C'est ainsi que le nombre de charge d'un atome neutre est par définition nul.

Le numéro atomique, quant à lui, caractérise l'élément en ce sens qu'il représente le nombre de protons du noyau d'un atome de cet élément, donnée qui définit fondamentalement l'élément chimique. Ainsi le numéro atomique d'un élément correspond bien à un nombre de charge, non pas celui de l'atome dans son entier, mais celui de son seul noyau.

**2.** Le tableau de classification périodique (TCP) des éléments repose sur la reproduction, à intervalles prévisibles, de propriétés chimiques analogues pour des éléments différents. Ces propriétés chimiques sont liées au cortège électronique des atomes de ces éléments, qui sont eux-mêmes tributaires du numéro atomique de ces éléments, soit encore du nombre de charge de leur noyau. L'ordonnancement des éléments dans le TCP repose donc à la base sur les caractéristiques électriques des noyaux atomiques.

La masse molaire des atomes, quant à elle, repose sur les isotopes de ces atomes (les nombres de neutrons et l'abondance relative de chacun dans la nature). La masse molaire n'est donc pas une affaire de charge électrique, mais uniquement de masse. Il se trouve que globalement, la masse molaire suit approximativement le même sens de variation que le numéro atomique : les nombres de protons et de neutrons croissent en général dans le même sens, aboutissant à des masses molaires qui croissent dans le même sens que le numéro atomique.

Mais ceci ne doit pas faire perdre de vue que la considération de base dans l'établissement du TCP tient à la charge électrique. Par ailleurs, plusieurs exceptions mettent en défaut l'affirmation proposée. Ainsi, par exemple, la masse molaire de l'argon est supérieure à celle du potassium, malgré le fait qu'il précède ce dernier dans le TCP. Idem pour le cobalt et le nickel, le tellure et l'iode, le thorium et le protactinium, l'uranium et le neptunium...

**3.** On a souvent tendance à dire qu'au sein d'un corps moléculaire, la liaison à l'œuvre entre les atomes est une liaison de covalence, tandis qu'au sein d'un corps ionique, c'est l'interaction coulombienne qui assurerait la cohésion des atomes. Il importe de bien comprendre que ces deux liaisons, quels que

soient les noms que l'on choisit de leur donner, reposent sur une même réalité : l'attraction mutuelle des particules dont les charges électriques sont de signes opposés.

Dans le cas des liaisons au sein d'une molécule, la mise en commun par deux atomes d'un électron externe chacun pour assurer une liaison de covalence n'est pas autre chose qu'un équilibre entre les attractions des noyaux pour les électrons, et les répulsions entre noyaux et entre nuages électroniques. La liaison de covalence elle-même ne résulte donc pas d'autre chose qu'une interaction coulombienne.

Réciproquement, ce que l'on qualifie de liaison ionique n'est en fait guère autre chose qu'une liaison de covalence polarisée à outrance. Vous avez vu que les atomes peuvent être caractérisés, entre autres choses, par leur électronégativité, grandeur quantifiant leur propension à accaparer les électrons de toute liaison covalente dans laquelle ils se trouvent engagés. Celle-ci peut en fait se résumer à la tendance qu'aura cet atome vis-à-vis des électrons, pour satisfaire à la règle de l'octet.

Il n'est que de citer l'exemple du solide ionique sans doute le plus connu au monde : le chlorure de sodium  $\text{NaCl}$  (sel de cuisine). L'atome neutre de chlore n'a besoin que d'un seul électron en plus pour satisfaire à la règle de l'octet : il est de ce fait très avide d'électrons, ce qui se traduit par une grande électronégativité. À l'opposé, l'atome neutre de sodium est encombré d'un électron de trop pour satisfaire à cette même règle. Il est donc naturellement poussé à s'en débarrasser, ce que traduit sa très faible électronégativité. Associé à un atome de chlore, ils vont établir une liaison de covalence, mais le nuage électronique la constituant sera beaucoup plus concentré autour de l'atome de chlore, qu'autour de celui de sodium. C'est ce que l'on entend par liaison polarisée.

Mais s'il est vrai que cette polarisation autorise une modélisation en termes d'attraction coulombienne plus simple que dans le cas d'une liaison non polarisée, il serait en revanche faux de croire que l'interaction fondamentale en jeu diffère.

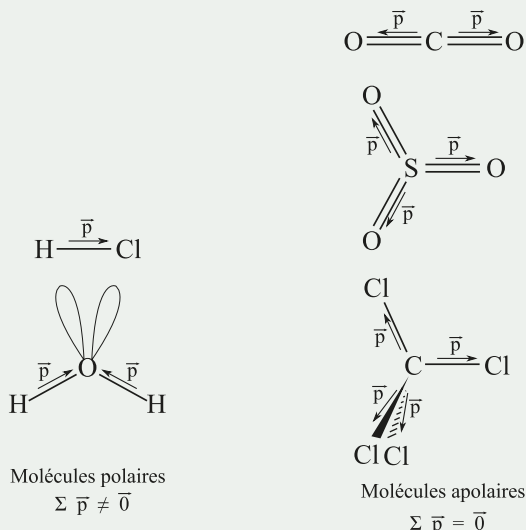
**4.** Une molécule est dite polaire lorsqu'elle possède un moment dipolaire électrique global, c'est-à-dire que les centres de charge respectifs (analogues du centre de masse, à ceci près que les particules sont pondérées non par leurs masses respectives mais par leurs charges électriques) des charges positives et négatives ne coïncident pas. Dans les faits, le moment dipolaire électrique d'une molécule est égal à la somme vectorielle des moments dipolaires électriques des différentes liaisons qu'elle comporte.

Ainsi, si l'on considère une molécule diatomique polarisée telle que le chlorure d'hydrogène  $\text{HCl}$ , elle est effectivement polaire, du fait même que l'unique liaison qu'elle comporte est polarisée.

Si l'on considère la molécule d'eau  $\text{H}_2\text{O}$ , celle-ci comporte deux liaisons polarisées, et la molécule présentant une géométrie coudée, les moments dipolaires respectivement associés à ces deux liaisons ne possèdent pas la même direction : leur somme ne peut donc pas aboutir au vecteur nul (tout comme deux forces non colinéaires ne peuvent pas se compenser entre elles), et la molécule d'eau est donc également polaire.

Prenons à présent le cas du dioxyde de carbone  $\text{CO}_2$  : celle-ci comporte deux doubles liaisons  $\text{C}=\text{O}$  extrêmement polarisées. Pourtant, la géométrie de cette molécule étant linéaire (le carbone ne comporte pas de doublet non liant), leurs moments dipolaires électriques s'opposent l'un à l'autre et, résultant chacun de la même liaison, ont même valeur et s'annulent donc entre eux. Ainsi le dioxyde de carbone, bien qu'étant le siège de liaisons très polarisées, se trouve être une molécule apolaire.

Il est intéressant de noter que cette capacité à être apolaire en dépit de liaisons polarisées se retrouve notamment dans toutes les situations où un atome central dépourvu de doublet non liant est entouré d'atomes identiques entre eux, sans que la géométrie soit nécessairement linéaire (trioxyde de soufre  $\text{SO}_3$ , dans laquelle le soufre est trigonal), ni même plane (tétrachlorométhane, dans laquelle le carbone est tétragonal) :



On remarque au passage que le soufre ne satisfait pas ici à la règle de l'octet. Vous verrez en effet en CPGE que, dès la troisième ligne du tableau de classification périodique, les atomes peuvent procéder à ce que l'on appelle une promotion de valence, qui leur permet de devenir hypervalents, c'est-à-dire de s'entourer de plus de quatre doublets d'électrons.

5. Nous avons vu que le spectre UV-visible n'apportait aucune information sur la structure de la molécule : trop grossier par rapport aux liaisons chimiques, il se contente de donner une signature de l'espèce absorbante. Mais cette signature, si elle est éventuellement identifiable (et encore, si l'on ignore tout de l'espèce qui l'a donnée, la bataille n'est que rarement gagnée), est en revanche illisible. Nous pourrions peut-être reconnaître la molécule si nous la voyons de nouveau, mais nous n'en connaissons toujours pas l'anatomie.

Tout au plus peut-on dire, face à un spectre présentant une particularité caractéristique d'une espèce, qu'il permet d'identifier cette espèce, et incidemment les groupements fonctionnels qu'elle porte, si l'on connaît déjà sa nature par ailleurs. Mais l'identification, pièce par pièce, de la structure d'une molécule, ne peut se faire qu'à travers ses spectres IR et RMN, souvent utilisés conjointement d'ailleurs.

## Du Tac au Tac

1. Notons  $x = \frac{n_{^{35}\text{Cl}}}{n_{\text{Cl}}}$  la fraction en quantité de matière (on parle de fraction molaire) du chlore 35 dans un échantillon de chlore naturel. Nous pouvons alors affirmer, puisque le chlore ne possède que deux isotopes stables, que si l'on prend une quantité de matière  $n_{\text{Cl}}$  de chlore, celle-ci sera constituée de :

- $n_{^{35}\text{Cl}} = x n_{\text{Cl}}$  de chlore 35 ;
- $n_{^{37}\text{Cl}} = n_{\text{Cl}} - n_{^{35}\text{Cl}} = (1-x)n_{\text{Cl}}$  de chlore 37.

La masse molaire atomique d'un élément chimique étant par définition le coefficient de proportionnalité entre la masse d'un échantillon d'atomes de cet élément pris à l'état naturel (c'est-à-dire incluant tous ses isotopes), et la quantité de matière d'atomes dans cet échantillon, il vient :

$$M_{\text{Cl}} = \frac{m_{\text{Cl}}}{n_{\text{Cl}}} = \frac{m_{^{35}\text{Cl}} + m_{^{37}\text{Cl}}}{n_{\text{Cl}}} = \frac{n_{^{35}\text{Cl}} M_{^{35}\text{Cl}} + n_{^{37}\text{Cl}} M_{^{37}\text{Cl}}}{n_{\text{Cl}}}$$

Soit encore :

$$M_{\text{Cl}} = \frac{n_{^{35}\text{Cl}}}{n_{\text{Cl}}} M_{^{35}\text{Cl}} + \frac{n_{^{37}\text{Cl}}}{n_{\text{Cl}}} M_{^{37}\text{Cl}}$$

Nous reconnaissons les fractions molaires du chlore 35 et du chlore 37, ce qui nous permet d'écrire :

$$M_{\text{Cl}} = x M_{^{35}\text{Cl}} + (1-x) M_{^{37}\text{Cl}}$$

Ce qui revient au final à :

$$x = \frac{M_{^{37}\text{Cl}} - M_{\text{Cl}}}{M_{^{37}\text{Cl}} - M_{^{35}\text{Cl}}} = \frac{1,5 \text{ g.mol}^{-1}}{2,0 \text{ g.mol}^{-1}} = 0,75$$

Nous en concluons que le chlore naturel est constitué de 75 % de chlore 35, et le complémentaire, soit 25 %, de chlore 37.

2. La structure électronique de l'atome neutre de fer, de numéro atomique  $Z = 26$ , est :

$$1s^2 2s^2 2p^6 3s^2 3p^6 4s^2 3d^6$$

Pour obtenir celle des ions associés au fer, on commence par vider les niveaux de  $n$  le plus élevé, c'est-à-dire  $n = 4$  ici. On en déduit les structures électroniques des ions  $\text{Fe}^{2+}$  et  $\text{Fe}^{3+}$  :

- pour  $\text{Fe}^{2+}$  :  $1s^2 2s^2 2p^6 3s^2 3p^6 3d^6$  ;
- pour  $\text{Fe}^{3+}$  :  $1s^2 2s^2 2p^6 3s^2 3p^6 3d^5$ .

**Remarque :** l'existence de l'ion  $\text{Fe}^{2+}$  qui correspond au vidage total du niveau  $4s$  était assez intuitive. Pour justifier l'existence de l'ion  $\text{Fe}^{3+}$ , il faut savoir qu'une sous-couche de type  $d$  pleine, ou à moitié pleine, confère une certaine stabilité à l'entité qui la possède. Ainsi, on constatera un certain nombre d'exceptions à la règle de Klechkowski qui s'expliqueront par ce biais. Notamment, l'atome neutre de chrome ( $Z=24$ ) admet une structure électronique  $1s^2 2s^2 2p^6 3s^2 3p^6 4s^1 3d^5$  au lieu de  $1s^2 2s^2 2p^6 3s^2 3p^6 4s^2 3d^4$  comme prévu par la règle de Klechkowski. Le même genre d'irrégularité se retrouve notamment pour les éléments or et argent. Vous pourrez vous en convaincre en examinant leurs structures électroniques respectives...

3. Il n'existe malheureusement pas de formule simple permettant de déterminer le nombre d'isomères de constitution que l'on peut développer à partir d'une formule brute donnée. Pour minimiser le risque d'en oublier, il est donc souhaitable de procéder avec un peu de méthode.

Dans un premier temps, il peut être bon de chercher quelques molécules à tâtons pour en dégager les caractéristiques communes, notamment l'éventuelle présence d'insaturations (c'est-à-dire de liaisons multiples et/ou de cycles). Il est bon de savoir que le nombre d'insaturations d'une molécule ne dépend que de sa formule brute ; ce point sera développé dans l'exercice « Vers la Prépa ». Dans le cas présent, on voit rapidement que la molécule comporte systématiquement une insaturation.

On peut donc déjà envisager 3 situations, chacune excluant les autres (une seule insaturation possible) :

- (a) La molécule comporte un cycle.
- (b) La molécule comporte une liaison  $\text{C}=\text{O}$ .
- (c) La molécule comporte une liaison  $\text{C}=\text{C}$ .

Pour chacune de ces hypothèses, on peut alors développer plusieurs scénarios, compte tenu des 3 atomes C et de l'unique atome O en jeu :

1. Cas du cycle : il faut au moins 3 atomes. On peut donc envisager :

- a) Un cycle à 3 C : dans ce cas il reste un O à liaisons simples, et des H. On aura donc un groupement  $-\text{OH}$  à placer sur l'un des C (qui du reste sont équivalents entre eux), et des H à placer sur les liaisons vacantes.
- b) Un cycle à 2 C et 1 O : il ne reste qu'un C et des H, donc un groupement méthyle et des H qu'il suffit de distribuer sur les liaisons vacantes, et qui ne donnent lieu à aucune équivoque.
- c) Un cycle à 3 C et 1 O : il ne reste alors qu'à distribuer les H sur les liaisons vacantes.

2. Cas de la double liaison  $\text{C}=\text{O}$  : l'oxygène est saturé, il reste donc 2 C et des H à placer. Deux possibilités seulement se présentent :

- a) Deux groupements méthyle.
- b) Un groupement éthyle et un H.

3. Cas de la double liaison  $\text{C}=\text{C}$  : le carbone restant ne peut guère former autre chose qu'un groupement méthyle. Quant à l'oxygène, il peut au choix :

- a) Se combiner avec 1 H pour former un groupement  $-\text{OH}$ . Le dernier C forme alors un groupement méthyle avec 3 H, et est associé à l'un ou l'autre de ceux de la liaison éthylénique. Le groupement  $-\text{OH}$  peut alors être fixé :
  - i) Sur le C tétragonal.
  - ii) Sur le C de la liaison éthylénique, qui porte déjà le groupement méthyle.
  - iii) Sur l'autre C de la liaison éthylénique. On doit alors distinguer deux diastéréoisomères : l'un Z, l'autre E.
- b) Établir un pont entre l'un des C de la liaison éthylénique et le groupement méthyle.

#### 1. Molécules cycliques

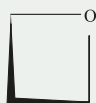
##### 1.(a) Cycle à 3 C :



##### 1.(b) Cycle à 2C, 1O :

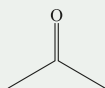


##### 1.(c) Cycle à 3C, 1O :



#### 2. Molécules avec liaison $\text{C}=\text{O}$

##### 2.(a) Deux méthyle :



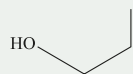
##### 2.(b) Éthyle et H :



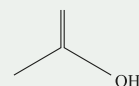
#### 3. Molécules avec liaison $\text{C}=\text{C}$

##### 3.(a) Avec groupement $-\text{OH}$ :

###### 3.(a).(i) Sur le C tétragonal :



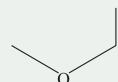
###### 3.(a).(ii) Sur le C trigonal méthylé :



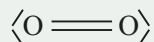
###### 3.(a).(iii) Sur le C trigonal non méthylé :



##### 3.(b) Avec liaison $-\text{C}-\text{O}-\text{C}-$ :

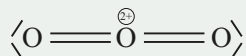


4. La molécule de dioxygène possède  $2 \times 6 = 12$  électrons de valence, donc 6 doublets à répartir. Le caractère divalent de l'oxygène pousse à établir une liaison double entre les deux atomes d'oxygène ; on n'a plus alors qu'à satisfaire à la règle de l'octet en ajoutant deux doublets non-liants sur chacun des atomes d'oxygène. On obtient ainsi le schéma de Lewis suivant :



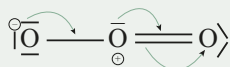
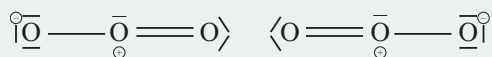
Chacun des deux atomes d'oxygène possède naturellement 6 électrons de valence ; dans cette molécule, ils sont chacun entouré de 6 électrons (4 venant des doublets non-liants et 1 pour chaque doublet liant). Aucune charge formelle n'apparaît donc.

Pour ce qui est de la molécule d'ozone, elle dispose de  $3 \times 6 = 18$  électrons au total, donc de 9 doublets. Par analogie avec ce qui précède, on pourrait être tenté d'établir des liaisons doubles entre les atomes d'oxygène, puis à compléter avec des doublets non liants pour satisfaire à la règle de l'octet sur les atomes terminaux. On aboutirait ainsi à la structure suivante, dans laquelle l'atome central porte une charge formelle  $2+$  non compensée :



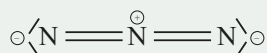
Cependant la formule de cet édifice est  $\text{O}_3^{2+}$ . On constate du reste qu'elle ne comporte que 8 doublets sur les 9 annoncés (logique, puisque 2 électrons manquent à l'appel).

Pour retrouver la bonne structure, remplaçons l'une des doubles liaisons par une simple et ajoutons des doublets non-liants en conséquence. On aboutit ainsi à l'une des structures suivantes, lesquelles sont totalement équivalentes (on a représenté en-dessous les déplacements internes d'électrons qui permettent de passer de l'une à l'autre des formes) :



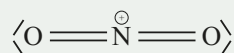
**Remarque :** on constate sur cet exemple que des charges formelles peuvent apparaître même dans un édifice globalement neutre. Il est donc important de bien vérifier, lorsque vous écrivez une structure de Lewis, si le nombre d'électrons dont est effectivement entouré chaque atome correspond bien à son nombre d'électrons de valence. Notons par ailleurs que plusieurs schémas de Lewis peuvent convenir pour une seule et même entité. Les différentes formes écrites sont alors dites liées par une relation de mésomérie, notion qui sera développée en CPGE.

Pour l'ion azoture  $\text{N}_3^-$  on a  $(3 \times 5) + 1 = 16$  électrons de valence à placer (5 par atome d'azote et un de plus traduisant la charge de l'ion), soit 8 doublets. L'azote N, élément de la deuxième ligne, doit satisfaire à la règle de l'octet ; la structure symétrique suivante, qui engage trois charges formelles, mène bien à une charge globale de  $-1$  :



Pour l'ion nitronium  $\text{NO}_3^+$  on a  $5 + (2 \times 6) - 1 = 16$  électrons de valence à placer, soit encore 8 doublets. En reprenant une

structure formellement analogue à la précédente, on aboutit au schéma suivant :



5. En règle générale, parvenu(e) à ce stade de votre formation, il est une poignée de connaissances que vous êtes supposé(e) avoir sur vous : l'acide éthanóïque est un acide carboxylique, la soude est une base, l'éthanol est un alcool... Parmi celles-ci, se trouve également le fait que les atomes de certains éléments ont une électronégativité particulièrement vigoureuse tandis que d'autres sont plutôt faibles en la matière. On retiendra en particulier que l'oxygène est l'élément le plus électronégatif, derrière le fluor (que du reste on ne rencontre que rarement), et qu'il est suivi des halogènes pris en descendant dans le tableau de classification périodique. Par ailleurs les atomes de carbone et d'hydrogène sont faiblement électronégatifs, l'hydrogène étant cependant encore plus faible que le carbone.

Ici, donc, avant même d'aller chercher le tableau de classification périodique, on pouvait d'ores et déjà affirmer que :

- (a) O-H est une liaison très polarisée, O portant une charge partielle négative et H son complément positif.
- (b) O-C est une liaison très polarisée, mais moins que la précédente, et O porte une charge partielle négative tandis que C porte son complément positif.
- (c) O-O est une liaison non polarisée, les deux O ayant la même propension à accaparer les électrons de toute liaison dans laquelle ils se trouvent engagés.
- (d) H-Br est une liaison très polarisée, mais moins que O-H, et Br porte une charge partielle négative tandis que H porte son complément positif.
- (e) C-H est une liaison peu polarisée, C portant une faible charge partielle négative et H son complément positif.
- (f) C-Cl est une liaison très polarisée, mais moins que O-C, et Cl porte une charge partielle négative tandis que C porte son complément positif.

Si l'on souhaite ensuite un classement précis de ces liaisons suivant leurs niveaux de polarisation, il devient cette fois nécessaire de mener une étude quantitative. Celle-ci reste cependant simple : il nous suffit de calculer pour chaque liaison la différence d'électronégativité entre l'atome le plus électronégatif, et celui qui l'est le moins. Plus cette différence est élevée, plus la liaison est polarisée. Nous n'avons plus qu'à classer les liaisons par valeur croissante de cette différence. Nous trouvons alors :

| Liaison      | O-H | O-C | O-O | H-Br | C-H | C-Cl |
|--------------|-----|-----|-----|------|-----|------|
| $ \Delta E $ | 1,2 | 0,8 | 0,0 | 0,8  | 0,4 | 0,6  |

Nous en concluons donc qu'en termes de polarisation, on a :



## Vers la prépa

1. Nous savons que la fraction massique d'un élément Y dans un échantillon ech d'une molécule A s'exprime :

$$x_Y = \frac{m_{Y,ech}}{m_{A,ech}} = \frac{n_{Y,ech} \times M_Y}{n_{A,ech} \times M_A} = \frac{n_{Y,ech}}{n_{A,ech}} \times \frac{M_Y}{M_A}$$

$$= N_Y \times \frac{M_Y}{M_A}$$

où  $N_Y$  représente le nombre d'atomes de Y par une molécule de A.

Nous en déduisons :

$$N_Y = x_Y \times \frac{M_A}{M_Y}$$

Les applications numériques donnent alors  $N_C = 4$ ,  $N_H = 8$  et  $N_O = 2$ , d'où une formule brute pour la molécule considérée  $C_4H_8O_2$ .

**2.** Le nombre d'insaturations d'une molécule est le nombre de doublets liants surnuméraires (liaisons multiples) et de fermetures (cycles) au sein de cette molécule. Il suffit pour en calculer le nombre de partir d'un atome de carbone (offrant 4 liaisons), et de compter pour chaque atome qu'on lui lie, combien de liaisons supplémentaires il permet de développer, dans l'hypothèse de liaisons simples et ne refermant pas la molécule sur elle-même uniquement. Ainsi par exemple :

- L'ajout d'un atome tétravalent (carbone notamment) ajoute 2 liaisons à la molécule (3 nouvelles liaisons, moins celle que sa mise en place supprime sur le radical auquel il est lié). C'est ainsi que pour un atome de carbone, 4 liaisons sont disponibles, mais l'ajout d'un deuxième atome fait passer ce nombre à 6, pour 3 atomes on passe à 8...
- L'ajout d'un atome trivalent (azote notamment) ajoute 1 liaison à la molécule (2 nouvelles liaisons, moins celle que sa mise en place supprime sur le radical auquel il est lié).
- L'ajout d'un atome divalent (oxygène notamment) n'ajoute aucune liaison à la molécule (il se contente de reconduire celle qu'il a utilisée pour se fixer sur le radical auquel il est lié).
- L'ajout d'un atome monovalent (hydrogène ou halogène) neutralise 1 liaison de la molécule (il la monopolise pour s'y fixer, sans en ouvrir de nouvelle).

Le nombre d'atomes monovalents est égal au nombre de liaisons disponibles sur les autres. Il s'ensuit que, pour une molécule ne comprenant aucune insaturation, on doit avoir l'égalité :

$$(2N_4 + 2) + N_3 = N_1$$

où  $N_4$ ,  $N_3$  et  $N_1$  désignent les nombres respectifs d'atomes tétravalents, trivalents et monovalents dans la molécule.

La présence de toute insaturation dans la molécule aura pour effet de neutraliser deux liaisons sur le radical (2 atomes qui seront liés chacun à l'autre une fois de plus, plutôt que d'offrir la place à des atomes monovalents). Une insaturation a donc le même effet, en termes de pertes de liaisons disponibles, que la fixation de deux atomes monovalents.

En notant  $N_i$  le nombre d'insaturations dans la molécule, la relation précédente devient :

$$2N_4 + 2 + N_3 = N_1 + 2N_i \Leftrightarrow N_i = 1 + N_4 + \frac{N_3}{2} - \frac{N_1}{2}$$

**3.** Dans le cas présent, on constate ainsi, avec  $N_4 = N_C = 4$ ,  $N_3 = 0$  et  $N_1 = N_H = 8$ , que l'on trouve  $N_i = 1$ .

**4.** Le spectre IR fourni présente cinq bandes exploitables, occupant les intervalles approximatifs suivants :

- Une entre 2 900 et 3 000  $\text{cm}^{-1}$ , forte.
- Une entre 1 700 et 1 800  $\text{cm}^{-1}$ , forte.
- Une entre 1 450 et 1 500  $\text{cm}^{-1}$ , forte.
- Une entre 1 150 et 1 350  $\text{cm}^{-1}$ , forte.
- Une entre 1 000 et 1 150  $\text{cm}^{-1}$ , forte.

On peut attribuer les origines de ces bandes, respectivement :

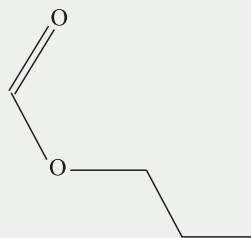
- À des liaisons  $C_{tet}-H$ , à l'œuvre entre 2 800 et 3 000  $\text{cm}^{-1}$ , seule bande forte dans ce domaine de nombres d'onde, parmi celles proposées.
- À une liaison  $C=O$ , bande forte à l'œuvre autour de 1 700  $\text{cm}^{-1}$  ; on peut ajouter qu'il s'agit vraisemblablement de celle d'un ester, seul cas, parmi ceux proposés, où la bande dépasse significativement les 1 700  $\text{cm}^{-1}$ .
- À des liaisons  $C_{tet}-H$ , à l'œuvre entre 1 415 et 1 470  $\text{cm}^{-1}$ , seule bande forte dans ce domaine de nombres d'onde, parmi celles proposées. Elle corrobore la première bande mentionnée ci-dessus.
- À une liaison  $C_{tet}-O$ , à l'œuvre entre 1 050 et 1 450  $\text{cm}^{-1}$ . La liaison  $C_{tet}-C_{tet}$  ne suffit pas à expliquer l'étendue de la bande, qui monte de manière visible au-delà des 1 300  $\text{cm}^{-1}$ , quand la liaison  $C_{tet}-C_{tet}$  est supposée ne pas dépasser 1 250  $\text{cm}^{-1}$ .
- À des liaisons  $C_{tet}-C_{tet}$ , seules proposées dans le domaine de nombres d'onde où s'inscrit cette dernière bande.

La présence d'une liaison  $C=O$  et d'une liaison  $C-O$  pour un unique groupement fonctionnel ne peut convenir qu'à un acide carboxylique ou à un ester. Cependant :

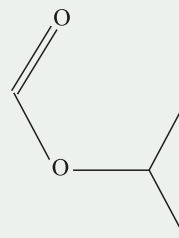
- Comme nous l'avons dit précédemment, la bande observée vers 1 700  $\text{cm}^{-1}$  indique plutôt un ester.
- S'il s'agissait d'une acide carboxylique, nous devrions également observer une bande large et forte (donc bien visible) entre 2 500 et 3 200  $\text{cm}^{-1}$ . Or nous n'observons rien de ce type.

Nous en déduisons qu'il s'agit d'un ester.

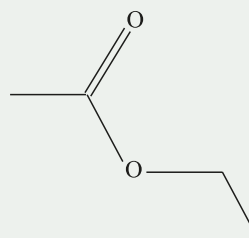
**5.** Il existe 4 isomères d'esters de formule brute  $C_4H_8O_2$  :



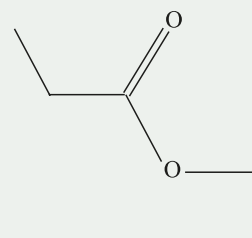
Méthanoate de propyle



Méthanoate de 1-méthyle éthyle



Éthanoate d'éthyle



Propanoate de méthyle



### 2.1 Mesure de l'abondance d'un échantillon de matière

#### — De quelles grandeurs dispose-t-on pour mesurer l'abondance d'un échantillon de matière ?

La chimie se propose d'étudier la matière et ses transformations. Elle ouvre la voie à d'innombrables applications qui, une fois mises en œuvre concrètement, vont poser la question de l'abondance de la matière. De quelle quantité de tel ou tel produit a-t-on besoin pour mener à bien un projet, un chantier, etc. De quelles quantités de réactif a-t-on besoin pour synthétiser ces produits ? À partir de quelle quantité telle substance améliore-t-elle ou au contraire compromet-elle la bonne santé d'un organisme vivant, le fonctionnement d'un écosystème, etc.

La mesure de cette abondance est donc cœur de la chimie appliquée. Concrètement, elle peut se faire à travers **3 grandeurs essentielles**.

- **La masse  $m$**  d'un corps, mesure à la fois :
  - l'impact de la force gravitationnelle sur ce corps (masse dite *grave*) ;
  - la résistance de ce corps à une modification de son vecteur vitesse par rapport à un référentiel galiléen, lorsqu'il subit l'exercice d'une force (masse dite *inerte*).

Bien qu'issues de deux lois distinctes (loi de la gravitation pour la première, 2<sup>e</sup> loi de Newton pour la seconde), aucune différence de valeur entre ces deux masses n'a à ce jour pu être mise en évidence, et nous parlerons simplement de LA masse en général, grandeur quantifiant le point auquel le corps considéré est **lourd**.

L'USI de masse est le kilogramme (kg). Cependant les échantillons employés dans un laboratoire de chimie étant en général plutôt de l'ordre du gramme, c'est le plus souvent cette dernière unité que sera utilisée.

- **Le volume  $V$**  d'un corps, mesure la quantité d'espace qu'occupe le corps en question. Il quantifie le point auquel le corps considéré est **encombrant**.

L'USI de volume est le mètre cube ( $m^3$ ). Cependant ici encore les échantillons manipulés dans le cadre d'un laboratoire atteignent rarement des volumes de cet ordre, et l'on rencontrera plus fréquemment le litre :

$$1 \text{ L} = 10^{-3} \text{ m}^3 = (10^{-1} \text{ m})^3 = 1 \text{ dm}^3$$

ou encore le millilitre :

$$1 \text{ mL} = 10^{-3} \text{ L} = 10^{-6} \text{ m}^3 = (10^{-2} \text{ m})^3 = 1 \text{ cm}^3$$

- **La quantité de matière  $n$**  d'un certain type d'objet (généralement des atomes, des ions ou des molécules) dans un corps, mesure le nombre de moles, c'est-à-dire le nombre de paquets de  $6,022 \cdot 10^{23}$  objets de ce type, dans le corps considéré.

La mole est une unité de dénombrement, définie dans le but de compter les entités atomiques ou moléculaires en paquets suffisamment grands pour que le nombre de ces paquets soit facilement manipulable.

Il s'agit d'une USI, dont le contenu est résumé par la constante d'Avogadro :

$$N_A = 6,022 \cdot 10^{23} \text{ mol}^{-1}$$

Cependant la valeur détaillée du nombre de molécules dans un échantillon de matière est à la fois pénible à manipuler (nombre énorme) et sa connaissance à l'unité près, totalement inutile.

## — Quelles grandeurs intermédiaires permettent de passer de l'une de ces grandeurs à une autre ?

Parmi les grandeurs évoquées ci-dessus, seule la quantité de matière nous intéresse réellement en chimie. En effet, une transformation chimique engage les molécules en **nombres** précis. C'est donc le dénombrement de ces molécules, non leur masse ou leur volume, qui vont conditionner les proportions dans lesquelles vont être consommés les réactifs et formés les produits.

Cependant, il est matériellement impossible de compter les molécules (trop petites et surtout beaucoup trop nombreuses) dans un échantillon de matière. Pour accéder à la quantité de matière d'une espèce, on va donc utiliser les 2 autres grandeurs (volume et/ou masse), accessibles à la mesure, et en déduire ensuite la quantité de matière correspondante, grâce à des **grandeurs intermédiaires**.

Ces grandeurs intermédiaires sont au nombre de 5.

- **La masse molaire  $M_X$**  d'une entité X, est le coefficient de proportionnalité entre la masse  $m_{X,\text{éch}}$  d'un échantillon de ces entités, et leur quantité de matière  $n_{X,\text{éch}}$  dans cet échantillon :

$$M_X = \frac{m_{X,\text{éch}}}{n_{X,\text{éch}}} \Leftrightarrow n_{X,\text{éch}} = \frac{m_{X,\text{éch}}}{M_X} \Leftrightarrow m_{X,\text{éch}} = n_{X,\text{éch}} \times M_X$$

- si l'échantillon considéré est constitué d'atomes appartenant tous au **même élément** chimique (cas d'atomes isolés), la masse molaire à prendre en compte est la **masse molaire atomique** de l'élément en question, qui intègre les abondances relatives des différents isotopes de celui-ci dans la nature. Les masses molaires atomiques sont consignées dans le **tableau de Mendeleïev** ;
- si l'échantillon considéré est constitué de molécules appartenant toutes à la **même espèce chimique** (cas d'un corps pur), il est nécessaire de **calculer la masse molaire moléculaire** de cette espèce. Celle-ci s'obtient en effectuant la somme des masses molaires atomiques des éléments qu'elle contient, chacune pondérée par le nombre d'atomes de son élément dans l'espèce considérée.
- Enfin, si l'échantillon considéré est un **mélange** (plusieurs types de molécules), on peut éventuellement définir une **masse molaire moyenne** de celui-ci s'appuyant sur les masses molaires des différentes espèces intervenant dans ce mélange, pondérées par leurs abondances respectives.

La masse molaire est donc utile lorsque l'on connaît la masse de l'échantillon et que l'on souhaite obtenir la quantité de matière associée (et réciproquement).

- **La masse volumique  $\rho_X$**  d'un corps pur X, est le coefficient de proportionnalité entre la masse  $m_{X,\text{éch}}$  d'un échantillon constitué de ce corps, et le volume  $V_{\text{éch}}$  occupé par cet échantillon :

$$\rho_X = \frac{m_{X,\text{éch}}}{V_{\text{éch}}} \Leftrightarrow V_{\text{éch}} = \frac{m_{X,\text{éch}}}{\rho_X} \Leftrightarrow m_{X,\text{éch}} = V_{\text{éch}} \times \rho_X$$

Les unités de masse et de volume utilisées varient selon les cas, et l'on ne saurait trop veiller à leur cohérence, ainsi qu'à la crédibilité des valeurs numériques obtenues.

La masse volumique est donc utile lorsque nous est fourni le volume d'un échantillon, en ce qu'elle nous permet de déterminer la masse de celui-ci, et ce faisant de nous ramener au cas précédent (masse, dont il ne reste plus qu'à déduire la quantité de matière associée au moyen de la masse molaire).

- **La concentration  $C_X$  en soluté X apporté** dans une solution, est le coefficient de proportionnalité entre la quantité de matière  $n_{X,\text{app}}$  de ce soluté dissout dans un échantillon de solution, et le volume  $V_{\text{solution}}$  de l'échantillon en question :

$$C_X = \frac{n_{X,\text{app}}}{V_{\text{solution}}} \Leftrightarrow V_{\text{solution}} = \frac{n_{X,\text{app}}}{C_X} \Leftrightarrow n_{X,\text{app}} = V_{\text{solution}} \times C_X$$

La concentration en soluté apporté est utile lorsque l'on souhaite connaître **la quantité de soluté à apporter** pour confectionner un certain volume de solution de concentration donnée.

- **La concentration effective  $[X]$  en une espèce solvatée X** dans une solution, est le coefficient de proportionnalité entre la quantité de matière  $n_{X,\text{eff}}$  de cette espèce effectivement présente dans un échantillon de solution, et le volume  $V_{\text{solution}}$  de ce dernier :

$$[X] = \frac{n_{X,\text{eff}}}{V_{\text{solution}}} \Leftrightarrow V_{\text{solution}} = \frac{n_{X,\text{eff}}}{[X]} \Leftrightarrow n_{X,\text{eff}} = V_{\text{solution}} \times [X]$$

La concentration effective est utile lorsque l'on souhaite connaître **la quantité effectivement présente** en une espèce solvatée dans un certain volume de solution.

La valeur de cette concentration effective coïncide souvent avec celle de la concentration en soluté apporté, mais pas toujours. Lorsqu'il se dissout dans le solvant, le soluté peut en effet :

- conserver son intégrité (soluté dit *moléculaire*), auquel cas une molécule de soluté apporté donne une molécule de soluté solvato ; les deux concentrations sont dans ce cas égales entre elles ;
  - se dissocier en un cation et un anion (soluté dit *ionique*), auquel cas on obtient 2 ions solvatés, dont les concentrations effectives sont égales à la concentration en soluté apporté ; cependant les espèces solvatées diffèrent du soluté de départ ;
  - se dissocier en un ou plusieurs cation(s) et/ou un ou plusieurs anion(s), auquel cas non seulement les espèces effectivement présentes diffèrent du soluté de départ, mais en outre leurs concentrations diffèrent de la concentration en soluté apporté.
- **Le volume molaire  $V_m$  (utilisable uniquement dans le cas d'un corps gazeux)**, est le coefficient de proportionnalité entre le volume  $V_{g,\text{éch}}$  occupé par un échantillon de gaz, et la quantité de matière  $n_{\text{mlc},\text{éch}}$  de molécules présentes dans cet échantillon :

$$V_m = \frac{V_{g,\text{éch}}}{n_{\text{mlc},\text{éch}}} \Leftrightarrow n_{\text{mlc},\text{éch}} = \frac{V_{g,\text{éch}}}{V_m} \Leftrightarrow V_{g,\text{éch}} = n_{\text{mlc},\text{éch}} \times V_m$$

Le volume molaire présente la spécificité de ne dépendre, en pratique, que de la température et de la pression du gaz, à l'exclusion notamment de sa nature chimique.

Dans le cadre du modèle du gaz parfait, on peut le calculer par la relation :

$$V_m = \frac{V_{g,éch}}{n_{mlc,éch}} = \frac{R \times T}{P}$$

avec :

- $R = 8,314$  USI, constante du gaz parfait ;
- $T$  la température absolue du gaz (en kelvin (K), sachant que  $T_{(K)} = T_{(°C)} + 273,15$ ) ;
- $P$  la pression (en pascal (Pa), sachant que  $1 \text{ bar} = 10^5 \text{ Pa}$ ) ;
- la valeur de  $V_m$  est alors obtenue en  $\text{m}^3 \cdot \text{mol}^{-1}$ .

Le volume molaire est donc utile lorsque l'on souhaite connaître **la quantité de gaz contenue** dans des conditions données de température et de pression.

**Remarque :** bien qu'il existe plusieurs grandeurs permettant de quantifier l'abondance d'un échantillon (masse, volume, quantité de matière), il importe de garder à l'esprit que ces grandeurs sont de natures différentes les unes des autres : elles sont certes proportionnelles entre elles, mais un kilogramme ne possède aucun équivalent absolu en litre, ni moins encore en moles. On évitera donc d'évoquer des « conversions » de l'une de ces grandeurs en une autre, ce terme étant réservé à un changement d'unité au sein d'un **même type de grandeur** : on peut convertir des kilogrammes en grammes et réciproquement, mais pas en litres par exemple. On pourra en revanche parler de **correspondance** entre deux grandeurs, à la rigueur d'équivalence, mais pas de conversion et SURTOUT PAS d'égalité.

## 2.2 Les méthodes physiques d'analyse d'une solution

### Qu'est-ce qu'une méthode d'analyse physique ?

L'un des objectifs majeurs de la chimie des solutions est de déterminer la concentration (on dit également le *titre*) d'une espèce chimique dans une solution (le plus souvent aqueuse, dans notre cas). Nous disposons pour ce faire de deux catégories de techniques :

- les **mesures physiques**, objet de la présente section, pour lesquelles on procède à la mesure, sur la solution, d'une **grandeur physique dépendant de la concentration effective** de l'espèce dosée dans cette solution. Connaissant la relation entre cette concentration et la grandeur physique en question, une mesure de celle-ci permet de remonter à la concentration recherchée. Ces méthodes ne consomment pas l'espèce dosée ;
- les **dosages chimiques**, également appelés **titrages**, consistent à **faire réagir l'espèce** dont on souhaite déterminer le titre (on dit également l'espèce à titrer, ou l'espèce titrée), **avec une autre espèce** (appelée espèce titrante) introduite en quantité connue, selon une réaction de stœchiométrie également connue. Ces méthodes sont destructives : l'espèce titrée est consommée au cours du titrage. Elles seront détaillées plus loin, dans le chapitre 4.

### Peut-on utiliser ces techniques pour déterminer directement la concentration d'une espèce en solution ?

Les méthodes d'analyse physique que nous allons voir présentent en général une précision assez médiocre en ce qu'elles sont tributaires de la qualité du matériel utilisé, de son entretien, ainsi

que des conditions dans lesquelles est menée la mesure. Il n'est ainsi pas rare de trouver des écarts entre valeur réelle et valeur mesurée (vérifiée, elle, avec une méthode chimique fiable) de 20 à 30 %, voire davantage.

En revanche, si l'on suit l'**évolution** de ces grandeurs au cours d'un titrage (on double donc une méthode chimique d'un suivi physique) on constate souvent des changements remarquables de ces grandeurs, à mesure que le titrage progresse. Ces changements découlent du changement de réactif limitant se produisant au franchissement de l'équivalence (cf. chapitre 4).

Ainsi, si l'on ne peut pas toujours se fier aveuglément aux valeurs prises par la grandeur mesurée, leur suivi permet en revanche d'obtenir la valeur du volume d'espèce titrante versé à l'équivalence avec une grande précision. Cette manière de procéder trouve tout son intérêt lorsque aucune des espèces engagées dans la réaction de dosage (titrée, titrante ou produits issus de la réaction de dosage) ne présente de couleur en solution aqueuse (ou bien au contraire si beaucoup d'entre elles sont colorées et qu'il n'est plus possible de distinguer les teintes les unes des autres). Comme nous devons pouvoir détecter la fin de la réaction de dosage, on peut alors substituer au critère colorimétrique un critère de changement d'allure dans l'évolution de la grandeur suivie.

## — Qu'est-ce que le pH ?

Le pH est une abréviation qui signifie **potentiel hydrogène**. Il s'agit d'une grandeur traduisant la quantité d'ions  $H^+$  disponibles par unité de volume dans une solution aqueuse. Les ions  $H^+$  se résument en fait à des protons, et ne peuvent exister seuls en solution aqueuse : ils ne s'y trouvent donc que combinés avec des molécules d'eau, sous forme d'ions oxonium  $H_3O^+$ . Le pH est donc une grandeur traduisant la concentration d'une solution aqueuse en ions  $H_3O^+$ .

Concrètement, il est défini pour de faibles concentrations ( $[H_3O^+] < 10^{-1} \text{ mol.L}^{-1}$  typiquement) par la relation :

|                                                             |                                  |                                      |
|-------------------------------------------------------------|----------------------------------|--------------------------------------|
| $pH = -\log \left( \frac{[H_3O^+]}{c_{\text{ref}}} \right)$ | pH<br>$[H_3O^+], c_{\text{ref}}$ | sans unité<br>en $\text{mol.L}^{-1}$ |
|-------------------------------------------------------------|----------------------------------|--------------------------------------|

**Remarque :** l'unité proposée ci-dessus pour les concentrations est la  $\text{mol.L}^{-1}$ . Notons cependant que la valeur du rapport  $\frac{[H_3O^+]}{c_{\text{ref}}}$ , sans unité, ne dépend pas du système d'unités adopté. La seule chose qui compte est que les valeurs de ces deux concentrations soient exprimées dans la même unité.

Comme dans le cas du quotient de réaction (cf. chapitre 3),  $c_{\text{ref}}$  est une concentration de référence dont la valeur est  $1 \text{ mol.L}^{-1}$ . Sa présence peut donc être omise, à condition de préciser que la valeur de  $[H_3O^+]$  est exprimée en  $\text{mol.L}^{-1}$ . On obtient alors la définition du pH sous la forme :

|                         |                  |                                             |
|-------------------------|------------------|---------------------------------------------|
| $pH = -\log ([H_3O^+])$ | pH<br>$[H_3O^+]$ | sans unité<br>valeur en $\text{mol.L}^{-1}$ |
|-------------------------|------------------|---------------------------------------------|

**Remarque :** bien que le pH n'ait pas d'unité, on exprime parfois ses variations en « unités de pH », notamment si l'on doit définir une échelle sur un graphique dont l'un des axes porte les valeurs du pH.

Il importe de noter quelques points concernant cette grandeur :

- Le pH est une **fonction décroissante** de  $[H_3O^+]$  : plus le milieu est acide (c'est-à-dire plus il contient d'ions oxonium), et plus les valeurs du pH seront faibles.
- Les variations du pH sont une **fonction logarithmique** de  $[H_3O^+]$ , c'est-à-dire que ses variations sont relativement lentes : il est nécessaire de multiplier  $[H_3O^+]$  par 10 pour abaisser le pH ne serait-ce que de 1 unité. Pour 2 unités de pH,  $[H_3O^+]$  doit être multipliée par 100, etc. Inversement, le pH augmente de 1 unité à chaque fois que  $[H_3O^+]$  est divisée par 10.

Notons par ailleurs la relation réciproque, permettant d'exprimer la valeur en  $\text{mol.L}^{-1}$  de  $[\text{H}_3\text{O}^+]$  dans une solution de pH connu :

|                                            |                          |                               |
|--------------------------------------------|--------------------------|-------------------------------|
| $[\text{H}_3\text{O}^+] = 10^{-\text{pH}}$ | pH                       | sans unité                    |
|                                            | $[\text{H}_3\text{O}^+]$ | valeur en $\text{mol.L}^{-1}$ |

## — Qu'est-ce qu'une solution acide, basique ou neutre ?

Les premiers contacts que vous avez eus avec le concept de pH, au collège puis dans les classes de 2<sup>nde</sup> et de 1<sup>re</sup>, ne portaient pas sur des **espèces chimiques** qui étaient des acides ou des bases, mais sur des **solutions aqueuses** auxquelles on associait les adjectifs « acide », « basique » ou « neutre ».

Vous aviez alors appris qu'une solution était :

- acide si son pH était inférieur à 7,0 ;
- neutre si son pH était égal à 7,0 ;
- basique si son pH était supérieur à 7,0.

Nous verrons au chapitre 3 que même dans une eau où n'ont été apportés aucun acide ni aucune base (autre que l'eau elle-même), la réaction d'autoprotolyse de l'eau génère malgré tout des ions oxonium et des ions hydroxyde. Nous verrons encore que ces ions sont alors présents en concentration égales, de valeur :  $[\text{H}_3\text{O}^+]_f = [\text{HO}^-]_f = 1,0 \cdot 10^{-7} \text{ mol.L}^{-1}$  à température ambiante. Le calcul du pH correspondant donne immédiatement :  $\text{pH} = 7,0$ .

Ainsi, une solution est dite neutre lorsqu'elle contient autant d'ions oxonium que d'ions hydroxyde.

**Remarque :** la valeur du pH d'une solution neutre dépend de la valeur de  $K_e$ , qui dépend uniquement de la température. Il en résulte que si la température varie, le pH d'une solution neutre varie également : si  $T$  augmente,  $K_e$  augmente, donc  $[\text{H}_3\text{O}^+]_f$  dans la solution neutre augmente. Ainsi le pH d'une solution neutre diminue-t-il lorsque la température augmente.

Imaginons à présent qu'un acide ait été introduit dans la solution. Sa réaction avec l'eau va fournir un surcroît d'ions oxonium. Ainsi  $[\text{H}_3\text{O}^+]_f$  va augmenter, donc le pH va diminuer et tomber en dessous de sa valeur de neutralité. La solution est alors dite acide, ce qui traduit le fait qu'un acide y a été dissous.

Prenons enfin le cas où c'est une base qui a été introduite en solution. Sa réaction avec l'eau va fournir des ions hydroxyde et c'est cette fois  $[\text{HO}^-]_f$  qui va augmenter. Ce faisant, nous déplaçons l'équilibre d'autoprotolyse de l'eau : comme le produit  $[\text{H}_3\text{O}^+]_f [\text{HO}^-]_f$  est constant (égal au produit ionique de l'eau,  $K_e$ ), alors  $[\text{H}_3\text{O}^+]_f$  doit diminuer, et le pH passe au-dessus de sa valeur de neutralité. La solution est alors dite basique, ce qui traduit le fait qu'une base y a été dissoute.

## — Comment mesure-t-on le pH d'une solution aqueuse ?

La mesure du pH d'une solution aqueuse peut se faire essentiellement de 3 façons. En allant du moins précis au plus précis, nous trouvons :

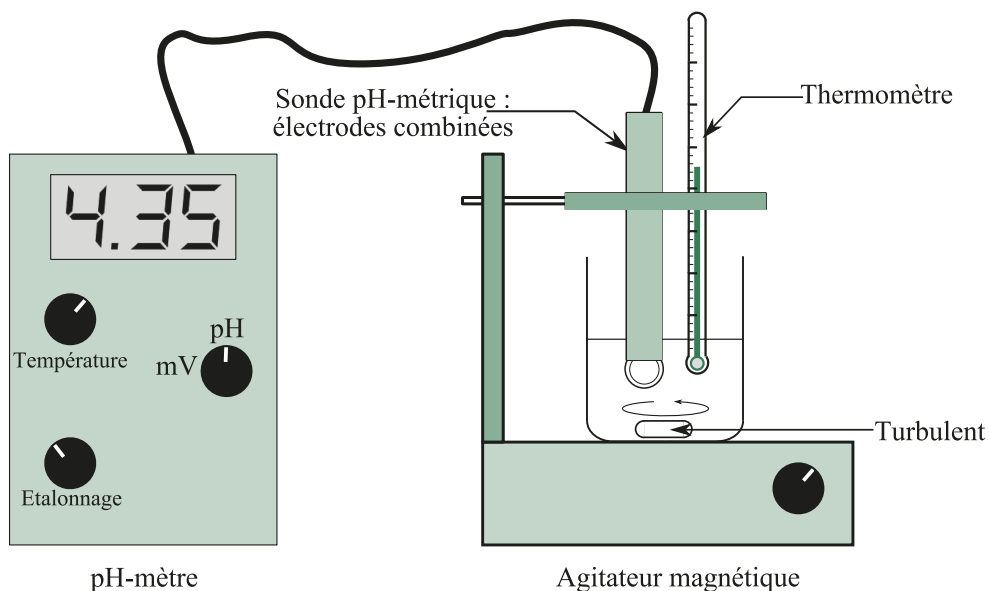
- **Les indicateurs colorés acido-basiques :** il s'agit d'acides ou de bases dont l'espèce conjuguée possède une couleur différente de la leur. Selon le pH du milieu dans lequel elles se trouvent, c'est donc la forme acide (si le pH est faible) ou basique (si le pH est élevé) qui dominera et conférera sa couleur à la solution. Il existe une zone, large de 1 à 3 unités de pH en général, dans laquelle les deux espèces coexistent en proportions comparables. Cette zone est appelée **zone de virage** et est propre à chaque indicateur coloré

acido-basique. La couleur d'une solution contenant un indicateur coloré acido-basique, et dont le pH se trouve dans la zone de virage de cet indicateur, est appelée **teinte sensible**. La teinte sensible d'un indicateur coloré résulte de la superposition en synthèse soustractive des couleurs des formes acide et basique de cet indicateur.

On peut ainsi, en ajoutant quelques gouttes d'un indicateur coloré à une solution aqueuse, savoir si le pH de celle-ci se trouve avant, dans ou bien après la zone de virage. Il s'agit donc d'une méthode très peu précise. Nous verrons cependant par la suite tout l'intérêt que ces indicateurs peuvent malgré tout présenter dans le domaine des titrages acido-basiques.

- **Le papier pH** : il s'agit d'une bandelette de papier très absorbante (type buvard), sur laquelle ont été déposés plusieurs indicateurs colorés acidobasiques, dont les zones de virage couvrent toutes les valeurs courantes de pH. En déposant une goutte de solution à analyser sur ce papier, celui-ci adopte une teinte, à laquelle un code de couleur (propre à chaque marque) fait correspondre une valeur spécifique du pH. On trouve ici une méthode permettant de connaître le pH avec une précision de 1 à 2 unités.
- **Le pH-mètre** : il s'agit d'un instrument constitué d'une électrode de verre et d'une électrode de référence (électrode au calomel saturé en général). Ces deux électrodes sont souvent rassemblées en une seule (on parle alors de pH-mètre à électrodes combinées). On peut alors, par mesure de la différence de potentiel électrique entre ces électrodes, remonter à la valeur du pH. Le pH-mètre est plus compliqué à mettre en œuvre que les deux méthodes précédentes (nécessité d'un étalonnage préalable de l'instrument, entre autres) mais possède une précision nettement meilleure (*cf.* plus bas).

Lorsque l'on mesure le pH d'une solution, celle-ci doit être agitée en permanence afin d'éviter d'éventuelles inhomogénéités de concentration. Le dispositif se représente alors de la manière suivante :



**Remarque** : vous savez qu'un potentiel électrique n'est défini qu'à une constante additive près : nous ne savons mesurer que des **différences de potentiel**. Lorsque l'on parle du potentiel d'une électrode, d'un couple ou d'une demi-pile, ce potentiel est en fait la différence de potentiel entre cette demi-pile et une autre demi-pile particulière, dont le potentiel est fixé par convention comme étant égal à 0 V. Cette demi-pile par rapport à laquelle sont référencés les potentiels des autres est la demi-pile à hydrogène, représentative du couple  $H^+/H_2$ . L'électrode au calomel saturé, qui met en jeu le couple  $Hg_2Cl_2/Hg$ , possède ainsi un potentiel de 0,27 V par rapport à cette électrode à hydrogène.

La différence de potentiel mesurée est ainsi une fonction affine de  $\log[\text{H}_3\text{O}^+]$ , donc également de  $-\log[\text{H}_3\text{O}^+]$ , donc du pH :  $\Delta V_{\text{mes}} = \alpha \times \text{pH} + \beta$ . En fixant pour l'appareil les valeurs de  $\alpha$  et  $\beta$ , celui-ci est alors en mesure de calculer la valeur du pH à partir de celle de  $\Delta V_{\text{mes}}$ . L'**étalonnage** du pH-mètre a précisément pour objectif d'ajuster les valeurs de  $\alpha$  et de  $\beta$ .

Les pH-mètres à vocation pédagogique sacrifient la précision à la robustesse. Bien que souvent affiché à 0,01 unité de pH près, les résultats qu'ils délivrent ne sont généralement fiables (comme le reconnaissent les notices techniques elles-mêmes) qu'à  $\pm 0,05$  unité de pH. La mesure affichée par un pH-mètre n'est donc garantie (et ne doit être donnée, dans le cadre de résultats expérimentaux) qu'à 0,1 unité de pH près.

**Remarque :** il existe également des pH-mètres professionnels, fiables à 0,01 unité de pH. Ceux-ci sont de fait plus fragiles et plus chers que ceux que vous utilisez en classe. Lors de l'utilisation d'un pH-mètre, on doit donc notamment veiller aux points suivants :

- Éviter les chocs, en particulier avec le barreau aimanté ; les électrodes ne doivent pas se trouver trop près du fond.
- Éviter le dessèchement des membranes des électrodes ; les électrodes ne doivent pas être laissées à l'air libre.
- Éviter les abrasions ; les électrodes doivent être essuyées, lorsqu'il y a lieu, avec un papier très doux : le papier Joseph.

## — Comment étalonne-t-on un pH-mètre ?

Le pH-mètre est un instrument qui se dérègle très facilement et doit donc être réétalonné à chaque nouvelle mise sous tension. Dans ce but, on le soumet successivement à deux solutions tampons (solutions de pH fixé et connu) qui vont permettre de fixer deux points de fonctionnement par lesquels devra passer la fonction affine, et ce faisant, de déterminer les valeurs des coefficients  $\alpha$  et  $\beta$  dans la relation  $\Delta V_{\text{mes}} = \alpha \text{pH} + \beta$  vue plus haut.

**Remarque :** rappelons qu'on appelle **solution tampon acido-basique** une solution dont le pH varie peu lors de l'ajout modéré d'acide ou de base, ainsi que lors d'une dilution.

L'étalonnage d'un pH-mètre se déroule en général de la manière suivante :

1. On sélectionne deux solutions tampons : une à peu près neutre (pH = 7 à température ambiante), l'autre différente de 3 à 4 unités de pH, dans la zone de pH où seront effectuées les mesures (vers 3-4 pour une solution acide, donc, et vers 10-11 pour une solution basique).
2. On allume le pH-mètre.
3. On immerge les électrodes dans la solution tampon neutre, et l'on ajuste la valeur affichée à l'aide du bouton de température.
4. On retire les électrodes, on les rince à l'eau distillée et on les essuie avec du papier Joseph.
5. On immerge les électrodes dans la seconde solution tampon, et l'on ajuste la valeur affichée à l'aide du bouton d'étalonnage (parfois noté  $\Delta \text{pH}$ ).

La valeur du pH varie peu tant que la température ne varie pas au-delà de  $\pm 5^\circ\text{C}$ . Une valeur approximative suffit donc (le thermomètre n'est représenté sur le schéma précédent que pour le principe). Nous allons voir qu'en revanche, la question est plus sensible dans le cas du conductimètre.

## — Qu'est-ce que la conductance d'une solution ?

Une solution s'obtient par le mélange d'au moins 2 corps purs :

- l'un liquide et présent en quantité généralement très supérieure à l'autre : **le solvant** (le plus souvent, pour nous : l'eau, ce qui vaudra à la solution le qualificatif de solution *aqueuse*).

- l'autre, initialement solide, liquide ou gazeux, introduit en quantité très inférieure à la quantité de solvant, et qui va se dissoudre dans le solvant : **le soluté**.

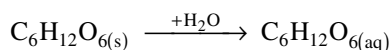
La dissolution consiste en la consommation du soluté, tandis qu'il se mélange au solvant. Elle traduit en fait à l'échelle macroscopique d'autres phénomènes microscopiques que sont :

- la **dissociation** du soluté (si celui-ci est ionique), dont les constituants se séparent les uns des autres ;
- la **solvatation** des produits issus de cette dissociation ;
- la **dispersion** de ces produits solvatés, qui tendent à se répartir de manière homogène dans la solution.

Or selon la nature du soluté considéré, la dissociation peut aboutir à 2 types de produits :

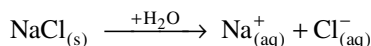
- soit le soluté est constitué de molécules qui conservent leur intégrité au cours du processus et les produits que nous trouverons à l'état solvaté seront simplement ces mêmes molécules, associées à des molécules de solvant (on parle dans ce cas de **soluté moléculaire**).

C'est par exemple le cas du glucose :



- soit le soluté est constitué d'un assemblage électriquement neutre d'ions (**soluté ionique**), ou encore de molécules très polarisées qui vont se dissocier en ions au contact du solvant, et dans ce cas les produits que nous trouverons à l'état solvaté seront les ions issus de cette dissociation.

C'est par exemple le cas du chlorure de sodium :



Dans ce dernier cas, la solution contient donc des ions libres de se déplacer dans la solution, autrement dit des porteurs de charges libres. La solution peut donc conduire le courant électrique et l'on constate en outre que l'intensité  $i$  du courant électrique dont une portion de cette solution va être le siège sous l'effet d'une tension  $u$  imposée à ses bornes, est directement proportionnelle à cette dernière. Autrement dit : la solution se comporte comme un conducteur ohmique, et l'on peut lui associer une résistance  $R = \frac{u}{i}$ .

Vous avez vu qu'il était également possible de travailler avec la conductance  $G = \frac{1}{R} = \frac{i}{u}$  d'un

conducteur ohmique. L'USI de cette grandeur sera donc l' $\Omega^{-1}$ , ou encore l'A.V $^{-1}$ , ou plus simplement le siemens (S).

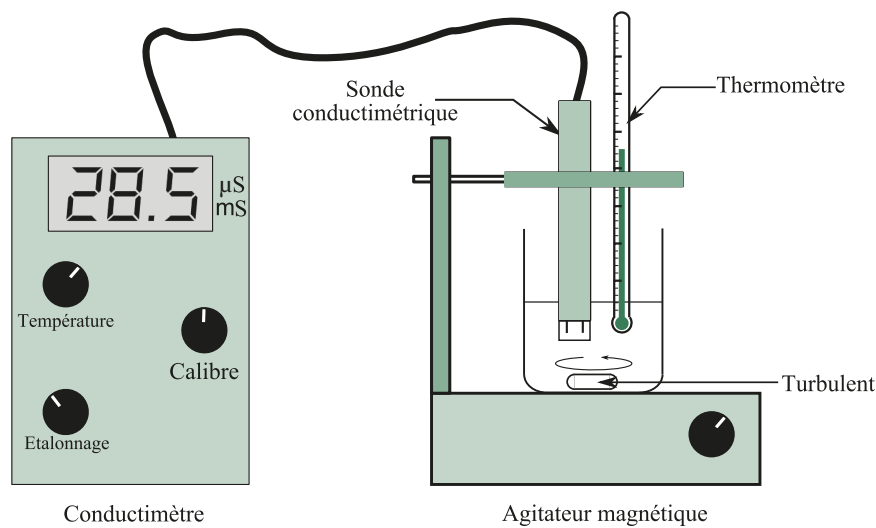
Or cette conductance dépend, entre autres, de la concentration des espèces conductrices dans la solution.

En résumé, une mesure de la conductance d'une portion de solution nous offre un moyen d'accéder à la concentration de cette solution en espèces conductrices.

## — Comment un conductimètre fonctionne-t-il ?

Le conductimètre est un appareil possédant deux électrodes de métal (deux plaques métalliques en regard l'une de l'autre), entre lesquelles il impose une tension électrique alternative  $u$ . Il mesure alors l'intensité  $i$  du courant électrique qui traverse l'espace entre les électrodes. Tout comme le pH-mètre, il est supposé mesurer une grandeur caractéristique de l'ensemble de la solution. Pour garantir l'homogénéité de celle-ci, on a donc soin

de la placer sous agitation préalablement à toute mesure. Le dispositif se présente de la manière suivante :



**Remarque :** comme nous allons le voir, la température a une influence majeure sur les mesures obtenues. Le travail doit donc être mené à température constante et, idéalement, sous contrôle permanent d’un thermomètre. Si cette condition n’est pas toujours réalisée dans les faits, elle doit au moins l’être sur le schéma de principe.

La conduction du courant électrique entre les plaques de la sonde est quantifiée par une grandeur appelée **conductance**  $G$ , et définie comme :

|                   |     |      |
|-------------------|-----|------|
| $G = \frac{i}{u}$ | $G$ | en S |
|                   | $i$ | en A |
|                   | $u$ | en V |

On constate expérimentalement que la conductance dépend de :

- la géométrie de la portion de solution comprise entre les électrodes ;
- la nature des ions présents dans cet espace ;
- la température.

**Remarque :** un autre paramètre, aussi fondamental que difficile à quantifier, est l’état d’oxydation des plaques : plus celles-ci sont oxydées, moins le courant électrique circulera facilement.

On peut en outre quantifier les influences respectives de la surface  $S$  des plaques en regard, et de la distance  $L$  séparant ces plaques :  $G$  est proportionnelle à la première, et inversement proportionnelle à la seconde. On appelle alors **conductivité**  $\sigma$  de la solution le facteur de proportionnalité restant :

|                          |          |                      |
|--------------------------|----------|----------------------|
| $G = \sigma \frac{S}{L}$ | $G$      | en S                 |
|                          | $\sigma$ | en $\text{S.m}^{-1}$ |
|                          | $S$      | en $\text{m}^2$      |
|                          | $L$      | en m                 |

**Remarque :** le facteur  $\frac{S}{L}$  est parfois noté  $k$ , et appelé *constante* ou *paramètre* de cellule. Elle est en réalité directement déterminée lors de l'étalonnage de l'instrument, et intègre donc l'éventuelle altération de la surface due à l'état des plaques.

La conductivité dépend uniquement de la nature des porteurs de charge, de leur concentration dans la solution étudiée, et de la température.

L'influence de la température ne peut être quantifiée simplement, mais est loin d'être négligeable. On aura donc toujours soin, lors des expériences fondées sur la conductimétrie, de travailler à température constante (en évitant notamment de manipuler près d'un radiateur, ou au soleil par exemple, sauf si l'on est sûr que ces derniers dispensent constamment le même rayonnement).

Les ordres de grandeur rencontrés en conductimétrie ( $G$  de l'ordre du mS, voire du  $\mu$ S, et la constante de cellule de l'ordre du cm, voire du mm), font que  $\sigma$  est souvent exprimée en  $\text{mS}\cdot\text{cm}^{-1}$ , ou autres unités certes adaptées mais n'appartenant pas au système international. Dans le cas où les concentrations des espèces solvatées sont faibles (inférieures ou de l'ordre de  $10^{-2} \text{ mol}\cdot\text{L}^{-1}$ , typiquement), on constate que la conductivité d'une solution est liée aux concentrations des ions  $X_i$  qu'elle contient par la **loi de Kohlrausch** :

|                                       |                 |                                                  |
|---------------------------------------|-----------------|--------------------------------------------------|
| $\sigma = \sum_i \lambda_{X_i} [X_i]$ | $\sigma$        | en $\text{S}\cdot\text{m}^{-1}$                  |
|                                       | $\lambda_{X_i}$ | en $\text{S}\cdot\text{m}^2\cdot\text{mol}^{-1}$ |
|                                       | $[X_i]$         | en $\text{mol}\cdot\text{m}^{-3}$                |

Les coefficients  $\lambda_{X_i}$  sont des constantes dépendant uniquement de la nature des ions  $X_i$  et de la température, et appelées **conductivités ioniques molaires** respectives des ions  $X_i$ . Leur USI est le  $\text{S}\cdot\text{m}^2\cdot\text{mol}^{-1}$ , mais l'usage de cette unité suppose des concentrations également exprimées en USI, c'est-à-dire en  $\text{mol}\cdot\text{m}^{-3}$ . Il s'agit d'un écueil courant, dont il convient de se méfier.

## — Comment étalonne-t-on un conductimètre ?

Dans les faits, le conductimètre mesure donc l'intensité du courant électrique à travers la portion de solution comprise entre ses armatures, sous l'effet de la tension électrique qu'il impose aux bornes de cette portion, et la rapporte à ladite tension pour en déduire la conductance. Le conductimètre affiche cependant la conductivité : les paramètres géométriques du problème encombreraient inutilement nos calculs et surtout, comme nous l'avons dit, l'état des plaques joue également un rôle dans la qualité de la conduction du courant électrique.

Il n'est donc pas question de chercher à calculer la constante de proportionnalité entre  $G$  et  $\sigma$ . Celle-ci est déterminée à chaque nouvelle mise sous tension de l'appareil par un étalonnage.

On procède de la manière suivante :

1. On note la valeur de la température du milieu dans lequel vont être effectuées les mesures.
2. On choisit une solution étalon (souvent une solution aqueuse de chlorure de potassium), dont on connaît précisément :
  - la nature ;
  - la concentration en soluté apporté (il est préférable que la valeur de cette concentration soit proche de celles auxquelles on va travailler) ;
  - la valeur de la conductivité à la température de travail.
3. De la valeur attendue pour la conductivité, on déduit le calibre adéquat pour réaliser l'étalonnage, et l'on place le conductimètre sur celui-ci.
4. On introduit la sonde conductimétrique dans la solution étalon.
5. On tourne le bouton d'étalonnage, jusqu'à ce que le conductimètre affiche la bonne valeur.

## — Qu'est-ce que l'absorbance d'une solution ?

Lorsque l'on projette de la lumière à travers une solution colorée, on constate que cette lumière ressort avec une intensité lumineuse moindre. La solution a absorbé une partie de la puissance que véhiculait cette lumière.

On quantifie l'importance de ce phénomène par une grandeur appelée **absorbance** de la solution. Celle-ci se définit, pour une radiation de longueur d'onde dans le vide  $\lambda$  fixée,

par la relation 
$$A(\lambda) = \log \left( \frac{I_i(\lambda)}{I_e(\lambda)} \right).$$

- $I_i(\lambda)$  représente la contribution de la radiation de longueur d'onde dans le vide  $\lambda$ , à l'intensité lumineuse **incidente**, c'est-à-dire à l'entrée de la solution.
- $I_e(\lambda)$  représente la contribution de la radiation de longueur d'onde dans le vide  $\lambda$ , à l'intensité lumineuse **émergente**, c'est-à-dire à la sortie de la solution.

Ainsi, l'absorbance représente la puissance de 10 du rapport entre les valeurs des intensités incidente et émergente :

- une absorbance  $A = 1$  témoigne d'une intensité émergente  $10^1$  fois plus faible que l'intensité incidente, autrement dit une radiation dont 90 % de la puissance a été absorbée par la solution ;
- une absorbance  $A = 2$  témoigne d'une intensité émergente  $10^2$  fois plus faible que l'intensité incidente, autrement dit une radiation dont 99 % de la puissance a été absorbée par la solution.

Pour cette raison, on évite en général de mesurer des absorbances supérieures à 2, puisqu'elles témoignent d'une puissance lumineuse émergente très faible, et donc sujette à de grandes incertitudes.

**Remarque :** nous constatons ainsi que les valeurs limites de l'absorbance sont  $A = 0$  si la solution n'absorbe absolument pas ( $I_e = I_i$ ) et  $A \rightarrow +\infty$  pour une solution parfaitement opaque ( $I_e \rightarrow 0^+$ ).

Si l'on étudie les variations de l'absorbance d'une solution en fonction de la longueur d'onde dans le vide  $\lambda$  de la radiation qui la traverse, on constate que le maximum d'absorption se produit lorsque la radiation utilisée est complémentaire de la couleur sous laquelle est vue la solution. Ainsi, dans le cadre d'un modèle simple de la lumière (résumée à des composantes bleue, verte et rouge), on a :

| Une solution apparaissant... | absorbe le...                                                                | Une solution apparaissant...          | absorbe le...                                       |
|------------------------------|------------------------------------------------------------------------------|---------------------------------------|-----------------------------------------------------|
| Rouge (diffuse le rouge)     | <b>Vert et le bleu</b> (pics d'absorbance à gauche et au milieu du spectre)  | Cyan (diffuse le vert et le bleu)     | <b>Rouge</b> (pic d'absorbance à droite du spectre) |
| Verte (diffuse le vert)      | <b>Rouge et le bleu</b> (pics d'absorbance au milieu et à droite du spectre) | Magenta (diffuse le rouge et le bleu) | <b>Vert</b> (pic d'absorbance au milieu du spectre) |
| Bleue (diffuse le bleu)      | <b>Rouge et le vert</b> (pics d'absorbance à gauche et à droite du spectre)  | Jaune (diffuse le rouge et le vert)   | <b>Bleu</b> (pic d'absorbance à gauche du spectre)  |

Ceci semble logique : la couleur sous laquelle nous apparaît une solution est la superposition des couleurs que cette solution n'absorbe pas. Ainsi plus une radiation est absorbée, moins elle sera diffusée, et donc moins elle contribuera à la couleur sous laquelle est vue la solution.

On constate en outre que, pour une radiation fixée, le phénomène d'absorption est d'autant plus important que :

- l'épaisseur  $l$  de solution traversée est plus grande ;
- la concentration de la solution en espèce colorée est plus grande.

Une étude quantitative permet de montrer que, pour de faibles concentrations (inférieures ou de l'ordre de  $10^{-2}$  mol.L<sup>-1</sup>), l'absorbance est liée à la concentration d'une espèce colorée X dans la solution et à l'épaisseur  $l$  de solution traversée par la relation suivante (dite **loi de Beer-Lambert**) :

|                                |                        |                                      |
|--------------------------------|------------------------|--------------------------------------|
| $A = \varepsilon(\lambda)l[X]$ | $A$                    | sans unité                           |
|                                | $\varepsilon(\lambda)$ | en m <sup>2</sup> .mol <sup>-1</sup> |
|                                | $l$                    | en m                                 |
|                                | $[X]$                  | en mol.m <sup>-3</sup>               |

$\varepsilon(\lambda)$  est appelé **coefficient d'extinction molaire** de la substance étudiée. Ce terme centralise toute la dépendance de l'absorbance d'une solution contenant cette substance vis-à-vis de la longueur d'onde dans le vide de la radiation traversant cette solution. Une fois la radiation de travail choisie, la valeur de ce coefficient est donc fixée et il n'apparaît plus que comme un coefficient de proportionnalité.

Cette dépendance de l'absorbance vis-à-vis de la concentration de la solution en espèce absorbante, nous offre ici encore un moyen d'accéder à la valeur de cette concentration en s'appuyant sur la mesure d'une grandeur physique qui lui est liée.

**Remarque :** ici encore, si nous souhaitons tout exprimer en USI pour répondre notamment à une longueur exprimée en mètres, la concentration doit être exprimée en mol.m<sup>-3</sup>. On peut cependant trouver des coefficients d'extinction molaire exprimés en L.mol<sup>-1</sup>.cm<sup>-1</sup>, ce qui suppose donc des valeurs de  $l$  et  $[X]$  exprimées respectivement en cm et en mol.L<sup>-1</sup>.

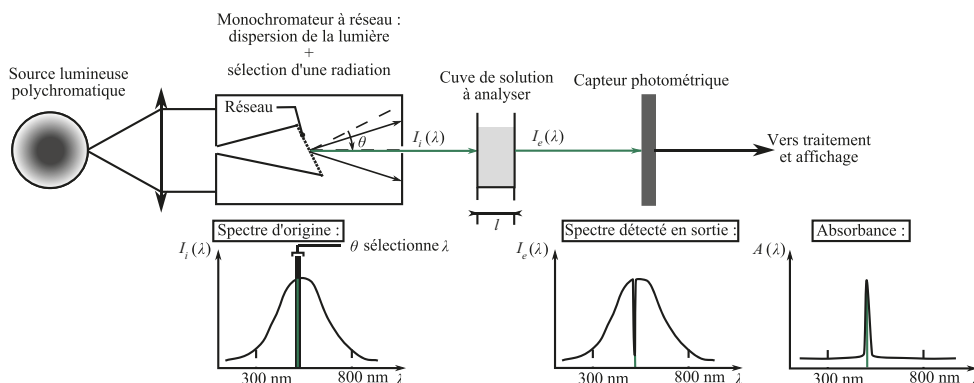
Le spectrophotomètre a déjà été évoqué dans le chapitre 1 (section 1.5, sur la spectroscopie). Il s'agit, donc, d'un appareil destiné à mesurer l'absorbance d'une solution à une certaine longueur d'onde. Pour ce faire, on dispose d'une source lumineuse polychromatique dont le spectre contient notamment toutes les radiations du domaine visible (longueurs d'onde dans le vide comprises entre  $\lambda_{\text{violet}} = 0,384 \mu\text{m}$  et  $\lambda_{\text{rouge}} = 0,768 \mu\text{m}$ ).

La lumière issue de cette source est dispersée par un réseau et chaque radiation est envoyée dans une direction différente.

En faisant pivoter le réseau, on fait également pivoter l'éventail des radiations. Une fente placée dans une direction précise permet ainsi à l'utilisateur de sélectionner l'une de ces radiations, de longueur d'onde dans le vide  $\lambda$  fixée.

L'intensité lumineuse  $I_i$  de cette radiation est mesurée par l'appareil au moment où l'on fait le blanc. Lorsque la radiation traverse une cuve de solution contenant des espèces solvatées, elle en ressort avec une intensité lumineuse moindre,  $I_e$ , que l'instrument mesure également. L'instrument délivre alors une absorbance calculée sur la base de ces mesures.

**Remarque :** la sensibilité des photorécepteurs utilisés pour mesurer  $I_e$  a ses limites. Aussi, si la valeur de  $A$  devient trop grande, le spectrophotomètre sature : il n'affiche plus qu'une valeur (souvent 2, correspondant à une intensité émergente égale à 1 % de l'intensité incidente), à un seul chiffre significatif. Le choix de la longueur d'onde de travail doit toujours répondre au compromis suivant : permettre des valeurs d'absorbance aussi élevées que possible (pour un maximum de précision sur les résultats), tout en évitant de faire saturer l'appareil.



L'étude spectrophotométrique d'une solution passe en général par le tracé de deux courbes :

- Le tracé du **spectre d'absorption** de la solution pour une concentration en soluté apporté donnée : on fait varier la longueur d'onde dans le vide  $\lambda$  de la radiation et l'on mesure à chaque fois l'absorption. On obtient ainsi la courbe représentative  $A(\lambda)$  sur la base de laquelle on sélectionnera la valeur  $\lambda_0$  de la longueur d'onde dans le vide à laquelle sera menée l'étude de la solution.

Cette valeur doit être telle que l'absorbance soit maximale, afin de conserver un maximum de précision tout en évitant de faire saturer le spectrophotomètre. Il est donc bon, préalablement à toute expérience reposant sur de la spectrophotométrie, de calculer l'ordre de grandeur des concentrations en espèces colorées auxquelles on peut s'attendre.

- Une fois sélectionnée la longueur d'onde de travail, on va mesurer l'absorbance de cette solution à diverses concentrations apportées en soluté. On obtient ainsi un nuage de points expérimentaux qui, loi de Beer-Lambert oblige, peut être modélisé en très bonne approximation par une droite passant par l'origine, d'équation :  $A = k[X]$ , où  $k = \varepsilon(\lambda_0)l$  s'identifie à la pente de cette droite (en  $\text{L} \cdot \text{mol}^{-1}$ ). La détermination de la pente de cette droite permet donc par la suite de déterminer une concentration à partir d'une mesure d'absorbance.

Dans le cas où l'on étudie la cinétique d'une réaction, une troisième courbe peut être tracée, mesurant l'évolution de l'absorbance de solution au cours du temps. On obtient ainsi la courbe représentative  $A(t)$ . Les valeurs de  $A$  permettent alors de remonter à la valeur de  $[X]$ , qui permet elle-même de déterminer l'avancement :

$$A = k[X] \quad [X] = [X]_i \mp \alpha_X x_V \Leftrightarrow x_V = \pm \frac{1}{\alpha_X} \left( [X]_i - \frac{A}{k} \right)$$

où  $x_V = \frac{x}{V}$  désigne l'avancement de la réaction par unité de volume, appelé **avancement volumique**, et  $\alpha_X$  est le nombre stœchiométrique portant sur l'espèce X. Le signe  $\pm$  dans l'expression finale dépend de la nature (réactif ou produit) de l'espèce X : + s'il s'agit d'un réactif, - s'il s'agit d'un produit.

Dans le cas d'une étude cinétique, la détermination de la vitesse de réaction à une date  $t_0$  s'obtient par mesure de la pente de la tangente à la courbe représentative de  $x(t)$  (ou  $x_V(t)$  si l'on cherche une vitesse volumique de réaction), à la date  $t_0$ . Pour travailler à partir de la courbe représentative de  $A(t)$ , deux méthodes sont possibles :

- Recalculer toutes les valeurs de  $x_V(t)$  à partir de celles de  $A(t)$  à l'aide de la relation vue ci-dessus. On doit alors tracer la courbe des variations de  $x_V(t)$ , et la pente de la tangente à cette courbe donnera la vitesse volumique de la réaction. Cette méthode est coûteuse en

temps puisqu'elle nécessite le calcul de toutes les valeurs de  $x$  et le tracé d'une nouvelle courbe, donc au minimum un codage dans un tableur.

- Dériver la relation vue ci-dessus par rapport au temps :

$$x_V = \pm \frac{1}{\alpha_X} \left( [X]_i - \frac{A}{k} \right) \Rightarrow v_V = \frac{dx_V}{dt} = \mp \frac{1}{k\alpha_X} \frac{dA}{dt}$$

Il suffit alors de tracer directement la tangente à la courbe  $A(t)$ , d'en mesurer la pente (dont la valeur est  $\frac{dA}{dt}$  à la date considérée), et enfin de multiplier celle-ci par  $\mp \frac{1}{k\alpha_X}$ .

**Remarque :** si l'espèce colorée est un réactif, elle est consommée au cours du temps et l'absorbance décroît. La pente de la tangente à la courbe représentative de  $A(t)$  est donc négative mais, multipliée par le coefficient  $-\frac{1}{k\alpha_X}$ , on aboutit finalement bien à une vitesse positive. S'il s'agit d'un produit, alors la pente est positive et, multipliée par  $+\frac{1}{k\alpha_X}$ , on trouve de nouveau une vitesse positive.

Si une solution contient plusieurs espèces absorbant la lumière, alors l'absorbance totale de la solution est la somme des absorbances que présenterait cette solution, si elle ne contenait à chaque fois qu'une seule de ces espèces en solution :

$$A_{\text{tot}} = \sum_k A_k$$

$A_{\text{tot}}, A_k$  sans unité

Il reste cependant un problème : l'espèce colorée n'est pas la seule à posséder des propriétés d'absorption de la lumière. Dans les faits, toute matière, même transparente, absorbe de la lumière, y compris et même surtout le plastique de la cuve contenant la solution, ainsi que le solvant lui-même. Il est donc nécessaire de réaliser une mesure permettant au spectrophotomètre de décompter la contribution de ces facteurs parasites, à la valeur totale de l'absorbance. Ceci revient en fait à réaliser un étalonnage. On dit qu'on « fait le blanc ».

## — Comment étalonne-t-on un spectrophotomètre ?

Étalonner un spectrophotomètre se dit également **faire le blanc**. La manipulation est très simple :

1. On sélectionne la longueur d'onde de travail (attention : certains modèles de spectrophotomètres utilisent différents filtres selon le domaine de longueur d'onde dans lequel on travaille ; le choix du filtre doit être réalisé au même moment que celui de la longueur d'onde de travail).
2. On introduit dans le spectrophotomètre une cuve contenant le solvant de travail (attention, ces cuves présentent en général deux faces striées ou dépolies pour poser les doigts, et deux autres lisses, à travers lesquelles passe la lumière ; une simple empreinte digitale sur celles-ci suffit à fausser les mesures).
3. On appuie sur le bouton portant la mention **blanc** ou **étalonnage** : le spectrophotomètre est à présent réglé.

L'utilisation du spectrophotomètre est alors des plus simples : on introduit à chaque fois une cuve de la solution à analyser, et on lit la valeur de l'absorbance sur l'écran de l'appareil.

**Remarque :** il est recommandé de refaire le blanc périodiquement, et obligatoirement si l'on change la longueur d'onde de travail.

## Halte aux idées reçues

Expliquer, pour chacune des affirmations suivantes, pourquoi elle est fausse :

1. Le volume molaire d'une espèce chimique ne peut être défini que si cette espèce est gazeuse.
2. Masse volumique et concentration massique sont deux appellations d'une même grandeur.
3. L'étude spectrophotométrique d'une solution doit toujours être menée à la longueur d'onde pour laquelle cette solution présente une absorbance maximale.

4. Si on multiplie la concentration apportée en un soluté ionique dans une solution par un facteur  $n$ , on multiplie également sa conductivité par  $n$ .

5. Une diminution de  $n$  unités du pH d'une solution témoigne d'une multiplication par  $n$  de la concentration en ions oxonium.

## Du Tac au Tac

(Exercices sans calculatrice)

1. On introduit, dans une fiole jaugée de volume  $V_0 = 200,0$  mL, une masse  $m_1 = 0,585$  g de chlorure de sodium NaCl, un volume  $V_2 = 50,0$  mL de solution aqueuse d'acide nitrique  $\text{HNO}_3$  à la concentration apportée en soluté  $c_2 = 1,00 \cdot 10^{-1}$  mol.L<sup>-1</sup>, et un volume  $V_3 = 480$  mL de chlorure d'hydrogène gazeux HCl. On complète ensuite avec de l'eau distillée jusqu'au trait de jauge.

Déterminer les concentrations effectives de toutes les espèces présentes dans ce mélange.

On admettra que toutes les espèces se dissocient totalement et qu'elles ne réagissent pas entre elles.

On donne les masses molaires suivantes :  $M_{\text{Na}} = 23,0$  g.mol<sup>-1</sup>,  $M_{\text{Cl}} = 35,5$  g.mol<sup>-1</sup>, et le volume molaire des gaz dans les conditions de l'expérience :  $V_m = 24,0$  L.mol<sup>-1</sup>.

2. On considère une solution aqueuse d'acide sulfurique  $\text{H}_2\text{SO}_4$  présentant les caractéristiques suivantes :

- masse volumique :  $\rho_{\text{solution}} = 1,3$  g.mL<sup>-1</sup>;
- masse molaire du soluté :  $M_{\text{H}_2\text{SO}_4} = 98$  g.mol<sup>-1</sup>;
- fraction massique du soluté dans la solution :  $x_{\text{m,H}_2\text{SO}_4} = 49$  %.

Déterminer la concentration molaire de cette solution en soluté apporté.

3. On considère un échantillon de gaz noble porté à une température de 25 °C.

Déterminer la valeur en °C à laquelle doit être portée la température pour doubler sa masse volumique à pression constante (on approximera le zéro absolu à -273 °C).

4. On considère une solution aqueuse d'acide chlorhydrique ( $\text{H}_3\text{O}_{(\text{aq})}^+; \text{Cl}_{(\text{aq})}^-$ ) de concentration en soluté apporté  $C$ , dont on mesure le pH et la conductivité. On obtient les valeurs  $\text{pH} = 2,00$  et  $\sigma = 4,26$  ms.cm<sup>-1</sup>.

Sachant que la conductivité ionique molaire des ions oxonium a pour valeur :

$$\lambda_{\text{H}_3\text{O}^+} = 35,0 \text{ mS.m}^2.\text{mol}^{-1}$$

déterminer celle des ions chlorure.

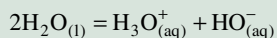
5. Un indicateur coloré acidobasique est un couple acide/base dont les membres présentent des couleurs différentes. Si l'on introduit une petite quantité de l'un ou de l'autre dans une solution, les proportions respectives des deux formes (et donc la couleur de la solution) seront conditionnées par le pH de la solution en question.

On considère deux échantillons contenant un même indicateur coloré acidobasique à la concentration en soluté apporté  $C$ , l'une dont le pH a été mené à la valeur  $\text{pH} = 1,0$ , l'autre à la valeur  $\text{pH} = 13,0$ , et l'on trace les spectres d'absorbance, respectivement  $A_1(\lambda)$  et  $A_{13}(\lambda)$ , de ces deux solutions.

Montrer que si l'on se place à une longueur d'onde où les absorbances des deux solutions sont égales entre elles, alors pour toute solution dans laquelle cet indicateur coloré a été introduit à la concentration  $C$ , l'absorbance sera la même quel que soit le pH.

## Vers la prépa

L'eau se dissocie spontanément en ions oxonium  $\text{H}_3\text{O}^+$  et hydroxyde  $\text{HO}^-$  selon la réaction dite d'*autoprotolyse de l'eau*, d'équation :



Ainsi l'eau dite « pure » contient-elle ces deux ions en concentrations égales.

On peut alors montrer qu'à l'équilibre chimique, le produit de ces concentrations conserve, à température donnée, la même valeur en toute circonstance ; cette valeur est appelée *produit ionique* de l'eau :

$$K_e = [\text{H}_3\text{O}^+]_f \times [\text{HO}^-]_f$$

La valeur de ce produit ionique dépend exclusivement de la température absolue  $T$ , selon la *loi de Van't Hoff* :

$$\frac{d \ln K_e}{dT} = \frac{\alpha}{T^2}$$

avec  $\alpha = 6,9 \cdot 10^3 \text{ K}$ .

Montrer que pour de faibles variations de températures, le pH de l'eau varie selon une fonction affine de la température, dont on précisera les caractéristiques et les limites de validité.

On rappelle qu'à  $T_0 = 298\text{K}$ , la valeur du pH est de  $\text{pH}_0 = 7,0$ , et l'on pourra utiliser l'approximation selon laquelle, pour  $\varepsilon \ll 1$  :

$$\frac{1}{1 + \varepsilon} \simeq 1 - \varepsilon$$

# Corrigés

## Halte aux idées reçues

1. Le volume molaire d'un gaz est défini comme le rapport du volume occupé par un gaz, à la quantité de matière d'entités gazeuses (qu'il s'agisse d'atomes ou de molécules) contenue dans ce volume. On peut parfaitement appliquer cette même définition pour un liquide ou un solide.

Pour l'eau liquide, par exemple, on a alors :

$$V_{m,H_2O} = \frac{V_{H_2O}}{n_{H_2O}} = \frac{V_{H_2O}}{\frac{m_{H_2O}}{M_{H_2O}}} = \frac{M_{H_2O}}{\rho_{H_2O}} = 18,0 \text{ mL.mol}^{-1}$$

Ceci a quelque chose de logique : l'eau a une masse de 18,0 g par mole de molécules d'eau considérée. Or, on sait que l'eau liquide occupe un volume de 1 mL par gramme d'eau considérée. De ce fait, une mole d'eau pèse 18,0 g, et occupe donc 18,0 mL, d'où cette valeur de  $18,0 \text{ mL.mol}^{-1}$  (notons au passage que ce volume molaire est très inférieur à celui des gaz, puisqu'en phase liquide la matière se trouve beaucoup plus condensée).

Pourquoi dans ce cas réserver l'usage du volume molaire aux gaz ? Tout simplement parce que le volume molaire tel qu'il est calculé ci-dessus repose sur deux grandeurs propres à l'espèce considérée :

- Sa masse molaire, qui se calcule très simplement comme somme des masses molaires atomiques des atomes constituant une molécule de l'espèce considérée. Il suffit donc, pour la connaître, d'avoir à disposition les masses molaires de la grosse centaine d'éléments dont est constitué l'univers, et qui figurent dans le tableau de classification périodique de Mendeleïev.
- Sa masse volumique qui, elle, ne peut se calculer simplement à partir de grandeurs relatives aux atomes la constituant. On a donc alors besoin de connaître individuellement la masse volumique de chaque espèce, sachant que s'il n'existe qu'une centaine d'éléments chimiques, leurs combinaisons donnent en revanche des millions d'espèces chimiques différentes.

Ce problème lié à la masse volumique rend donc le volume molaire sans grand intérêt dans le cas général. Mais dans le cas particulier des gaz, on constate le petit miracle que résume la loi d'Avogadro-Ampère : le volume molaire ne dépend pas de la nature du gaz considéré : qu'il s'agisse d'hélium, de dihydrogène, de néon, de dioxygène, d'argon, de dioxyde de carbone, de vapeur d'eau, de diazote... Tous ont le même

volume molaire, qui dépend uniquement de la température et de la pression.

Dans le cas d'un gaz, on n'a donc plus besoin d'écrire  $V_{m,X} = \frac{V_X}{n_X}$ , mais seulement  $V_m = \frac{V_X}{n_X}$ , où  $V_m$  a, pour une pression et une température données, la même valeur quelle que soit la nature chimique de X.

Cette simplification considérable fait que le volume molaire, est (mais dans le cas des gaz uniquement) une façon commode de déterminer une quantité de matière.

On peut ajouter que le volume molaire d'un gaz a pour expression, si celui-ci satisfait à la loi du gaz parfait :  $V_m = \frac{V_X}{n_X} = \frac{RT}{P}$ . Ceci permet de souligner deux points importants :

- Cette expression permet de calculer le volume molaire du gaz (supposé parfait), à partir de sa pression et de sa température. Attention toutefois : si  $R$  est fournie en USI, alors  $P$  doit être exprimée en pascals et  $T$  en kelvins, et le volume molaire sera obtenu en  $\text{m}^3.\text{mol}^{-1}$ .
- Si un gaz répond mal au modèle du gaz parfait (situation assez rare), son volume molaire diffère également de l'expression fournie ci-dessus, et l'on commence à perdre le bénéfice du volume molaire identique pour tous les gaz.

2. Si l'on utilise les formules sans réfléchir, par exemple dans le cas présent  $\frac{m}{V}$ , on peut effectivement être tenté d'assimiler masse volumique et concentration massique. Les deux sont définies par le rapport d'une masse à un volume, donc ont la même unité.

Si l'on regarde d'un peu plus près, on constate que concentration massique et masse volumique sont deux grandeurs quantifiant deux relations de proportionnalités différentes :

- La masse d'un échantillon de corps pur est proportionnelle au volume qu'occupe cet échantillon : si un échantillon de ce corps occupe deux fois plus de place qu'un autre, il est deux fois plus lourd. Idem pour trois, quatre... fois. La masse volumique de ce corps est le coefficient de proportionnalité entre le volume et la masse d'un échantillon de ce corps. On peut sous-titrer son unité comme kg (de corps pur), par mètre cube (de ce même corps pur) :

$$\rho_X = \frac{m_X}{V_X}$$

- La masse de soluté apporté dans un échantillon de solution, est proportionnelle au volume de cet échantillon : si l'on prélève un échantillon de cette solution, dont le volume est deux fois plus important qu'un autre, alors cet échantillon contient une masse de soluté apporté deux fois plus importante. La concentration massique de cette solution en ce soluté est le coefficient de proportionnalité entre le volume de solution prélevé et la masse de soluté apporté qui vient avec. On peut sous-titrer son unité comme kg (de soluté X apporté) par mètre cube (de solution) :

$$c_X = \frac{m_X}{V_{\text{solution}}}$$

3. En règle générale, on cherche à obtenir des résultats avec la plus grande précision possible. Or la précision est d'autant plus limitée que les valeurs mesurées sont plus petites. Il est donc *a priori* intéressant de mener une étude spectrophotométrique à la longueur d'onde à laquelle l'absorbance manifeste les valeurs les plus élevées, donc à la longueur d'onde du maximum d'absorption de l'espèce colorée sur laquelle repose le suivi.

Cependant, il ne faut pas perdre de vue que l'absorbance quantifie la perte de lumière occasionnée par la traversée d'une solution. Si elle adopte de grandes valeurs, cela signifie que le faisceau de lumière récupéré est très peu visible. Or le capteur chargé d'évaluer la lumière résiduelle ne dispose pas non plus d'une précision infinie et, si cette lumière est trop faible, sa réponse tendra à se confondre avec le bruit (lumineux et/ou électronique) inhérent à tout appareil de mesure. Dans la pratique, ceci se manifeste par la saturation du spectrophotomètre à une certaine valeur (souvent 2,000).

Aussi doit-on s'assurer, avant de sélectionner une longueur d'onde de travail, que la concentration de l'espèce suivie dans la solution en question est au moins égale à la valeur maximale de la concentration qui sera rencontrée au cours de l'étude. Dans le cas contraire, on risque de voir le spectrophotomètre saturer en cours de route.

Si l'on constate que le spectrophotomètre sature à cette concentration maximale pour la longueur d'onde utilisée, alors on doit opter pour une autre longueur d'onde de travail qui offrira le compromis recherché.

Notons au passage que l'on pourrait imaginer de travailler simplement à des concentrations moins élevées. Le problème étant que la spectrophotométrie est souvent utilisée pour des suivis cinétiques, et que les cinétiques en question dépendent des concentrations des réactifs en solution. Ainsi, si l'espèce colorée est un réactif, diminuer sa concentration initiale dans le milieu réactionnel signifierait aussi ralentir la réaction et possiblement sortir de tel ou tel contexte expérimental.

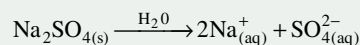
4. La relation décrite par l'énoncé est une relation de proportionnalité, entre la conductivité d'une solution et sa concentration apportée en soluté.

Plusieurs nuances doivent être apportées à cette affirmation. En premier lieu, on peut tout-à-fait étudier la conductivité d'une solution dans laquelle ont été dissous plusieurs solutés. Dans ce cas la conductivité n'est pas une fonction linéaire de

la concentration apportée en soluté, mais une fonction affine par rapport à chacune des concentrations apportées en soluté.

Si à présent nous admettons qu'un seul soluté ionique a été dissous, nous savons que la conductivité peut s'exprimer comme la somme des produits des concentrations effectives de chacun des ions résultant de la dissociation de ce soluté, par un coefficient qui lui est propre : la conductivité ionique molaire de cet ion. Or la concentration effective de chacun des ions libérés en solution par un soluté ionique est proportionnelle à la concentration apportée en ce soluté.

Si par exemple on considère la dissolution du sulfate de sodium dans l'eau à raison d'une concentration apportée en soluté  $c_0$  :



Les ions obtenus ont pour concentrations effectives respectives  $[\text{Na}^+] = 2c_0$  et  $[\text{SO}_4^{2-}] = c_0$ . L'expression de la conductivité de la solution serait alors :

$$\sigma = \lambda_{\text{Na}^+} [\text{Na}^+] + 2 \times \lambda_{\text{SO}_4^{2-}} [\text{SO}_4^{2-}] = 2 \left( \lambda_{\text{Na}^+} - \lambda_{\text{SO}_4^{2-}} \right) c_0$$

(Le facteur 2 sur les ions sulfate provient d'une convention selon laquelle les conductivités ioniques molaires des ions polychargés sont données en considérant qu'elles seront multipliées par leur charge dans l'expression de la conductivité.) Donc on trouve effectivement une relation de proportionnalité entre conductivité  $\sigma$  et concentration  $c_0$  apportée en soluté. Mais il n'importe pas seulement de retenir la partie simple d'une propriété : encore faut-il en connaître les limites. Or la propriété énoncée ci-dessus n'est valable que pour des concentrations au plus de l'ordre de  $10^{-2} \text{ mol.L}^{-1}$ .

En résumé, l'affirmation ci-dessus est vraie mais seulement moyennant deux conditions : un seul soluté dissous dans la solution, et à raison d'au plus  $10^{-2} \text{ mol}$  de soluté par litre de solution.

5. Contrairement à l'affirmation précédente, qui devait surtout être clarifiée et cadrée, celle-ci est catégoriquement fausse. On sait en effet que le pH d'une solution aqueuse est défini comme l'opposé du logarithme décimal de sa concentration effective en ions oxonium, exprimée en  $\text{mol.L}^{-1}$  :

$$\text{pH} = -\log[\text{H}_3\text{O}^+]_{\text{mol/L}} \Leftrightarrow [\text{H}_3\text{O}^+]_{\text{mol/L}} = 10^{-\text{pH}}$$

On constate donc qu'une diminution de  $n$  unités de pH entraîne une baisse de la concentration en ions oxonium d'un facteur multiplicatif  $10^{-n}$ , donc une décroissance considérablement plus importante qu'une simple division par  $n$ . Il ne s'agit pas de diviser  $[\text{H}_3\text{O}^+]$  par 2, 3, 4..., mais par 100, 1 000, 10 000...

## Du Tac au Tac

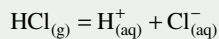
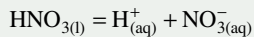
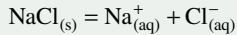
(Exercices sans calculatrice)

1. Commençons par déterminer les quantités de matière apportées en solutés :

- Le chlorure de sodium est apporté sous forme solide, et l'on nous donne la masse  $m_1$  de l'échantillon introduit. Nous en déduisons la quantité de matière  $n_{1,\text{app}}$  apportée en chlorure de sodium NaCl introduit :  $n_{1,\text{app}} = \frac{m_1}{M_{\text{NaCl}}} = 10,0 \text{ mmol}$ .

- L'acide nitrique est apporté déjà mélangé avec de l'eau. Ceci ne nous empêche pas de calculer la quantité de matière  $n_{2,app}$  apportée en acide nitrique dans le volume  $V_2 : n_{2,app} = c_2 V_2 = 5,00 \text{ mmol}$ .
- Le chlorure d'hydrogène est apporté sous forme gazeuse, et l'on nous donne le volume de gaz  $V_3$  dissout dans le mélange. S'agissant d'un gaz, nous pouvons utiliser le volume molaire des gaz et nous obtenons ainsi la quantité de matière  $n_{3,app}$  apportée en chlorure d'hydrogène dans le volume  $V_3 : n_{3,app} = \frac{V_3}{V_m} = 20,0 \text{ mmol}$ .

Une fois le mélange réalisé, chacun de ces solutés se dissocie. Pour déterminer les concentrations effectives, on doit considérer les espèces issues de la dissociation de ces solutés, soit ici :



Nous devons alors déterminer la quantité de matière de chaque espèce dissoute. Tous les nombres stœchiométriques étant égaux entre eux, il vient immédiatement que :

- Pour les espèces issues du chlorure de sodium, on a :  $n_{\text{Na}^+,1} = n_{\text{Cl}^-,1} = n_{1,app} = 10,0 \text{ mmol}$ .
- Pour les espèces issues de l'acide nitrique, on a :  $n_{\text{H}^+,2} = n_{\text{NO}_3^-,2} = n_{2,app} = 5,00 \text{ mmol}$ .
- Pour les espèces issues du chlorure d'hydrogène, on a :  $n_{\text{H}^+,3} = n_{\text{Cl}^-,3} = n_{3,app} = 20,0 \text{ mmol}$ .

Les espèces solvatées sont donc au final, en prenant en compte pour chacune ses différentes sources :

- Les ions sodium :  $n_{\text{Na}^+} = n_{\text{Na}^+,1} = 10,0 \text{ mmol}$ .
- Les ions chlorure :  $n_{\text{Cl}^-} = n_{\text{Cl}^-,1} + n_{\text{Cl}^-,3} = 30,0 \text{ mmol}$ .
- Les ions hydrogène :  $n_{\text{H}^+} = n_{\text{H}^+,2} + n_{\text{H}^+,3} = 25,0 \text{ mmol}$ .
- Les ions nitrate :  $n_{\text{NO}_3^-} = n_{\text{NO}_3^-,2} = 5,00 \text{ mmol}$ .

On trouve enfin les concentrations effectives, en rapportant ces quantités de matière au volume au volume  $V_0 = 200,0 \text{ mL}$  qui leur est offert :

- $[\text{Na}^+] = \frac{n_{\text{Na}^+}}{V_0} = 5,00 \cdot 10^{-2} \text{ mol.L}^{-1}$ .
- $[\text{Cl}^-] = \frac{n_{\text{Cl}^-}}{V_0} = 1,50 \cdot 10^{-1} \text{ mol.L}^{-1}$ .
- $[\text{H}^+] = \frac{n_{\text{H}^+}}{V_0} = 1,25 \cdot 10^{-1} \text{ mol.L}^{-1}$ .
- $[\text{NO}_3^-] = \frac{n_{\text{NO}_3^-}}{V_0} = 2,50 \cdot 10^{-2} \text{ mol.L}^{-1}$ .

2. Nous pouvons partir de la définition de la concentration en soluté apporté, en distinguant bien :

- les grandeurs relatives au soluté ( $\text{H}_2\text{SO}_4$ ) ;
- celles relatives à la solution proprement dite.

Nous avons ainsi :

$$C_{\text{H}_2\text{SO}_4} = \frac{n_{\text{H}_2\text{SO}_4}}{V_{\text{solution}}}$$

Notre objectif est d'exprimer cette concentration à partir des différentes données fournies par l'énoncé. Or ces dernières portent toutes sur des grandeurs massiques (masse volumique, masse molaire, fraction massique). Il semble donc approprié de déplacer le problème, de la quantité de matière de soluté à la masse correspondante :

$$n_{\text{H}_2\text{SO}_4} = \frac{m_{\text{H}_2\text{SO}_4}}{M_{\text{H}_2\text{SO}_4}}$$

En injectant cette expression dans celle de la concentration vue plus haut, il vient :

$$C_{\text{H}_2\text{SO}_4} = \frac{1}{M_{\text{H}_2\text{SO}_4}} \times \frac{m_{\text{H}_2\text{SO}_4}}{V_{\text{solution}}}$$

Les autres grandeurs se rapportant à la solution, et vu la présence du volume de solution au dénominateur, il peut être bienvenu de basculer sur des grandeurs relatives à la seule solution. On peut ainsi exprimer la masse de soluté dans l'échantillon à partir de la masse de solution correspondante :

$$x_{m,\text{H}_2\text{SO}_4} = \frac{m_{\text{H}_2\text{SO}_4}}{m_{\text{solution}}} \Leftrightarrow m_{\text{H}_2\text{SO}_4} = x_{m,\text{H}_2\text{SO}_4} \times m_{\text{solution}}$$

En injectant ce résultat dans l'expression précédente, nous obtenons :

$$C_{\text{H}_2\text{SO}_4} = \frac{x_{m,\text{H}_2\text{SO}_4}}{M_{\text{H}_2\text{SO}_4}} \times \frac{m_{\text{solution}}}{V_{\text{solution}}}$$

Nous reconnaissons alors l'expression de la masse volumique de la solution dans la seconde fraction figurant dans le membre de droite, et l'expression devient finalement :

$$C_{\text{H}_2\text{SO}_4} = \frac{x_{m,\text{H}_2\text{SO}_4}}{M_{\text{H}_2\text{SO}_4}} \times \rho_{\text{solution}}$$

L'application numérique nous donne alors :

$$C_{\text{H}_2\text{SO}_4} = \frac{0,49}{98 \text{ g.mol}^{-1}} \times 1,3 \cdot 10^3 \text{ g.L}^{-1} = 6,5 \text{ mol.L}^{-1}$$

3. Commençons par exprimer la masse volumique du gaz :

$$\rho_{\text{gaz}} = \frac{m_{\text{gaz}}}{V_{\text{gaz}}}$$

La masse de l'échantillon étant fixée, nous ne pouvons modifier la masse volumique de celui-ci qu'en modifiant le volume qu'il occupe.

Nous savons que le volume occupé par un gaz dépend des autres paramètres caractéristiques (vous apprendrez à les appeler *grandeurs d'état*) de ce gaz via la loi du gaz parfait :

$$P \times V_{\text{gaz}} = n_{\text{gaz}} \times R \times T \Leftrightarrow V_{\text{gaz}} = \frac{n_{\text{gaz}} \times R}{P} \times T$$

Nous pouvons donc exprimer la dépendance de la masse volumique du gaz vis-à-vis de la température comme :

$$\rho_{\text{gaz}} = \frac{m_{\text{gaz}} \times P}{n_{\text{gaz}} \times R} \times \frac{1}{T}$$

La masse volumique varie donc en raison inverse de la température absolue du gaz. Pour doubler la première, il est donc nécessaire de diviser la seconde par 2. Nous savons que 25 °C correspondent à  $T = 298 \text{ K}$  ; cette température doit donc être abaissée à  $\frac{298 \text{ K}}{2} = 149 \text{ K}$ , soit encore  $-124 \text{ °C}$ .

4. La solution contenant les ions chlorure et oxonium, nous pouvons exprimer sa conductivité à l'aide de la loi de Kohlrausch :

$$\sigma = \lambda_{\text{H}_3\text{O}^+} [\text{H}_3\text{O}^+] + \lambda_{\text{Cl}^-} [\text{Cl}^-]$$

Par ailleurs, l'électroneutralité de la solution impose que  $[\text{Cl}^-] = [\text{H}_3\text{O}^+]$ , et cette dernière concentration peut se déterminer simplement à partir de la définition du pH :

$$[\text{H}_3\text{O}^+] = C_{\text{réf}} \times 10^{-\text{pH}}$$

En combinant ces deux expressions, nous trouvons l'égalité :

$$\sigma = (\lambda_{\text{H}_3\text{O}^+} + \lambda_{\text{Cl}^-}) \times C_{\text{réf}} \times 10^{-\text{pH}}$$

d'où nous extrayons finalement la grandeur demandée :

$$\lambda_{\text{Cl}^-} = \sigma \times \frac{10^{+\text{pH}}}{C_{\text{réf}}} - \lambda_{\text{H}_3\text{O}^+}$$

Il ne reste plus qu'à mener l'application numérique, avec une vigilance particulière concernant les unités :

$$\begin{aligned} \lambda_{\text{Cl}^-} &= \frac{4,26 \text{ mS}}{10^{-2} \text{ m}} \times \frac{10^{2,00}}{10^3 \text{ mol.m}^{-3}} - 35,0 \text{ mS.m}^2.\text{mol}^{-1} \\ &= 7,6 \text{ mS.m}^2.\text{mol}^{-1} \end{aligned}$$

5. Lorsque l'on introduit un indicateur coloré acidobasique dans une solution, les formes acide et basique sont présentes dans des proportions variables selon le pH du milieu :

- dans un milieu très acide (donc très riche en protons), c'est la forme acide qui domine ;
- dans un milieu très basique (donc très pauvre en protons), c'est la forme basique qui domine.

On peut en tirer plusieurs conclusions :

- La conservation de la matière indique que les quantités de matière respectives en forme acide et en forme basique

sont complémentaires : si le pH augmente, des molécules de forme acide passeront sous forme basique, s'il diminue ce sera l'inverse. Mais en fin de compte, la quantité totale de molécules d'indicateur, formes acide et basique confondues, restera la même. En ramenant ces quantités de matière au volume de solution, on obtient l'égalité :

$$[\text{AH}] + [\text{A}^-] = C$$

- Dans un milieu à pH = 1,0 (très acide, donc), la concentration en base sera très inférieure à celle en acide et l'on aura :

$$[\text{AH}] \approx C \quad \text{et} \quad [\text{A}^-] \approx 0$$

Il s'ensuit que l'absorbance de cette solution s'exprimera, d'après la loi de Beer-Lambert :

$$A_1(\lambda) = A_{\text{AH}}(\lambda) = k_{\text{AH}}(\lambda) \times C$$

- Dans un milieu à pH = 13,0 (très basique), la concentration en acide sera très inférieure à celle en base et l'on aura :

$$[\text{A}^-] \approx C \quad \text{et} \quad [\text{AH}] \approx 0$$

Il s'ensuit que l'absorbance se résumera cette fois à :

$$A_{13}(\lambda) = A_{\text{A}^-}(\lambda) = k_{\text{A}^-}(\lambda) \times C$$

Si nous nous plaçons à une longueur d'onde  $\lambda_0$  telle que  $A_1(\lambda_0) = A_{13}(\lambda_0)$ , nous obtenons :

$$k_{\text{AH}}(\lambda_0) \times C = k_{\text{A}^-}(\lambda_0) \times C \Leftrightarrow k_{\text{AH}}(\lambda_0) = k_{\text{A}^-}(\lambda_0)$$

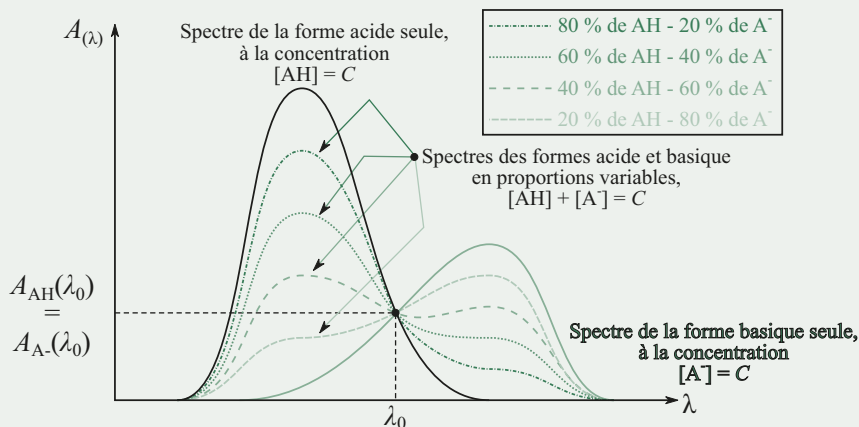
Dans ces conditions, pour une solution de pH quelconque et contenant donc les formes acide et basique de l'indicateur dans des proportions variables, l'absorbance totale s'exprime :

$$\begin{aligned} A_{\text{tot}}(\lambda_0) &= A_{\text{AH}}(\lambda_0) + A_{\text{A}^-}(\lambda_0) \\ &= k_{\text{AH}}(\lambda_0) \times [\text{AH}] + k_{\text{A}^-}(\lambda_0) [\text{A}^-] \end{aligned}$$

Les constantes de la loi de Beer-Lambert étant égales entre elles (nous noterons  $k$  cette valeur commune), et en exprimant  $[\text{A}^-] = C - [\text{AH}]$ , il vient :

$$A_{\text{tot}}(\lambda_0) = k \times [\text{AH}] + k \times (C - [\text{AH}]) = k \times C$$

Ce point est appelé **point isobestique** du spectre.



## Vers la prépa

Nous savons que la valeur du pH est directement liée à celle de  $[H_3O^+]_f$ . Nous pouvons donc déplacer le problème de l'expression du pH en fonction de la température, à celui de l'expression de  $[H_3O^+]_f$  en fonction de la température.

Les ions oxonium et hydroxyde étant en concentrations égales d'après l'énoncé, nous pouvons écrire :

$$[HO^-]_f = [H_3O^+]_f$$

d'où nous déduisons une nouvelle expression du produit ionique de l'eau :

$$K_e = [H_3O^+]_f^2 \Leftrightarrow [H_3O^+]_f = \sqrt{K_e}$$

Le problème se résume donc désormais à l'expression de  $K_e$  en fonction de  $T$ . Or cette dépendance est contenue dans la relation fournie par l'énoncé. En intégrant celle-ci par rapport à la variable  $T$ , nous trouvons :

$$\ln K_e = -\frac{\alpha}{T} + \text{cte}$$

Nous pouvons isoler  $K_e$  :

$$K_e = A \times e^{-\frac{\alpha}{T}}$$

où  $A = e^{\text{cte}}$  est une constante déterminée par une condition particulière. Dans le cas présent nous connaissons la valeur du pH (donc celle de  $[H_3O^+]_f$ , donc celle de  $K_e$ ) pour  $T_0 = 298 \text{ K}$ .

En notant  $K_{e,0}$  cette valeur, il vient :

$$K_{e,0} = A \times e^{-\frac{\alpha}{T_0}} \Leftrightarrow A = K_{e,0} \times e^{+\frac{\alpha}{T_0}}$$

En injectant cette expression dans l'expression générique de  $K_e$ , nous trouvons :

$$K_e(T) = K_{e,0} \times e^{\alpha \left( \frac{1}{T_0} - \frac{1}{T} \right)}$$

De là nous pouvons déduire l'expression de  $[H_3O^+]_f$  :

$$[H_3O^+]_f(T) = \sqrt{K_e(T)} = \sqrt{K_{e,0}} \times e^{\frac{\alpha}{2} \left( \frac{1}{T_0} - \frac{1}{T} \right)}$$

Enfin nous trouvons l'expression du pH :

$$\begin{aligned} \text{pH}(T) &= -\log([H_3O^+]_f) \\ &= -\log(\sqrt{K_{e,0}}) - \frac{\alpha}{2} \left( \frac{1}{T_0} - \frac{1}{T} \right) \times \log(e) \end{aligned}$$

Dans le cas particulier où  $T = T_0$ , cette égalité donne, d'après l'énoncé :

$$\text{pH}_0 = -\log(\sqrt{K_{e,0}}) = 7,0$$

Notons alors  $\Delta T = T - T_0$  l'écart à la température  $T_0$  ; nous avons ainsi :

$$T = T_0 + \Delta T$$

Dans ces conditions,  $\frac{1}{T}$  s'exprime :

$$\frac{1}{T} = \frac{1}{T_0 + \Delta T} = \frac{1}{T_0} \times \frac{1}{1 + \frac{\Delta T}{T_0}}$$

L'expression du pH devient alors :

$$\text{pH}(T) = \text{pH}_0 - \frac{\alpha \times \log(e)}{2T_0} \times \left( 1 - \frac{1}{1 + \frac{\Delta T}{T_0}} \right)$$

Pour de faibles écarts de température, autrement dit si  $\Delta T \ll T_0$ , nous pouvons mener l'approximation :

$$\frac{1}{1 + \frac{\Delta T}{T_0}} \approx 1 - \frac{\Delta T}{T_0}$$

L'expression précédente devient alors :

$$\text{pH}(T) = \text{pH}_0 - \frac{\alpha \times \log(e)}{2T_0} \times \frac{\Delta T}{T_0}$$

En menant l'application numérique avec les valeurs proposées par l'énoncé, nous trouvons :

$$\text{pH}(T) = 7,0 - 5,0 \times \frac{\Delta T}{T_0}$$

Ou si l'on préfère, en explicitant enfin  $\Delta T = T - T_0$  et  $T_0 = 298 \text{ K}$  :

$$\text{pH}(T) = 12 - 1,7 \cdot 10^{-2} \text{ K}^{-1} \times T$$

Nous constatons donc que le pH de l'eau varie de 0,017 unité de pH par kelvin de variation de température.

Nous comprenons donc que pour quelques kelvin de variation (typiquement, si la température est de 20 °C plutôt que les 25 officiellement requis), le pH de l'eau peut encore être considéré comme sensiblement égal à 7,0 (décalage de moins de 0,1 unité de pH).

Pour des écarts plus importants, en revanche, ce modèle doit être remis en cause à deux égards :

- la variation devient de moins en moins négligeable par rapport à la valeur réelle : par exemple pour 30 K d'écart de température, le pH varie cette fois de 0,5 unité ;
- le modèle affine ci-dessus suppose des variations de température faibles devant les 298 K de référence, donc il serait inapproprié de l'utiliser en l'état pour des écarts supérieurs à, justement, 30 K.

### 3.1 Modélisation d'une transformation chimique

#### — Quelle est la différence entre une transformation, une réaction et une équation chimiques ?

Un système est le siège d'une transformation chimique dès lors que la nature chimique des espèces qui le constituent évolue. Cette nature étant définie par la structure des molécules, une transformation chimique repose donc sur une modification de ladite structure, autrement dit une réorganisation des liaisons de covalence. Ces liaisons s'appuyant exclusivement sur le cortège électronique des atomes, une transformation chimique n'engage donc aucune altération des noyaux, ce qui entraîne une **conservation du nombre d'atomes de chaque élément chimique**.

Pour modéliser une transformation chimique on pose comme hypothèse que celle-ci repose sur le réarrangement des atomes de certaines espèces (appelées **réactifs**), pour en former de nouvelles (appelées **produits**). On associe ainsi à la transformation chimique (expérimentale) un modèle appelé **réaction chimique**. Cette réaction est quant à elle formalisée par une équation, appelée **équation de réaction**.

Dans cette équation :

- les **réactifs** figurent à **gauche**, séparés par des signes + ;
- les **produits** figurent à **droite**, séparés par des signes + ;
- la transformation est symbolisée par une **flèche** allant des réactifs vers les produits, ou un signe « = » si la réaction aboutit à un équilibre chimique.

Outre la nature des espèces qu'elle engage, une équation de réaction doit également rendre compte des **proportions dans lesquelles les réactifs et les produits sont respectivement consommés et formés** au cours de la transformation que l'on modélise.

On appelle alors :

- **stœchiométrie** d'une réaction, la donnée de ces proportions ;
- **nombres stœchiométriques**, les nombres traduisant ces proportions dans l'équation de réaction, qui doivent être **les plus petits entiers possibles**.

**Remarque :** vous emploierez parfois, en CPGE, des nombres stœchiométriques fractionnaires, voire négatifs. Il s'agit d'artifices de calcul permettant de simplifier certaines expressions mathématiques.

Dans les faits, naturellement, ce sont toujours des nombres entiers et positifs de molécules et/ou d'ions qui sont consommés ou formés.

Par ailleurs, les modèles de réaction que nous utilisons supposant que les produits d'une réaction sont bâtis à partir des réactifs qu'elle engage, les quantités de produits formées doivent répondre directement aux quantités de réactifs consommées.

Pour que le modèle soit cohérent, l'équation de la réaction doit donc intégrer cette contrainte, en respectant des **lois de conservation** garantissant qu'elle est **équilibrée**, c'est-à-dire que l'on trouve en **quantités égales** du côté des réactifs d'une part, et du côté des produits d'autre part :

- le nombre d'atomes de chaque élément ;
- la charge électrique globale.

**Remarque :** le fait que la charge électrique globale soit conservée entre les membres de gauche et de droite de l'équation de réaction ne signifie absolument pas que cette charge doive être nulle. S'il est vrai que le système chimique doit rester électriquement neutre, cette électroneutralité peut tout-à-fait s'appuyer sur des espèces spectatrices, qui par définition ne figurent pas dans l'équation de réaction.

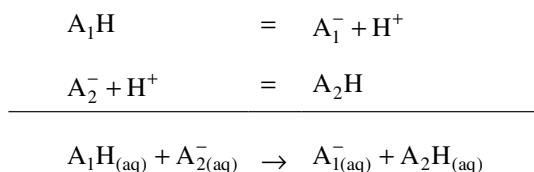
Enfin, une équation de réaction doit également préciser les **états physiques** (solide, liquide, gazeux ou aqueux) dans les conditions de l'expérience, des espèces qu'elle engage (en indice des espèces).

## — Combien de réactions chimiques existe-t-il ?

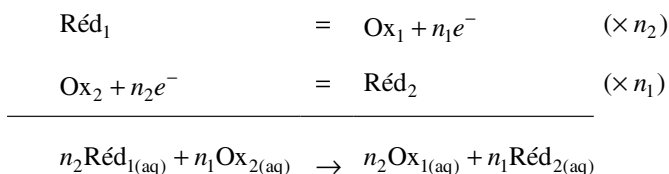
L'univers comptant (en l'état actuel de nos connaissances) un nombre fini d'atomes, on ne peut techniquement envisager qu'un nombre fini d'espèces chimiques, et donc également un nombre fini de transformations entre ces espèces. Mais *a priori* cela fait quand même beaucoup, et nous serions d'un optimisme coupable en espérant les répertorier toutes (surtout en moins de 300 pages).

On constate cependant que toutes les réactions observées appartiennent à l'une ou l'autre des **4 grandes catégories** suivantes :

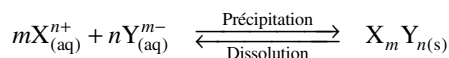
- les réactions acido-basiques, qui reposent sur le transfert d'un ion  $\text{H}^+$  entre :
  - une espèce donneuse appelée **acide** (ici notation générique  $\text{A}_1\text{H}$ ), et qui par suite va donner sa base conjuguée (couple  $\text{A}_1\text{H} / \text{A}_1^-$ ),
  - une espèce receveuse appelée **base** (ici notation générique  $\text{A}_2^-$ ), et qui par suite va donner son acide (couple  $\text{A}_2\text{H} / \text{A}_2^-$ ).



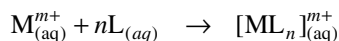
- les réactions d'oxydoréduction, qui reposent sur le transfert d'un ou plusieurs électrons entre :
  - une espèce donneuse appelée **réducteur** (notation générique  $\text{Réd}_1$ ), et qui par suite va donner son oxydant conjugué (couple  $\text{Ox}_1 / \text{Réd}_1$ ),
  - une espèce receveuse appelée **oxydant** (notation générique  $\text{Ox}_2$ ), et qui par suite va donner son réducteur conjugué (couple  $\text{Ox}_2 / \text{Réd}_2$ ).



- les réactions de précipitation (et leur réciproque : réaction de dissolution), qui reposent sur l'association de cations et d'anions pour former un solide électriquement neutre :



- les réactions de complexation, qui reposent sur l'association d'un cation métallique avec une ou plusieurs molécules appelées **ligands** par des liaisons de coordination, pour former un édifice appelé complexe :

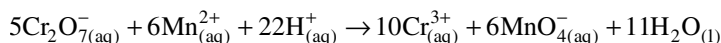


les liaisons de coordination diffèrent des liaisons de covalence, et sont décrites par une théorie appelée théorie du champ cristallin. Les complexes possèdent des propriétés particulières et constituent donc des entités à part, ce qui vaut à ces réactions un statut également à part. Post-baccalauréat, elles sont souvent étudiées en même temps que les réactions de précipitation, dont elles sont formellement assez proches.

## — Quelle relation existe-t-il entre les quantités consommées en réactifs et formées en produits lorsque la réaction a lieu ?

Nous avons vu que la stœchiométrie d'une réaction indiquait les proportions dans lesquelles les réactifs et produits qu'elle engage étaient respectivement consommés et formés. Il est intéressant de noter que ces proportions lient entre elles les quantités de matière de ces espèces, à travers un dénominateur commun : le nombre de fois où la réaction a eu lieu.

Ainsi l'équation de réaction suivante, par exemple :



peut-elle être lue littéralement comme une recette de cuisine :

| Pour réaliser cette réaction, réunir :         | Mijoter à feu doux, et recueillir en fin de cuisson : |
|------------------------------------------------|-------------------------------------------------------|
| 5 ions dichromate $Cr_2O_7^{2-}$ pas trop mûrs | 10 ions chrome (III) $Cr^{3+}$                        |
| 6 beaux ions manganèse (II) $Mn^{2+}$          | 6 ions permanganate $MnO_4^{-}$                       |
| 22 ions $H^+$ bien fermes                      | 11 molécules d'eau $H_2O$                             |

Et ce, chaque fois que la réaction aura lieu.

On peut ainsi dire que si une espèce est engagée dans l'équation de réaction avec un nombre stœchiométrique  $\alpha$ , alors la quantité consommée (si l'espèce considérée est un réactif) ou formée (s'il s'agit d'un produit) sera égale à  $\alpha$  fois le nombre de fois où la réaction aura eu lieu.

Les quantités de molécules (et donc le nombre de fois où la réaction a lieu entre ces molécules) étant, nous le savons, très importantes, on préférera travailler en termes de quantités de matière (mole, mon amie). Et il devient donc particulièrement pratique de s'appuyer sur la quantité de matière (autrement dit le nombre de moles) de fois où la réaction a eu lieu.

C'est ce nombre de moles de fois où la réaction a eu lieu que l'on appelle **avancement** (généralement noté  $x$ ) de la réaction.

**Remarque :** en CPGE, vous trouverez parfois encore l'avancement représenté par la lettre grecque  $\xi$  (prononcer « ksi »).

On peut ainsi exprimer, sur cette base :

- la quantité de matière **consommée en un réactif** lorsque la réaction a eu lieu  $x$  moles de fois :

$$n_{R, \text{conso}}(x) = r \times x$$

- la quantité de matière **formée en un produit** lorsque la réaction a eu lieu  $x$  moles de fois :

$$n_{P, \text{formé}}(x) = p \times x$$

$r$  et  $p$  désignant leurs nombres stœchiométriques respectifs.

De manière générale, on peut ainsi mettre n'importe quelle quantité de matière consommée en un réactif ou formée en un produit, en relation avec n'importe quelle autre :

$$x = \frac{n_{R_1, \text{conso}}}{r_1} = \frac{n_{R_2, \text{conso}}}{r_2} = \dots = \frac{n_{P_1, \text{formé}}}{p_1} = \frac{n_{P_2, \text{formé}}}{p_2} = \dots$$

## — Comment exprimer la quantité de matière effectivement présente en une espèce dans le milieu réactionnel, en fonction de l'avancement ?

Ayant exprimé les quantités de matière consommées et formées dans le milieu réactionnel, respectivement en réactifs et en produits, il suffit d'exprimer la conservation de la matière au cours du processus de réaction :

- pour un réactif, ce qui est effectivement présent est égal à ce qui se trouvait au départ dans le milieu réactionnel, **minoré** de ce que la réaction en a consommé :

$$n_R(x) = n_{R,i} - n_{R, \text{conso}}(x) \Leftrightarrow n_R(x) = n_{R,i} - r \times x$$

- pour un produit, ce qui est effectivement présent est égal à ce qui se trouvait au départ dans le milieu réactionnel, **majoré** de ce que la réaction en a formé :

$$n_P(x) = n_{P,i} + n_{P, \text{formé}}(x) \Leftrightarrow n_P(x) = n_{P,i} + p \times x$$

## — Comment effectuer un bilan de matière dans le cas d'une réaction totale ?

On appelle **bilan de matière** l'inventaire des quantités de matière des différentes espèces engagées dans une réaction, lorsque celle-ci prend fin. Déterminer quand au juste une réaction prend fin peut être assez complexe lorsque celle-ci aboutit à un équilibre (notamment dans le cas des équilibres acido-basiques).

En effet, les équilibres chimiques résultent du fait que les produits d'une réaction peuvent parfois (souvent) réagir entre eux pour redonner les réactifs de départ. À la réaction consommant les réactifs et formant les produits (réaction dite **directe**), va donc se superposer une

seconde, consommant les produits et régénérant les réactifs (réaction dit **inverse**). À ce jeu-là, les réactifs ne sont donc jamais consommés en totalité et l'évolution macroscopique du système semble bloquée malgré le fait qu'il reste encore de tous les réactifs nécessaires. Nous étudierons ce point en détail plus loin.

Cependant dans le cas où les produits sont indifférents les uns aux autres, seule la réaction directe a lieu, et l'évolution macroscopique du système obéit alors à une règle simple : elle se poursuit tant que tous les réactifs sont présents, et **s'arrête si et seulement si l'un au moins d'entre eux parvient à épuisement**.

Du point de vue méthodologique, nous travaillerons en 3 temps :

1. Supposer tour à tour que chacun des réactifs  $R_j$  est épuisé, et en déduire pour chacun la **valeur maximale  $x_{\max,j}$**  autorisée à l'avancement par le réactif en question.
2. Repérer, parmi tous les réactifs, celui ou ceux qui impose(nt) la plus petite valeur maximale  $x_{\max,j^*}$  et en déduire :
  - la nature de ce(s) réactif(s), appelé(s) **réactif(s) limitant(s)** (les autres étant alors qualifiés de réactifs **en excès**) ;
  - la **valeur maximale  $x_{\max}$  effectivement imposée** à l'avancement par ce(s) réactif(s) limitant(s), qui donc précise combien de fois la réaction a pu se faire avant de devoir s'arrêter parce que ce(s) réactif(s) lui faisai(en)t défaut.
3. Injecter cette valeur dans les expressions des quantités de matière des différentes espèces, et en déduire pour chacune quelle quantité de matière est présente dans le milieu réactionnel lorsque la réaction prend fin, ce qui nous donne alors le fameux inventaire (« bilan de matière ») invoqué en ouverture.

En pratique, nous commencerons par exprimer la quantité de matière de chaque espèce sous l'une des deux formes suivantes :

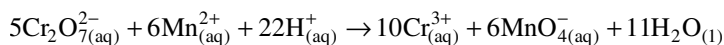
- Pour un réactif  $R_j$  intervenant dans l'équation de réaction avec le nombre stœchiométrique  $r_j$ , initialement présent en quantité  $n_{R_j,i}$  :

$$n_{R_j}(x) = n_{R_j,i} - n_{R_{\text{cons},j}} = n_{R_j,i} - r_j x$$

- Pour un produit  $P_k$  intervenant dans l'équation de réaction avec le nombre stœchiométrique  $p_k$ , initialement présent en quantité  $n_{P_k,i}$  :

$$n_{P_k} = n_{P_k,i} + n_{P_{\text{form},k}} = n_{P_k,i} + p_k x$$

Prenons par exemple la réaction d'équation :



Notons  $n_1$  et  $n_2$  les quantités de matière respectives d'ions dichromate  $\text{Cr}_2\text{O}_7^{2-}$  et manganèse (II)  $\text{Mn}^{2+}$ . Nous pouvons ainsi exprimer, pour un avancement  $x$  :

- $n_1(x) = n_{1,i} - 5x$  ;
- $n_2(x) = n_{2,i} - 6x$ .

En supposant que les ions chrome (III)  $\text{Cr}^{3+}$  (quantité de matière  $n_3$ ) et permanganate  $\text{MnO}_4^{-}$  (quantité de matière  $n_4$ ) sont initialement absents du milieu réactionnel ( $n_{3,i} = n_{4,i} = 0$ ), nous pouvons alors également exprimer :

- $n_3(x) = 10x$  ;
- $n_4(x) = 6x$ .

Vous avez appris au lycée à dresser un tableau d'avancement, permettant une vision synoptique des évolutions des différentes quantités de matière en fonction de l'avancement. Dans le cas présent, ce tableau prend la forme :

|       | $5\text{Cr}_2\text{O}_7^{2-}(\text{aq}) + 6\text{Mn}^{2+}(\text{aq}) + 22\text{H}^+(\text{aq}) \rightarrow 10\text{Cr}^{3+}(\text{aq}) + 6\text{MnO}_4^-(\text{aq}) + 11\text{H}_2\text{O}(\text{l})$ |                  |       |         |        |         |
|-------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------|-------|---------|--------|---------|
| $E_i$ | $n_{1,i}$                                                                                                                                                                                             | $n_{2,i}$        | Excès | 0       | 0      | Solvant |
| $x$   | $n_{1,i} - 5x$                                                                                                                                                                                        | $n_{2,i} - 6x$   | Excès | $10x$   | $6x$   | Solvant |
| $E_f$ | $n_{1,i} - 5x_f$                                                                                                                                                                                      | $n_{2,i} - 6x_f$ | Excès | $10x_f$ | $6x_f$ | Solvant |

Ce type d'outil peut encore vous servir en CPGE, même si son usage y est moins systématique (vous mènerez le même travail de fond, mais sans forcément recourir à ce tableau qui finalement n'est jamais qu'un accessoire de rangement).

L'objet de ce tableau est de donner un descriptif factuel de l'état du système chimique, en particulier de son état final. Dans la dernière ligne, c'est donc bien la valeur **finale** de l'avancement (notée  $x_f$ ) qui doit figurer. La valeur finale n'est pas toujours égale à  $x_{\max}$ . En particulier pour certains équilibres (que nous allons voir un peu plus loin), elle y est significativement inférieure.  $x_f$  ne peut donc être assimilée à  $x_{\max}$ , que si la réaction est totale (ou à la rigueur qu'elle aboutit à un équilibre très marqué vers la droite).

Supposons donc que la réaction soit totale, c'est-à-dire qu'elle se poursuive aussi longtemps qu'il reste des molécules de chaque réactif dans le milieu réactionnel. Calculons alors l'avancement maximal autorisé à la réaction par le réactif  $R_j$ , c'est-à-dire la valeur  $x_{\max,j}$  de l'avancement qui annule la quantité de matière du réactif  $R_j$ .

On trouve facilement que :

$$x_{\max,j} = \frac{n_{R_j,i}}{r_j}$$

$$\frac{x_{\max,j} n_{R_j,i}}{r_j} \quad \begin{array}{l} \text{en mol} \\ \text{sans unité} \end{array}$$

En reprenant l'exemple précédent (et en supposant que la réaction se fait bien de la gauche vers la droite), nous trouvons ainsi :

- $x_{\max,1} = \frac{n_{1,i}}{5} ;$
- $x_{\max,2} = \frac{n_{2,i}}{6}.$

Le réactif qui s'épuise en premier, et qui donc va par cet épuisement provoquer la fin de la réaction, est appelé **réactif limitant** (les autres étant alors qualifiés de réactifs en excès). Il s'agit donc de celui qui autorise la réaction à se faire le moins de fois, et impose donc **la plus petite valeur maximale**  $x_{\max,j}$  à l'avancement.

Plusieurs réactifs peuvent être limitants, pour peu qu'ils s'éteignent simultanément, autrement dit qu'ils imposent une même valeur maximale à l'avancement (et que cette valeur commune soit la plus petite parmi celles associées aux différents réactifs, bien sûr). Il est

intéressant de noter qu'en écrivant l'égalité de ces valeurs maximales, on trouve ainsi que deux réactifs sont limitants seulement si :

$$x_{\max,1} = x_{\max,2} \Leftrightarrow \frac{n_{1,i}}{r_1} = \frac{n_{2,i}}{r_2} \Leftrightarrow \frac{n_{1,i}}{n_{2,i}} = \frac{r_1}{r_2}$$

En d'autres termes, deux réactifs s'éteignent simultanément si leurs quantités de matière initiales sont introduites dans un rapport (ou si l'on préfère, dans des *proportions*) précisément égal au rapport de leurs nombres stœchiométriques respectifs. Ces proportions particulières sont de ce fait appelées **proportions stœchiométriques**, et l'on retiendra que deux réactifs introduits en proportions stœchiométriques s'éteignent forcément pour une même valeur de l'avancement.

Ainsi, dans le cas de la réaction vue plus haut, prenons par exemple :

|         | $n_{1,i}$ (mmol) | $x_{\max,1} = \frac{n_{1,i}}{5}$ | $n_{2,i}$ (mmol) | $x_{\max,2} = \frac{n_{2,i}}{6}$ |
|---------|------------------|----------------------------------|------------------|----------------------------------|
| Cas n°1 | 12,0             | 2,40                             | 12,0             | <b>2,00</b>                      |
| Cas n°2 | 10,0             | <b>2,00</b>                      | 12,0             | <b>2,00</b>                      |
| Cas n°3 | 10,0             | <b>2,00</b>                      | 14,4             | 2,40                             |

Nous avons alors :

- **Cas n° 1** : les ions manganèse (II) s'épuisent dès que  $x = 2,00$  mmol, tandis que les ions dichromate pourraient aller jusqu'à  $x = 2,40$  mmol. Les ions manganèse (II) constituent donc le réactif limitant, et la réaction prend fin lorsque  $x = x_{\max,2} = 2,00$  mmol.
- **Cas n° 2** : les ions dichromate et les ions manganèse (II) s'épuisent simultanément, lorsque  $x = 2,00$  mmol. Ils ont été introduits en proportions stœchiométriques et sont donc tous les deux limitants, et la réaction prend fin lorsque  $x = x_{\max,1} = x_{\max,2} = 2,00$  mmol.
- **Cas n° 3** : les ions dichromate s'épuisent dès que  $x = 2,00$  mmol, tandis que les ions manganèse (II) pourraient aller jusqu'à  $x = 2,40$  mmol. Les ions dichromate constituent donc le réactif limitant, et la réaction prend fin lorsque  $x = x_{\max,1} = 2,00$  mmol.

Une fois connue la valeur maximale de l'avancement, on peut réaliser un bilan de matière, c'est-à-dire calculer la quantité de matière de chaque espèce présente dans le milieu réactionnel en fin de réaction. Il suffit pour ce faire de reprendre l'expression de chaque quantité de matière en fonction de l'avancement, et d'y substituer  $x$  par la valeur  $x_{\max}$  trouvée.

Prenons par exemple le cas n° 1 de l'exemple précédent. Nous savons qu'en fin de réaction,  $x_{\max} = x_{\max,2} = 2,00$  mmol. Nous pouvons alors en déduire les valeurs finales (en supposant les produits initialement absents du milieu réactionnel :  $n_{3,i} = n_{4,i} = 0$ ) :

- $n_{1,f} = n_{1,i} - 5x_{\max} = 2,00$  mmol ;
- $n_{2,f} = n_{2,i} - 6x_{\max} = 0$  ;
- $n_{3,f} = 10x_{\max} = 20,0$  mmol ;
- $n_{4,f} = 6x_{\max} = 12,0$  mmol.

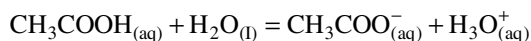
Répetons encore une fois que le fait d'assimiler la valeur finale  $x_f$  de l'avancement, à la valeur maximale  $x_{\max}$  imposée par le(s) réactif(s) limitant(s), vaut uniquement dans le cas d'une réaction totale.

## — La valeur maximale autorisée à l'avancement par les quantités initiales de réactifs est-elle toujours atteinte ?

Durant la plus grande partie de votre scolarité au lycée, les transformations chimiques que vous avez étudiées étaient supposées totales. C'est-à-dire qu'elles se poursuivaient aussi longtemps que chacun des réactifs nécessaires à sa réalisation était encore présent dans le milieu réactionnel.

La réalité est en fait toute autre : la plupart des transformations aboutissent, au bout d'une durée plus ou moins longue, à une situation dans laquelle coexistent tous les produits de la transformation, mais également tous les réactifs (en proportions variables selon les situations). Chaque quantité de matière se trouve alors figée à une valeur fixe et le système n'est plus (macroscopiquement parlant) siège d'aucune évolution, sans pourtant qu'aucun réactif ne soit arrivé à épuisement.

Par exemple, lors de la mise en solution de 0,010 mol d'acide éthanóique ( $\text{CH}_3\text{COOH}$ ) dans 1,0 L d'eau, se produit la réaction d'équation :



Une analyse rapide nous permet de calculer une valeur maximale de l'avancement, égale à  $x_{\text{max}} = 1,0 \cdot 10^{-2}$  mol. Or l'expérience montre que la valeur finale de l'avancement se stabilise à  $4,2 \cdot 10^{-4}$  mol, soit presque 25 fois moins que la valeur maximale calculée dans l'hypothèse d'une réaction totale.

On constate donc que la valeur finale  $x_f$  de l'avancement d'une réaction n'est pas nécessairement égal à la valeur maximale  $x_{\text{max}}$  que lui autorisent les réactifs engagés dans cette réaction, et l'on aura bien soin par la suite de différencier *a priori* ces valeurs.

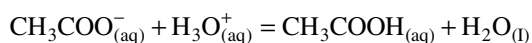
## — 3.2 Interprétation des équilibres chimiques et détermination de la composition d'un système à l'équilibre

### — Une réaction à l'équilibre est-elle réellement bloquée ?

En réalité, à l'échelle microscopique, la réaction se poursuit en permanence. Se pose alors la question de savoir ce qui l'empêche de se poursuivre jusqu'au bout. Il peut être intéressant d'observer ce qu'il se passe si l'on met les seuls produits de cette réaction en présence les uns des autres. Dans le cas de la réaction vue précédemment, il s'agirait par exemple d'introduire dans un même becher :

- de l'éthanoate de sodium (solide ionique, dont la dissolution en solution aqueuse donne des ions éthanoate) ;
- de l'acide chlorhydrique (solution aqueuse de chlorure d'hydrogène, dont la dissolution en solution aqueuse donne des ions  $\text{H}^+$ , qui se combinent aussitôt avec de l'eau pour former des ions oxonium  $\text{H}_3\text{O}^+$ ).

On peut alors observer que les ions  $\text{H}_3\text{O}^+$  sont consommés : ils réagissent donc avec les ions éthanoate. L'équation de cette réaction donne alors :



On constate ainsi que les produits formés au cours d'une transformation peuvent parfois réagir entre eux, et redonner les réactifs de départ. Cette réaction en sens inverse va donc consommer les produits au fur et à mesure de leur formation et régénérer les réactifs, empêchant par là-même la première réaction de se faire jusqu'à leur consommation complète.

La réaction de départ n'est donc en réalité pas bloquée : elle se poursuit mais ne connaît pas de fin puisque les réactifs, sans cesse régénérés par la réaction inverse, ne s'épuisent jamais. On dit de l'équilibre chimique qu'il s'agit d'un **équilibre dynamique**.

La réaction pouvant se faire dans les deux sens, chaque espèce peut être considérée à la fois comme consommée et formée, et les notions de réactif et de produit peuvent devenir confuses. C'est pourquoi on définit certaines conventions. Ainsi, dans une équation de réaction, on appelle :

- **réactif** toute espèce chimique placée à gauche du signe « = » de l'équation ;
- **produit** toute espèce chimique placée à droite du signe « = » de l'équation ;
- **sens direct** de l'équation le sens de la gauche vers la droite ;
- **sens indirect** (ou sens inverse) de l'équation le sens de la droite vers la gauche.

## — Comment interpréter l'existence de l'état d'équilibre, en termes de cinétique chimique ?

Ainsi décrit, on peut supposer que l'équilibre chimique s'établit lorsque réactifs et produits sont présents en quantités comparables. La réalité est un peu plus complexe. En effet, les deux réactions ne se font pas *a priori* à la même vitesse (cf. chapitre 5 pour plus de détails à ce sujet). On peut même ajouter que plus la réaction se fait vite dans le sens direct (c'est-à-dire plus les réactifs réagissent facilement ensemble), plus elle se fait lentement dans le sens indirect (il sera d'autant plus difficile de remonter le cours d'une réaction, que celle-ci s'est faite avec plus de facilité). La réciproque est naturellement vraie aussi.

Ainsi, si la réaction est très rapide dans le sens direct, on n'observera que peu de réactif à l'équilibre. En effet, dès qu'une molécule de réactif apparaît, elle est consommée rapidement et celui-ci n'existe donc que peu de temps dans le milieu réactionnel. Le produit, au contraire, consommé par la réaction inverse à un rythme assez lent, sera très présent.

Réciproquement, si la réaction est très lente dans le sens direct, on observera beaucoup de réactif à l'équilibre. En effet, lorsqu'une molécule de réactif apparaît, elle prend son temps pour être consommée. Le produit, au contraire, consommé par la réaction inverse à un rythme rapide, sera peu présent.

## — Quels sont les facteurs influençant la composition d'un système à l'équilibre chimique ?

L'état d'équilibre résulte donc en fait d'un compromis entre deux réactions inverses l'une de l'autre, et de cinétiques *a priori* différentes. La composition d'un système à l'équilibre chimique dépend donc de tout facteur susceptible d'influencer les cinétiques de ses réactions directe et inverse.

Nous savons par exemple que la cinétique d'une réaction est accélérée par :

- **L'augmentation de la concentration des réactifs** : ainsi, si l'on augmente les concentrations des réactifs, la réaction directe va se faire plus vite, provoquant une augmentation des quantités de produits formées. Réciproquement, l'introduction dans le milieu réactionnel d'un excédent de produit favorisera la régénération des réactifs. Il est intéressant de noter que deux réactions réciproques voient leurs vitesses varier en sens inverses l'un de l'autre à mesure que le système évolue : si les réactifs sont très présents au départ et les produits peu, la réaction directe est rapide et l'inverse est lente. Mais au fur et à mesure les quantités de réactifs baissent tandis que celles de produits augmentent ; la réaction directe ralentit et la réaction inverse accélère. Cette évolution se poursuit jusqu'à ce que les quantités de réactifs et de produits permettent aux deux réactions de se dérouler à la même vitesse, situation correspondant à l'équilibre où toute consommation d'une quantité de réactifs entraîne la consommation simultanée d'une quantité équivalente de produit.
- **L'augmentation de température** : si  $T$  augmente, chacune des deux réactions (directe et inverse) va être accélérée, mais pas forcément dans les mêmes mesures. Concrètement, on peut montrer qu'une augmentation de température provoque un déplacement d'équilibre dans le sens endothermique de la réaction.
- **La présence d'un catalyseur** : tout catalyseur d'une réaction catalyse également la réaction inverse (ce point, très intéressant, nécessite pour être expliqué de détailler la notion d'état de transition, concept qui commence à sérieusement dépasser le cadre de cet ouvrage). En outre, l'accélération qui en résulte est la même pour les deux réactions. Donc, contrairement aux cas précédents, l'éventuelle présence d'un catalyseur **ne modifie pas** l'état d'équilibre auquel aboutit un système chimique (c'est même la caractéristique principale d'un catalyseur).



## — Y a-t-il quelque chose de conservé en toutes circonstances à l'équilibre, quelles que soient les conditions initiales ?

Selon ce qui précède, on peut se demander s'il existe bien un équilibre, puisque tant de facteurs semblent susceptibles de modifier l'état final du système qui en est le siège. Pourtant, si l'on mène à plusieurs reprises une même réaction aboutissant à un équilibre, en veillant à reconduire à chaque fois les mêmes conditions expérimentales, on constate qu'à chaque fois la composition finale du système est la même : l'équilibre est totalement reproductible et déterministe. Il existe donc bien quelque chose qui se conserve au cours du processus.

Toute la difficulté est d'identifier ce qui précisément est conservé, et de trouver le lien existant entre les compositions à l'équilibre de deux systèmes sièges d'une même réaction, mais réalisée pour chacun à partir de conditions initiales différentes.

Commençons par faire la part des choses. Nous venons de voir que la composition finale d'un système siège d'un équilibre chimique dépendait de la température du milieu, ainsi que des concentrations initiales des différentes espèces engagées dans cet équilibre.

Les rôles de ces deux facteurs peuvent être démontrés (vous le verrez en deuxième année de CPGE) mais réclament tout un arsenal de thermochimie dépassant largement le cadre de ce livre.

- **Concernant l'influence des concentrations initiales des différentes espèces** : ces démonstrations mettent en lumière une grandeur, dont la valeur à l'équilibre est invariablement la même pour une température donnée et ce, quelle que soit la composition initiale du système. Cette grandeur est appelée **quotient de réaction** de l'équilibre considéré, et sa valeur à l'équilibre est appelée **constante d'équilibre**.

- **Concernant l'influence de la température** : qualitativement, ils aboutissent au résultat mentionné plus haut (déplacement de l'équilibre dans le sens endothermique, consécutivement à une augmentation de température).

## — Comment exprime-t-on le quotient de réaction d'un équilibre chimique ?

Considérons un équilibre chimique, engageant plusieurs espèces (réactifs et produits). On attribue à chacune de ces espèces une contribution, appelée **activité** de cette espèce, dont l'expression **dépend de l'état** dans lequel se trouve l'espèce en question :

- S'il s'agit d'une espèce formant une **phase solide ou liquide**, ou du solvant, cette contribution est égale à 1.
- S'il s'agit d'une **espèce X solvatée** en faible concentration ( $\ll 1 \text{ mol.L}^{-1}$ ), cette contribution est égale au rapport  $\frac{[X]}{c_{\text{ref}}}$ , où  $[X]$  est la concentration effective de cette espèce en solution, et  $c_{\text{ref}}$  une concentration de référence, égale à  $1 \text{ mol.L}^{-1}$ .

**Remarque** : vous verrez en CPGE que si  $X$  est une espèce gazeuse, cette contribution prend encore une expression différente. Elle est égale au rapport  $\frac{P_X}{P_{\text{ref}}}$ , où  $P_X$  est la pression partielle de cette espèce

dans le milieu réactionnel, et  $P_{\text{ref}}$  une pression de référence, égale à 1 bar. La pression partielle est la pression qu'exercerait cette espèce, si elle était seule dans l'enceinte. On peut facilement montrer, à l'aide de la loi du gaz parfait, qu'elle s'exprime  $P_X = \frac{n_X}{n_{\text{tot}}} P_{\text{tot}}$ , où  $n_{\text{tot}}$  désigne la quantité de matière totale d'espèces gazeuses, et  $P_{\text{tot}}$  la pression totale dans le milieu réactionnel.

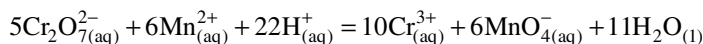
On appelle alors **quotient de réaction**, noté  $Q_r$ , le rapport de deux termes :

- Un numérateur, constitué du **produit des activités** des produits, **chacune élevée à une puissance égale à son nombre stœchiométrique** dans l'équation de réaction.
- Un dénominateur, constitué du **produit des activités** des réactifs, **chacune élevée à une puissance égale à son nombre stœchiométrique** dans l'équation de réaction.

Notons que les activités évoquées dans la définition ci-dessus sont sans unité (1, ou rapport de deux concentrations). Par suite, le quotient de réaction n'a pas d'unité non plus.

**Remarque** : comme nous l'avons vu précédemment pour le pH, si l'on choisit d'exprimer toutes les concentrations en  $\text{mol.L}^{-1}$ , alors tous les termes  $c_{\text{ref}}$  ont pour valeur 1 dans cette unité, et peuvent être occultés dans l'expression finale. Ceci permet définir la contribution d'une espèce solvatée au quotient de réaction, de manière plus simple : il s'agit de la valeur en  $\text{mol.L}^{-1}$  de la concentration de cette espèce, privée de son unité. De manière analogue, il est possible de se dispenser du terme  $P_{\text{ref}}$  si l'on prend soin de toujours exprimer la pression partielle en bar, puisque  $P_{\text{ref}}$  a pour valeur 1 dans cette unité. Attention cependant : il ne s'agit que d'artifices de calcul, et le quotient de réaction demeure une grandeur sans unité.

Reprenons l'exemple de la réaction vue en début de chapitre :



En choisissant d'exprimer toutes les concentrations en  $\text{mol.L}^{-1}$ , le quotient de réaction associé aurait pour expression :

$$Q_r = \frac{[\text{Cr}^{3+}]^{10} [\text{MnO}_4^{2-}]^6}{[\text{Cr}_2\text{O}_7^{2-}]^5 [\text{Mn}^{2+}]^6 [\text{H}^+]^{22}}$$

## — Quelles propriétés ce quotient de réaction possède-t-il ?

Ainsi que nous l'avons dit, l'intérêt du quotient de réaction est que la valeur qu'il prend à l'équilibre ne dépend finalement que de la température du milieu dans lequel a lieu ledit équilibre. Ceci à l'exclusion notamment des quantités initiales de réactifs et de produits. Cette valeur à l'équilibre est appelée **constante d'équilibre**. On la note usuellement  $K$ , éventuellement dotée d'un indice rappelant le type d'équilibre auquel elle se rapporte ( $K_A$  pour l'équilibre d'un acide avec l'eau,  $K_S$  pour un équilibre de solubilité, par exemple).

On peut ainsi écrire qu'à l'état final, lorsque l'équilibre est atteint, le système évolue toujours de manière à ce que son quotient de réaction adopte cette valeur privilégiée :  $Q_{r,f} = K$ .

Si nous reprenons l'exemple précédent, nous pouvons ainsi écrire que :

$$\frac{[\text{Cr}^{3+}]_f^{10} [\text{MnO}_4^-]_f^6}{[\text{Cr}_2\text{O}_7^{2-}]_f^5 [\text{Mn}^{2+}]_f^6 [\text{H}^+]_f^{22}} = K$$

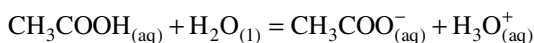
La suite des opérations est alors simple sur le principe : si l'on connaît la valeur de  $K$  et que l'on explicite la concentration de chaque espèce en solution, en fonction notamment de sa concentration initiale et de la valeur finale  $x_f$  de l'avancement, on finit par aboutir à une équation dont la seule inconnue est  $x_f$ . La résolution de cette équation permet de déterminer la valeur finale  $x_f$  de l'avancement, et par suite de déterminer les concentrations de toutes les espèces à l'équilibre.

Dans certains cas, cette équation est impossible à résoudre (comme dans l'exemple donné ci-dessus, où l'on aboutit à une équation en  $x_f$  du... 33<sup>e</sup> degré).

**Remarque :** dans une réaction de ce type, les ions  $\text{H}^+$  sont souvent apportés en très large excès. Il s'ensuit que la quantité de ces ions réagissant est négligeable devant leur quantité initiale dans le milieu réactionnel, et que l'on peut alors considérer leur concentration comme pratiquement constante. Cette idée d'annuler l'influence d'un réactif sur une grandeur en introduisant celui-ci en large excès se trouve souvent en Chimie, tant sur les questions d'équilibre que sur les questions de cinétique.

Dans d'autres cas, le problème est tout à fait accessible, soit rigoureusement, soit en utilisant certaines approximations comme vous apprendrez à le faire, notamment pour les équilibres acide/base engageant plusieurs couples simultanément.

Prenons par exemple la réaction vue plus haut, d'équation :



Notons  $K_A$  la constante d'équilibre de cette réaction (valeur connue :  $K_A = 10^{-4,8}$  à 298 K), et  $c_0 = 1,0 \cdot 10^{-2} \text{ mol.L}^{-1}$  la concentration apportée en acide éthanoïque dans cette solution. Notons encore  $x_f$  la valeur finale de son avancement, et  $V$ , le volume de solution dans lequel se déroule la réaction.

On peut aisément montrer qu'à l'état final (lorsque l'équilibre est atteint) :

- $[\text{CH}_3\text{COOH}]_f = c_0 - \frac{x_f}{V}$  ;
- $[\text{CH}_3\text{COO}^-]_f = \frac{x_f}{V}$  ;
- $[\text{H}_3\text{O}^+]_f = \frac{x_f}{V}$ .

La contribution de l'eau, solvant, est égale à 1. La constante d'équilibre s'exprime alors :

$$K_A = \frac{[\text{CH}_3\text{COO}^-]_f [\text{H}_3\text{O}^+]_f}{[\text{CH}_3\text{COOH}]_f} = \frac{\left(\frac{x_f}{V}\right)^2}{c_0 - \frac{x_f}{V}}$$

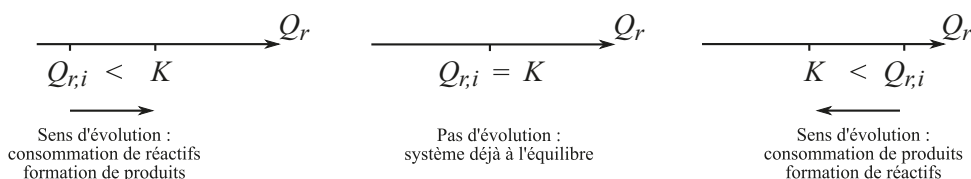
On aboutit alors à une équation du deuxième degré en  $\frac{x_f}{V}$ , dont la résolution nous permet d'obtenir la valeur finale de cette grandeur, et avec elle celles des différentes concentrations.

## — Comment prévoir dans quel sens va se faire l'évolution d'un système chimique placé hors équilibre ?

Nous venons de voir qu'à l'équilibre le quotient de réaction prend toujours une valeur privilégiée notée  $K$  et appelée constante d'équilibre. Cependant, au début d'une expérience, c'est à l'expérimentateur/trice qu'il revient de décider quelles quantités de chaque espèce il/elle met en présence dans le milieu réactionnel. Dans le cas général, la valeur initiale  $Q_{r,i}$  du quotient de réaction diffère de sa valeur finale  $Q_{r,f} = K$ .

Trois cas peuvent alors se présenter :

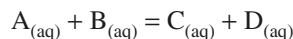
- Si  $Q_{r,i} < K$  : la valeur du quotient de réaction est inférieure à sa valeur d'équilibre. Le système va donc évoluer dans le sens d'un accroissement de celui-ci. Ainsi le numérateur du quotient de réaction va-t-il devoir augmenter, et corrélativement le dénominateur va diminuer (si les produits se forment davantage, c'est fatalement que les réactifs ont été davantage consommés). Autrement dit, la réaction se fait dans le **sens direct**.
- Si  $Q_{r,i} > K$  : la valeur du quotient de réaction est supérieure à sa valeur d'équilibre. Le système va donc cette fois évoluer dans le sens d'une diminution de celui-ci, c'est-à-dire qu'il va consommer les produits et former les réactifs : la réaction se fait dans le **sens indirect**.
- Si  $Q_{r,i} = K$ , le système est déjà à l'équilibre et aucune évolution ne se manifeste.



## — Quand peut-on considérer qu'une réaction est pratiquement totale ?

Lorsqu'une réaction aboutit à un équilibre mais que la valeur de sa constante est particulièrement élevée, il est nécessaire, pour que le quotient de réaction puisse atteindre celle-ci, que l'activité d'au moins l'un des réactifs tende pratiquement vers 0. Ainsi la valeur finale de l'avancement va-t-elle être très proche de sa valeur maximale et la réaction être considérée comme pratiquement totale : elle est alors dite **quantitative**.

On peut intuitivement caractériser éventuellement quantitativement une réaction. Prenons par exemple le cas simple suivant :



On considère ici simplement deux réactifs et deux produits, tous solvatés et affectés de nombres stœchiométriques tous égaux à 1.

Nous supposons en outre que A et B ont été introduits en quantités égales et notées  $n_0$  et les produits initialement absents du système chimique. Un tableau d'avancement nous permet de voir rapidement qu'à l'équilibre :

- $n_{C,f} = n_{D,f} = x_f$  ;
- $n_{A,f} = n_{B,f} = n_0 - x_f$ .

Nous en déduisons l'expression du quotient de réaction à l'équilibre, en notant  $V$  le volume de la solution :

$$Q_r(x = x_f) = \frac{[C]_f[D]_f}{[A]_f[B]_f} = \frac{\frac{x_f}{V} \times \frac{x_f}{V}}{\frac{n_0 - x_f}{V} \times \frac{n_0 - x_f}{V}} \Leftrightarrow Q_{r,f} = \left( \frac{x_f}{n_0 - x_f} \right)^2$$

Comme, à l'équilibre,  $Q_{r,f} = K$ , nous pouvons dire que  $x_f$  répond à l'équation :

$$\left( \frac{x_f}{n_0 - x_f} \right)^2 = K \Leftrightarrow x_f(1 + \sqrt{K}) = n_0\sqrt{K}$$

Comme par ailleurs la valeur maximale autorisée à l'avancement par les quantités initiales de A et B est  $x_{\max} = n_0$ , nous en déduisons le taux d'avancement :

$$\tau = \frac{x_f}{x_{\max}} = \frac{\sqrt{K}}{1 + \sqrt{K}} \Leftrightarrow 1 - \tau = \frac{1}{1 + \sqrt{K}}$$

Donc, pour que  $\tau \simeq 1$ , c'est-à-dire que  $1 - \tau \ll 1$ , il suffit, dans ce cas, que  $1 + \sqrt{K} \gg 1$ , c'est-à-dire que  $\sqrt{K} \gg 1$ . En chimie, on considère en général qu'une valeur domine sur une autre lorsque le quotient de la première par la deuxième est supérieur à 10, voire à 100. Ici, la condition devient donc  $\sqrt{K} \geq 10$ , c'est-à-dire  $K \geq 10^2$  (avec  $\sqrt{K} \geq 100$ , on aurait la condition  $K \geq 10^4$ ).

Il importe de garder à l'esprit que nous ne venons de montrer ce critère que pour un cas très particulier (deux réactifs introduits en proportions stœchiométriques, deux produits, tous solvatisés, avec des nombres stœchiométriques égaux à 1). Il n'est malheureusement pas possible d'établir un critère univoque pour savoir si une réaction est quantitative ou non. En effet, l'expression du quotient de réaction dépend de la stœchiométrie de la réaction envisagée, de l'état des espèces qu'elle engage et de leurs quantités de matières initiales respectives. L'équation  $Q_r(x_f) = K$ , dont on peut déduire la valeur de  $x_f$ , n'aura donc pas toujours la même expression. L'expression de  $x_f$  elle-même n'aura donc pas toujours la même forme, non plus que celle du taux d'avancement.

On considère cependant souvent qu'une réaction dont la constante d'équilibre est supérieure à  $10^4$  est quantitative. Ceci est une approximation convenable dans de nombreux cas, mais il convient de rester prudent. Des nombres stœchiométriques élevés, notamment, peuvent exacerber la contribution de certaines concentrations au quotient de réaction, et mettre ce critère en défaut.

## 3.3 Équilibres acido-basiques

### — Qu'est-ce qu'une réaction acido-basique ?

Une très grande partie des transformations rencontrées dans la nature résultent d'un échange d'ions  $H^+$ , autrement dit de protons (attention : rien de nucléaire, seulement un atome d'hy-

drogène privé d'électron). Leur importance en fait donc des acteurs de premier plan dans l'étude de la chimie.

On appelle alors **réaction acido-basique** toute réaction chimique au cours de laquelle on peut mettre en évidence un échange de proton (ion  $H^+$ ), entre une espèce qui le cède et une espèce qui le capte.

Dans ce cadre, on appelle alors :

- **acide au sens de Brönsted** toute espèce capable de céder un ou plusieurs protons  $H^+$ . Dans le cas où elle ne peut en céder qu'un seul, on parle de monoacide ; dans le cas contraire, de polyacide.

Un acide est donc une **espèce donneuse** de particule (en l'occurrence, de proton). Un proton ainsi capable de quitter son acide d'origine est dit **labile** ;

- **base au sens de Brönsted** toute espèce capable de capter un ou plusieurs protons  $H^+$ . Dans le cas où elle ne peut en capter qu'un seul, on parle de monobase ; dans le cas contraire, de polybase.

Une base est donc une **espèce receveuse** de particule (en l'occurrence, de proton).

## — Qu'est-ce qu'un couple acide/base ?

Un acide qui vient de céder un proton se retrouve avec une place vacante qui lui permet potentiellement de recevoir un proton, ce qui en fait une base au sens de Brönsted. La base ainsi obtenue est appelée **base conjuguée** de l'acide de départ.

Réciproquement, une base qui vient de capter un proton se retrouve avec un proton plus ou moins labile qu'elle peut potentiellement céder, ce qui en fait un acide au sens de Brönsted.

L'acide ainsi obtenu est appelé **acide conjugué** de la base de départ.

On appelle alors **couple acide/base** le couple formé par un acide et sa base conjuguée.

Par convention, l'acide est toujours donné le premier.

## — Qu'est-ce que la force d'un acide ou d'une base ?

La propension à céder un proton varie d'un acide à l'autre. On dit ainsi d'un acide qu'il est plus fort qu'un autre s'il cède son proton plus facilement.

De même, toutes les bases ne présentent pas la même avidité de protons. On dit alors d'une base qu'elle est plus forte qu'une autre si elle capte plus facilement un proton. Nous verrons ultérieurement comment il est possible de quantifier ces facilités.

Il est intéressant de noter que plus un acide est fort, plus sa base conjuguée est faible : la facilité de cession d'un proton entraîne immédiatement la difficulté de sa rétention.

Inversement, plus une base est forte, plus son acide est faible : très avide de protons, la base ne les cédera pas aisément.

Ce concept de force d'un acide ou d'une base ne doit pas être confondu avec ce que l'on appelle acide fort ou base forte dans l'absolu. Ainsi, on appelle **acide fort** (respectivement **base forte**), un acide (respectivement une base) réagissant totalement avec l'eau. La base conjuguée (respectivement l'acide conjugué) obtenu n'a alors de base (respectivement d'acide) que le nom, en ce sens que, conjugué(e) d'une espèce réagissant totalement avec l'eau, il lui est impossible de capter (respectivement de céder) un proton. Une telle espèce est dite **indifférente**.

Les acides forts constituent davantage une exception qu'une règle, et il est bon de retenir les plus courants d'entre eux ; on peut notamment citer : le chlorure d'hydrogène  $HCl$ , l'acide

nitrique  $\text{HNO}_3$ , ainsi que la première acidité de l'acide sulfurique  $\text{H}_2\text{SO}_4$  (ce dernier peut céder ses deux protons par réaction avec l'eau, mais seule la première réaction est totale, tandis que la seconde aboutit à un équilibre).

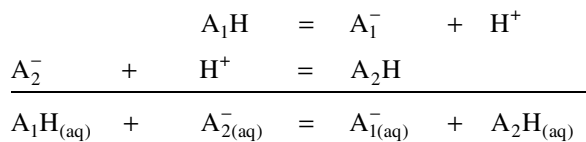
Parmi les bases fortes, on peut citer la soude  $\text{NaOH}$ , la potasse  $\text{KOH}$  ou l'ion éthanolate  $\text{C}_2\text{H}_5\text{O}^-$ .

**Remarque :** la soude et la potasse donnent en fait directement, lors de leur dissolution dans l'eau, des ions hydroxyde  $\text{HO}^-$ . Il n'est pas vraiment possible de dire si ceux-ci réagissent avec l'eau, étant donné qu'alors  $\text{H}_2\text{O}$  donnerait  $\text{HO}^-$  en cédant son proton, tandis que  $\text{HO}^-$  deviendrait  $\text{H}_2\text{O}$  en gagnant un. Reste cependant que l'introduction de ces ions hydroxyde peut être vue comme si une base avait été introduite, et avait capté le proton d'un certain nombre de molécules d'eau.

## — Comment formaliser l'échange de proton ?

Nous venons de distinguer certaines espèces susceptibles de céder un ou plusieurs protons, et d'autres susceptibles d'en capter. Si nous plaçons ainsi l'acide d'un couple et la base d'un autre en présence l'un de l'autre, il est possible que le premier cède un proton à la seconde.

Chacun réagira selon sa demi-équation acido-basique, et la combinaison des deux donnera une équation bilan (traduisant la réalité, donc les états doivent cette fois être précisés) :



On obtient ainsi la base conjuguée  $\text{A}_1^-$  de l'acide de départ  $\text{A}_1\text{H}$ , et l'acide conjugué  $\text{A}_2\text{H}$  de la base de départ  $\text{A}_2^-$ . Ceux-ci peuvent éventuellement réagir également entre eux pour redonner les réactifs de départ : le système chimique est alors le siège d'un équilibre chimique.

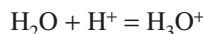
## — L'eau possède-t-elle des propriétés acido-basiques ?

Tout serait pour le plus simple, si le solvant avec lequel nous travaillons n'était lui-même doté de propriétés acido-basiques. En effet, on constate expérimentalement que l'eau est capable :

- de céder un proton : elle constitue donc un acide au sens de Brönsted. Sa base conjuguée est l'ion hydroxyde  $\text{HO}^-$ , avec lequel elle forme le couple  $\text{H}_2\text{O}/\text{HO}^-$  :



- de capter un proton : elle constitue donc également une base au sens de Brönsted. Son acide conjugué est l'ion oxonium  $\text{H}_3\text{O}^+$ , avec lequel elle forme le couple  $\text{H}_3\text{O}^+/\text{H}_2\text{O}$  :



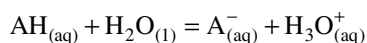
On dit d'une telle espèce qu'elle est **amphotère**, ou bien encore qu'il s'agit d'un **ampholyte**. Cette propriété se retrouve également chez d'autres espèces, par exemple :

- L'ion hydrogénocarbonate, base du couple  $(\text{H}_2\text{O}, \text{CO}_2)/\text{HCO}_3^-$ , et acide du couple  $\text{HCO}_3^-/\text{CO}_3^{2-}$ .
- L'ion hydrogénosulfate, base du couple  $\text{H}_2\text{SO}_4/\text{HSO}_4^-$ , et acide du couple  $\text{HSO}_4^-/\text{SO}_4^{2-}$ .

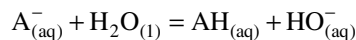
Cette propriété de l'eau va entraîner deux conséquences majeures :

1. Lors de la mise en solution aqueuse d'un acide (respectivement d'une base), l'eau va se comporter comme une base (respectivement un acide), et réagir avec l'espèce introduite :

- Cas de la mise en solution aqueuse d'un acide :

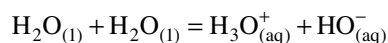


- Cas de la mise en solution aqueuse d'une base :



Nous verrons plus loin que le classement des forces des acides et des bases se fait précisément par rapport à la facilité avec laquelle les uns et les autres réagissent avec l'espèce de référence que constitue l'eau.

2. L'eau peut réagir sur elle-même, suivant la réaction d'équation :



Cette réaction est appelée **autoprotolyse de l'eau**. Nous constatons ainsi que de l'eau, si « pure » soit-elle, contient systématiquement des ions oxonium et hydroxyde issus de cette réaction de dissociation de l'eau.

La constante de cet équilibre est appelée **produit ionique de l'eau** et a pour valeur, à la température  $T = 298 \text{ K}$  :  $K_e = 1,0 \cdot 10^{-14}$ .

On peut déterminer précisément la valeur des concentrations en ions oxonium et hydroxyde. Un tableau d'avancement permet d'écrire, en notant  $V$  le volume du système étudié :

$$[\text{H}_3\text{O}^+] = [\text{HO}^-] = \frac{x}{V}. \text{ Nous en déduisons l'expression du quotient de réaction : } Q_r = \left(\frac{x}{V}\right)^2,$$

avec  $\frac{x}{V}$  en  $\text{mol.L}^{-1}$ . On obtient ainsi, à l'équilibre :  $K_e = \left(\frac{x_f}{V}\right)^2$ , d'où  $\frac{x_f}{V} = \sqrt{K_e} = 1,0 \cdot 10^{-7} \text{ mol.L}^{-1}$ .

On a donc :

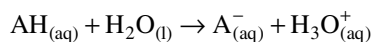
$$[\text{H}_3\text{O}^+]_f = [\text{HO}^-]_f = \frac{x_f}{V} = 1,0 \cdot 10^{-7} \text{ mol.L}^{-1}$$

Ainsi, si l'eau n'est pas à proprement parler pure, les concentrations en ions oxonium et hydroxyde y sont cependant très faibles devant les concentrations usuelles.

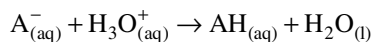
## — Comment savoir quelle forme d'un couple acide/base domine dans une solution, à pH fixé ?

Considérons une solution aqueuse de pH fixé. Lorsqu'un acide se trouve dans une telle solution, l'équilibre est fixé par deux réactions inverses l'une de l'autre :

- La réaction de l'acide avec l'eau :



- La réaction de la base conjuguée ainsi formée avec les ions oxonium (ceux issus de la première réaction, plus ceux éventuellement déjà présents dans le milieu) :

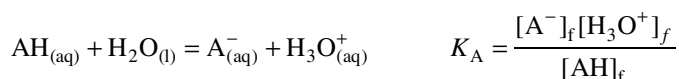


L'eau, en tant que solvant, est toujours présente en quantité quasi illimitée. Les proportions respectives d'acide AH et de sa base conjuguée,  $\text{A}^-$ , sont donc fixées par la concentration en ions oxonium (autrement dit le pH) de la solution.

Ceci correspond à une intuition assez évidente :

- Pour de faibles valeurs du pH (milieu très acide), le milieu est riche en protons. Il est donc difficile à l'acide de céder le sien et l'acide domine sur la base.
- Dans une certaine zone de pH, l'acide et sa base conjuguée coexistent en proportions comparables : le milieu est juste assez riche en protons pour que les réactions directe (acide avec l'eau) et inverse (base conjuguée avec les ions oxonium) aient même vitesse pour des quantités comparables des réactifs qu'elles engagent.
- Pour de grandes valeurs du pH (milieu très basique), le milieu est pauvre en protons. Il est donc très facile à l'acide de céder son proton, et difficile à la base d'en capter un. C'est donc cette fois la base qui domine sur l'acide.

La zone de coexistence peut être déterminée en exprimant la constante de cet équilibre. Nous allons voir qu'elle constitue une grandeur primordiale pour la suite de l'étude, aussi la noterons-nous  $K_A$ .



Nous constatons alors qu'en effet :

- Si  $[\text{H}_3\text{O}^+]_{\text{f}}$  est élevée, le rapport  $[\text{A}^-]_{\text{f}}/[\text{AH}]_{\text{f}}$  est faible : l'acide domine sur la base.
- Si  $[\text{H}_3\text{O}^+]_{\text{f}}$  est faible, le rapport  $[\text{A}^-]_{\text{f}}/[\text{AH}]_{\text{f}}$  est élevé : la base domine sur l'acide.

**Remarque :** notons ici que le produit ionique de l'eau,  $K_e$ , peut être vu comme la constante d'acidité du couple  $\text{H}_2\text{O}/\text{HO}^-$ .

Dans les faits on préfère souvent travailler avec des grandeurs logarithmiques, déduites de celles ci-dessus. Concrètement, on désigne le pendant logarithmique de la constante  $K_A$  sous le nom de  $\text{p}K_A$ , défini comme :

|                                                                   |                    |            |
|-------------------------------------------------------------------|--------------------|------------|
| $\text{p}K_A = -\log K_A \Leftrightarrow K_A = 10^{-\text{p}K_A}$ | $K_A, \text{p}K_A$ | sans unité |
|-------------------------------------------------------------------|--------------------|------------|

On constate alors que, en remplaçant  $K_A$  par son expression en fonction des concentrations des espèces intervenant dans l'équilibre, on obtient :

|                                                                                                                |                                                                                          |                                      |
|----------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------|--------------------------------------|
| $\text{pH}_{\text{f}} = \text{p}K_A + \log\left(\frac{[\text{A}^-]_{\text{f}}}{[\text{AH}]_{\text{f}}}\right)$ | $\text{pH}_{\text{f}}, \text{p}K_A$<br>$[\text{A}^-]_{\text{f}}, [\text{AH}]_{\text{f}}$ | sans unité<br>en $\text{mol.L}^{-1}$ |
|----------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------|--------------------------------------|

Nous constatons ainsi que la limite entre la zone où l'espèce  $\text{A}^-$  est majoritaire et celle où c'est AH qui est majoritaire se situe à  $\text{pH} = \text{p}K_A$  (cas où  $[\text{A}^-]_{\text{f}} = [\text{AH}]_{\text{f}}$ ). Dans les faits on considère qu'une espèce domine significativement sur une autre lorsque sa concentration dans le milieu est au moins 10 fois supérieure à celle de l'autre. On écrira ainsi que :

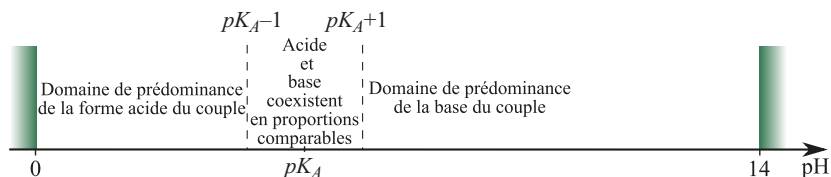
- AH domine sur  $\text{A}^-$  si :

$$\frac{[\text{A}^-]_{\text{f}}}{[\text{AH}]_{\text{f}}} \leq \frac{1}{10} \Leftrightarrow \log\left(\frac{[\text{A}^-]_{\text{f}}}{[\text{AH}]_{\text{f}}}\right) \leq -1 \Leftrightarrow \text{pH} \leq \text{p}K_A - 1$$

- $\text{A}^-$  domine sur AH si :

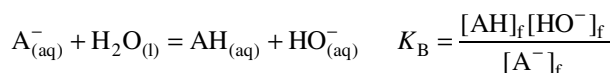
$$\frac{[\text{A}^-]_{\text{f}}}{[\text{AH}]_{\text{f}}} \geq 10 \Leftrightarrow \log\left(\frac{[\text{A}^-]_{\text{f}}}{[\text{AH}]_{\text{f}}}\right) \geq 1 \Leftrightarrow \text{pH} \geq \text{p}K_A + 1$$

Les domaines de pH ainsi délimités sont appelés **domaines d'Henderson**.



**Remarque :** on trouve parfois, dans certains ouvrages, des critères de prédominance plus exigeants. Une espèce n'est alors considérée prédominante sur une autre que si leurs concentrations sont dans un rapport 100, voire 1 000. Les limites de pH à  $pK_A \pm 1$  passent alors à  $pK_A \pm 2$ , voire  $pK_A \pm 3$ .

On utilise également, quoique moins souvent, le  $pK_B$ , relatif à l'équilibre d'une base avec l'eau :



Le passage au logarithme permet alors de montrer que :

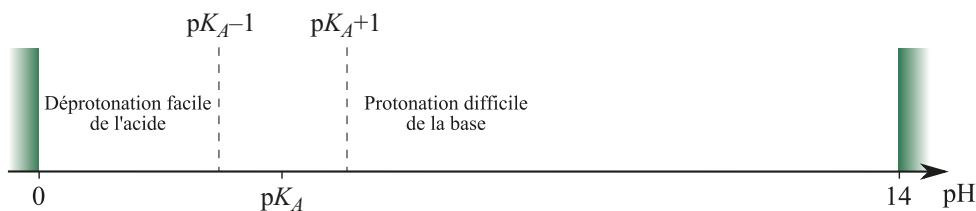
$$pOH = pK_B + \log \left( \frac{[AH]_f}{[A^-]_f} \right)$$

en notant  $pOH = -\log[HO^-]$  et  $pK_B = -\log K_B$ . On peut par ailleurs facilement démontrer que  $pOH = pK_e - pH$  et  $pK_B = pK_e - pK_A$ , avec  $pK_e = -\log K_e$ . On retrouve alors la même relation que précédemment.

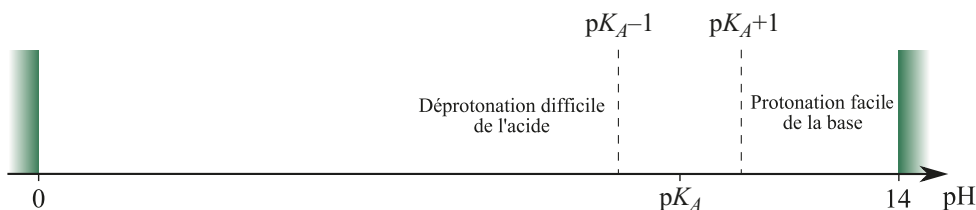
## — Comment évalue-t-on la force d'un acide ?

Nous venons de voir que le  $pK_A$  d'un couple acide/base constitue en fait la valeur limite du pH, où s'inverse le rapport espèce dominante/espèce dominée. Selon que l'acide est très fort (et donc sa base conjuguée très faible) ou l'inverse, le  $pK_A$  n'aura pas la même valeur. En effet :

$pK_A$  peu élevé : acide très peu faible, base très faible



$pK_A$  élevé : acide très faible, base très peu faible



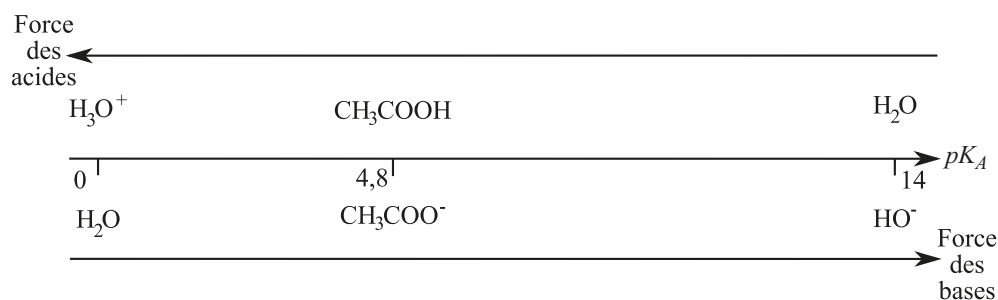
- Si l'acide est très fort, il cède son proton facilement, même dans un milieu où les protons abondent (c'est-à-dire dont le pH est faible). La base, très faible, sera quant à elle très difficile à protoner, même en milieu acide. On en déduit que la valeur du  $pK_A$  est faible.
- Inversement, si l'acide est très faible, il cède son proton difficilement, même dans un milieu où les protons sont rares (c'est-à-dire dont le pH est élevé). La base, très forte, sera quant à elle très facile à protoner, même en milieu basique. On en déduit que le  $pK_A$  a une valeur élevée.

Ainsi, bien que cela puisse paraître paradoxal, un acide est d'autant plus fort que son  $pK_A$  est plus faible, et qu'il lui est donc plus difficile de rester sous forme acide dans une solution. Plus généralement, la force d'une espèce peut être vue comme la facilité de celle-ci à se transformer pour donner son espèce conjuguée.

## — Comment classe-t-on les couples acide/base ?

Nous disposons donc à présent d'un critère permettant de mesurer la force d'un acide : la facilité avec laquelle il réagit avec l'eau pour aboutir à un équilibre, que quantifie sa constante d'acidité. Tous les acides dont la réaction avec l'eau aboutit ainsi à un équilibre sont qualifiés d'**acides faibles**. Nous venons cependant de voir qu'au sein même de ces acides un classement de force pouvait être établi.

On peut alors tracer le diagramme suivant, figurant les couples acide/base, leurs  $pK_A$  et les forces respectives de leurs membres respectifs :



On constate que le domaine des  $pK_A$  est limité :

- dans les faibles valeurs, par la réaction de l'ion oxonium avec l'eau. La constante de cet équilibre vaut bien évidemment 1, d'où un  $pK_A = 0$  ;
- dans les hautes valeurs, par la réaction de l'eau avec l'eau. La constante de cet équilibre vaut  $K_e$ , d'où un  $pK_A = pK_e = 14,0$  à température ambiante.

Ainsi, dans l'eau :

- L'ion oxonium est le plus fort de tous les acides.
- L'eau est le plus faible de tous les acides.
- L'ion hydroxyde est la plus forte de toutes les bases.
- L'eau est la plus faible de toutes les bases.

On constate donc que l'eau (qui n'en est plus à un paradoxe près), solvant essentiel doté de la propriété exceptionnelle d'être à la fois un acide et une base, trouve également le moyen d'être le plus mauvais représentant de chacune de ces deux catégories.

## — Qu'en est-il des acides forts ?

Pourtant nous savons qu'il existe des acides dont la réaction avec l'eau est totale : les **acides forts**. Mais lorsqu'elle a lieu, la libération de leurs protons forme autant d'ions oxonium. Lorsque l'on introduit un acide fort dans l'eau, tout se passe donc comme si l'on avait en fait introduit directement une quantité égale d'ions oxonium. On dit qu'il y a **nivellement des acides forts par l'eau**.

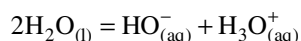
Réciproquement, la mise en solution d'une base forte se soldera par la transformation de toutes les molécules de cette base en leur acide conjugué, avec formation d'autant d'ions hydroxyde. Tout se passe donc comme si l'on avait en fait introduit directement une quantité égale d'ions hydroxyde. On dit qu'il y a **nivellement des bases fortes par l'eau**.

## — Peut-on ne pas mélanger un acide et une base dans une solution aqueuse ?

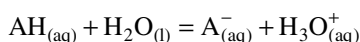
Nous avons vu que l'eau était à la fois un acide et une base. Vouloir ne pas mettre d'acide ou de base dans de l'eau est donc en soi un non-sens. Du reste, la seule situation où une solution aqueuse ne contient qu'un seul acide et une seule base est donc le cas de l'eau dite « pure ». Sitôt que l'on introduit un acide AH dans de l'eau, on a, en toute rigueur, deux acides (AH et H<sub>2</sub>O), ainsi qu'une base (H<sub>2</sub>O).

Chaque acide peut donc *a priori* réagir avec la base présente, et nous nous retrouvons avec deux réactions concurrentes dans la production d'ions oxonium :

- L'autoprotolyse de l'eau (notée APE dans la suite) :



- La réaction de l'autre acide avec l'eau (notée RAE dans la suite) :



Pourtant, vous avez appris au lycée à prendre en compte uniquement la seconde réaction pour déterminer la concentration finale en ions oxonium. Ceci semble, qualitativement au moins, légitime. Nous avons en effet vu que la concentration en ions oxonium issus de la seule APE était de l'ordre de 10<sup>-7</sup> mol.L<sup>-1</sup>. Cette faible valeur résulte de la faible valeur de la constante d'équilibre associée,  $K_e$ .

On peut en fait démontrer qu'en toute rigueur, à l'issue de la mise en solution d'un acide apporté en concentration  $c_A$  et de constante d'acidité  $K_A$ , les concentrations en ions oxonium issus de ces deux réactions peuvent s'exprimer en fonction de la concentration totale en ions oxonium (nous avons omis l'indice f pour alléger des écritures déjà chargées mais toutes les concentrations sont comme toujours considérées à l'état final, c'est-à-dire à l'équilibre) :

- Concentration des ions oxonium issus de la seule autoprotolyse de l'eau :

$$K_e = [\text{HO}^-]_{\text{tot}} [\text{H}_3\text{O}^+]_{\text{tot}} = [\text{HO}^-]_{\text{APE}} [\text{H}_3\text{O}^+]_{\text{tot}} = [\text{H}_3\text{O}^+]_{\text{APE}} [\text{H}_3\text{O}^+]_{\text{tot}}$$

$$\Leftrightarrow [\text{H}_3\text{O}^+]_{\text{APE}} = \frac{K_e}{[\text{H}_3\text{O}^+]_{\text{tot}}}$$

puisque les seuls ions hydroxyde présents sont ceux issus de l'APE, qui du reste produit autant d'ions hydroxyde que d'ions oxonium.

- Concentration des ions oxonium issus de la seule réaction de l'acide AH avec l'eau :

$$K_A = \frac{[A^-][H_3O^+]_{\text{tot}}}{[AH]} = \frac{[A^-][H_3O^+]_{\text{tot}}}{c_A - [A^-]} = \frac{[H_3O^+]_{\text{RAE}}[H_3O^+]_{\text{tot}}}{c_A - [H_3O^+]_{\text{RAE}}}$$

$$\Leftrightarrow [H_3O^+]_{\text{RAE}} = \frac{K_A c_A}{K_A + [H_3O^+]_{\text{tot}}}$$

puisque la réaction de l'acide avec l'eau produit autant de molécules de base  $A^-$  que d'ions oxonium.

Or la concentration totale en ions oxonium résulte de la somme de ces deux concentrations, d'où :

$$[H_3O^+]_{\text{tot}} = \frac{K_e}{[H_3O^+]_{\text{tot}}} + \frac{K_A c_A}{K_A + [H_3O^+]_{\text{tot}}}$$

Cette équation se développe en polynôme du troisième degré, qu'il n'est pas question de résoudre, mais sur laquelle nous pouvons mener certaines approximations.

En particulier, pour pouvoir négliger la contribution de l'autoprotolyse de l'eau à la concentration totale en ions oxonium, il suffit que  $[H_3O^+]_{\text{APE}} \ll [H_3O^+]_{\text{tot}}$ , c'est-à-dire que :

$$\frac{K_e}{[H_3O^+]_{\text{tot}}} \ll [H_3O^+]_{\text{tot}} \Leftrightarrow [H_3O^+]_{\text{tot}} > \sqrt{\frac{K_e}{10}}$$

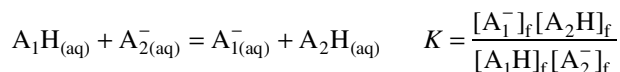
en considérant qu'un terme domine sur un autre s'il lui est supérieur au moins d'un facteur 10. L'équation ci-dessus réclame donc en fait que  $\text{pH} < 6,5$  à 298 K. Ainsi, pour pouvoir négliger la quantité d'ions oxonium produite par l'autoprotolyse de l'eau devant celle produite par l'autre réaction, il suffit que le pH à l'équilibre adopte une valeur inférieure à 6,5.

La démarche ci-dessus est en fait très générale et permet de déterminer, lorsque plusieurs réactions sont en concurrence pour la production d'une espèce (en l'occurrence, les ions oxonium), dans quelle mesure l'une sera plus prodigue qu'une autre et dominera donc dans la production de cette espèce. On peut alors mener les calculs dans l'hypothèse où cette réaction est seule à œuvrer, ce qui les simplifie considérablement. Les résultats obtenus de cette manière sont souvent tributaires d'approximations, toujours légitimées *a posteriori* cependant.

La réaction dominante est appelée **réaction prépondérante**. Nous allons voir à présent comment la déterminer.

## — Comment calculer la constante d'équilibre d'une réaction acido-basique ?

Considérons la réaction de l'acide  $A_1H$  d'un couple, avec la base  $A_2^-$  d'un autre, de  $\text{p}K_A$  respectifs  $\text{p}K_{A,1}$  et  $\text{p}K_{A,2}$ . Écrivons son équation, ainsi que la constante de l'équilibre auquel elle aboutit :



Nous pouvons réécrire cette constante en multipliant le numérateur et le dénominateur par  $[\text{H}_3\text{O}^+]_f$ , d'où il ressort que :

$$K = \frac{[\text{A}_1^-]_f [\text{H}_3\text{O}^+]_f}{[\text{A}_1\text{H}]_f} \times \frac{[\text{A}_2\text{H}]_f}{[\text{A}_2^-]_f [\text{H}_3\text{O}^+]_f} = \frac{K_{\text{A},1}}{K_{\text{A},2}} = 10^{\text{p}K_{\text{A},2} - \text{p}K_{\text{A},1}}$$

## — Comment savoir quelle réaction fixe prioritairement les concentrations à l'équilibre des espèces qu'elle engage ?

Lorsque plusieurs acides et plusieurs bases (y compris l'eau) sont introduits dans un même milieu, plusieurs réactions sont susceptibles de se dérouler (chaque acide peut réagir avec chaque base).

Si une espèce est engagée dans plusieurs de ces équilibres, la règle est simple : **sa concentration est fixée par celle des réactions où elle se trouve engagée, qui possède la constante d'équilibre la plus élevée**. Cette réaction est appelée **réaction prépondérante**. Dans le cas où les réactions envisagées sont de type acido-basique, ceci signifie que le premier équilibre à être fixé est celui pour lequel  $K = 10^{\text{p}K_{\text{A},2} - \text{p}K_{\text{A},1}}$  possède la plus grande valeur. Or nous constatons que la valeur de  $K$  sera d'autant plus importante que :

- la valeur de  $\text{p}K_{\text{A},2}$  est plus élevée, c'est-à-dire que la base engagée est plus forte ;
- la valeur de  $\text{p}K_{\text{A},1}$  est plus faible, c'est-à-dire que l'acide engagé est plus fort.

Ainsi, si plusieurs réactions acido-basiques sont possibles, la réaction qui fixe prioritairement les concentrations des espèces qu'elle engage est celle de **l'acide le plus fort sur la base la plus forte** (idée qui du reste semble assez logique).

Exemples :

1. Dans le cas où seule l'eau est présente, un seul acide et une seule base sont présents (les plus faibles, puisque l'eau est le plus faible des acides et la plus faible des bases). L'autoprotolyse de l'eau est donc la seule réaction possible, de constante d'équilibre égale au produit ionique de l'eau :  $10^{-\text{p}K_e}$ , avec  $\text{p}K_e = -\log K_e = 14$  à 298 K.
2. Dans le cas où sont présents l'eau et un acide AH apporté, on peut envisager deux réactions :
  - L'autoprotolyse de l'eau, de constante d'équilibre égale au produit ionique de l'eau :  $10^{-\text{p}K_e}$ .
  - La réaction de l'acide AH avec l'eau, de constante d'équilibre  $K_A = 10^{-\text{p}K_A}$ . Le  $\text{p}K_A$  étant inférieur à  $\text{p}K_e$ ,  $-\text{p}K_A > -\text{p}K_e$  et nous constatons que cette réaction prime forcément sur l'autoprotolyse de l'eau.
3. Dans le cas où sont présentes l'eau et une base  $\text{A}^-$  apportée, on peut envisager deux réactions :
  - L'autoprotolyse de l'eau, de constante d'équilibre égale au produit ionique de l'eau :  $10^{-\text{p}K_e}$ .
  - La réaction de la base  $\text{A}^-$  avec l'eau, de constante d'équilibre  $K_B = 10^{-\text{p}K_B} = 10^{\text{p}K_A - \text{p}K_e}$ . Le  $\text{p}K_A$  étant supérieur à 0,  $\text{p}K_A - \text{p}K_e > -\text{p}K_e$  et nous constatons que cette réaction prime forcément sur l'autoprotolyse de l'eau.

4. Dans le cas où sont présents l'eau, un acide  $A_1H$  et une base  $A_2^-$  apportés, on peut envisager quatre réactions :

- L'autoprotolyse de l'eau, de constante d'équilibre égale au produit ionique de l'eau :  $10^{-pK_e}$ .
- La réaction de l'acide  $A_1H$  avec l'eau, de constante d'équilibre  $K_{A,1} = 10^{-pK_{A,1}}$ .
- La réaction de la base  $A_2^-$  avec l'eau, de constante d'équilibre  $K_{B,2} = 10^{-pK_{B,2}} = 10^{pK_{A,2} - pK_e}$ .
- La réaction de l'acide  $A_1H$  avec la base  $A_2^-$ , de constante d'équilibre  $K = 10^{pK_{A,2} - pK_{A,1}}$ .

Tous les  $pK_A$  étant compris entre 0 et  $pK_e$ , nous constatons que cette dernière réaction prime forcément sur les trois autres.

Une fois que ce premier état d'équilibre a fixé les concentrations des différentes espèces qu'il engage (donc celles de  $A_1H$ ,  $A_1^-$ ,  $A_2H$  et  $A_2^-$ ), on peut déterminer les concentrations des autres espèces engagées dans les autres équilibres. Sur le fond, cette façon de faire consiste en réalité à considérer que les modifications apportées à cette valeur par une réaction de constante d'équilibre significativement inférieure, ne modifieront que très faiblement la valeur déjà fixée.

**Remarque :** dans les faits, on peut montrer que les divers équilibres ne peuvent être considérés comme indépendants les uns des autres que si les constantes qui leur sont associées diffèrent d'au moins deux ordres de grandeur. Dans le cas contraire, on doit considérer ces équilibres simultanément, ce qui complique significativement le problème et ne sera pas traité ici.

## 3.4 Équilibres d'oxydoréduction

### — Qu'est-ce qu'une réaction d'oxydoréduction ?

Une autre grande catégorie de transformations rencontrées dans la nature sont celles engageant des échanges d'électrons. Contrairement au cas des protons, qui étaient des ions (donc des entités monoatomiques à part entière), les particules échangées cette fois sont donc bien des « morceaux d'atomes », extraits d'un cortège électronique ou intégrant celui-ci.

On appelle ainsi **réaction d'oxydoréduction** toute réaction chimique au cours de laquelle on peut mettre en évidence un échange d'électrons, entre une espèce qui les cède et une espèce qui les capte.

On appelle alors :

- **réducteur** une espèce capable de céder un ou plusieurs électrons. Un réducteur est donc une **espèce donneuse** de particule (en l'occurrence, d'électron) ;
- **oxydant** une espèce capable de capter un ou plusieurs électrons. Un oxydant est donc une **espèce receveuse** de particule (en l'occurrence, d'électron).

### — Qu'est-ce qu'un couple oxydant/réducteur ?

Un oxydant qui vient de capter un ou plusieurs électrons se retrouve avec un ou plusieurs électron(s) qu'il peut *a priori* céder, ce qui en fait potentiellement un réducteur. Le réducteur ainsi obtenu est appelé **réducteur conjugué** de l'oxydant de départ.

Réciproquement, un réducteur qui vient de céder un ou plusieurs électron(s) se retrouve avec une lacune électronique, ce qui en fait un oxydant potentiel. L'oxydant ainsi obtenu est appelé **oxydant conjugué** du réducteur de départ.

On appelle alors **couple oxydant/réducteur** le couple formé par un oxydant et son réducteur conjugué. Par convention, l'oxydant est toujours donné le premier.

**Remarque :** dans le cas des couples acide/base, l'espèce donneuse (acide) est placée en premier, la receveuse (base) en second. Ici, pour un couple oxydant/réducteur, c'est l'espèce receveuse (oxydant) qui est placée en premier : ceci provient du fait qu'historiquement on a d'abord cru qu'il s'agissait d'un transfert de charge positive, qui dans ce cas aurait donc été cédée par l'oxydant et captée par le réducteur.

## — Comment formaliser l'échange d'électrons entre un oxydant et un réducteur simples ?

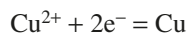
Si nous plaçons le réducteur d'un couple et l'oxydant d'un autre en présence l'un de l'autre, il est possible que le premier cède des électrons au second.

De la même manière que pour les réactions acido-basiques, on commence par établir les demi-équations des couples concernés. Ces demi-équations, appelées **demi-équations électroniques**, mettent en lumière l'échange d'électrons sous la forme :



où Ox et Red représentent respectivement un oxydant et son réducteur conjugué, tandis que  $n$  représente le nombre d'électrons (ici symbolisés par l'écriture  $\text{e}^-$ ) échangés.

La demi-équation électronique du couple  $\text{Cu}^{2+}/\text{Cu}$  s'écrit ainsi :



**Remarque :** dans le cas des réactions acido-basiques, les demi-équations engageaient à chaque fois un unique proton. Si un acide peut en céder (ou une base, en capter) plusieurs, on parle de polyacide (ou de polybase), et l'on considère plusieurs réactions successives, du fait que chacune des espèces intermédiaires peut exister dans l'eau pour peu que le pH y soit favorable. Les oxydants et réducteurs, en revanche, ne peuvent pas toujours gagner ou perdre un seul électron à la fois. Dans l'exemple du cuivre traité ci-dessus, on pourrait penser à traiter successivement les couples (fictifs)  $\text{Cu}^{2+}/\text{Cu}^+$  et  $\text{Cu}^+/\text{Cu}$ , dont les membres échangeraient à chaque fois un seul électron. Mais l'ion  $\text{Cu}^+$  est instable en solution aqueuse quelles que soient les conditions. On considère donc directement le couple  $\text{Cu}^{2+}/\text{Cu}$ , qui engage deux électrons.

## — Qu'en est-il pour les couples plus élaborés ?

Parmi les couples oxydant/réducteur, on trouve beaucoup de couples du type cation métallique/métal solide ( $\text{Ag}^+/\text{Ag}$ ,  $\text{Fe}^{2+}/\text{Fe}$ , ...). Dans ce cas, l'établissement de la demi-équation électronique se fait comme pour le couple  $\text{Cu}^{2+}/\text{Cu}$ , en ajoutant le nombre d'électrons manquant à l'oxydant pour retrouver le réducteur.

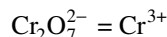
Nous allons cependant voir que les réactions d'oxydoréduction ont une portée beaucoup plus large. On trouve ainsi par exemple le couple ion dichromate/ion chrome (III),  $\text{Cr}_2\text{O}_7^{2-}/\text{Cr}^{3+}$ , pour lequel nous constatons qu'un simple ajout d'électrons ne suffira pas à obtenir une demi-équation équilibrée, du fait notamment de la différence entre les nombres d'atomes de chrome, et de la présence des atomes d'oxygène.

**Remarque :** on observe sur ce couple que l'oxydant est un anion, et le réducteur un cation. On retiendra donc qu'un oxydant, bien qu'avidé d'électrons, peut présenter une charge négative et qu'un réducteur, bien qu'apte à céder des électrons, peut être porteur d'une charge positive. Ceci provient du fait qu'au sein de l'ion dichromate, les atomes de chrome sont liés à de nombreux atomes d'oxygène (7 au total). Ceux-ci, très électronégatifs, accaparent donc l'essentiel des électrons de valence des atomes de chrome. Formellement, tout se passe alors comme si l'ion dichromate comportait en fait deux ions  $\text{Cr}^{6+}$ , entourés d'ions  $\text{O}^{2-}$ .

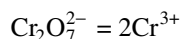
Il est donc nécessaire d'envisager une méthode d'équilibrage des demi-équations électroniques plus générale. Voici comment l'on procède pour une réaction se déroulant en milieu aqueux (présence d'eau) et acide (présence d'ions  $\text{H}^+$ ) (attention : les étapes doivent être menées dans l'ordre) :

1. On commence par équilibrer les éléments autres que l'oxygène O et l'hydrogène H.
2. On équilibre ensuite l'élément oxygène O en ajoutant des molécules d'eau  $\text{H}_2\text{O}$ .
3. On équilibre ensuite l'élément hydrogène en ajoutant des ions  $\text{H}^+$ .
4. On équilibre enfin les charges électriques, en ajoutant des électrons  $\text{e}^-$ .

**Exemple :** Dans le cas du couple  $\text{Cr}_2\text{O}_7^{2-}/\text{Cr}^{3+}$ , nous partons ainsi avec la demi-équation :



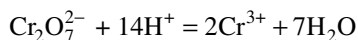
1. Le seul élément autre que O et H est le chrome Cr. On constate que deux atomes Cr sont présents à gauche, pour un seul à droite. On commence par équilibrer cet élément en ajustant les nombres stœchiométriques, et la demi-équation devient :



2. L'élément O est présent 7 fois à gauche, et pas une seule à droite. On ajoute donc autant de molécules d'eau que nécessaire pour équilibrer cet élément. À raison d'un atome d'oxygène par molécule d'eau, la demi-équation devient ainsi :



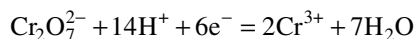
3. L'élément H était absent à gauche comme à droite au départ. Mais l'ajout des molécules d'eau fait que nous nous trouvons finalement avec 14 atomes H à droite, et pas un seul à gauche. On ajoute donc autant d'ions  $\text{H}^+$  que nécessaire pour équilibrer cet élément, et la demi-équation devient :



4. Nous dénombrons alors la charge globale :

- À gauche :  $2\ominus$  pour l'ion dichromate,  $14\oplus$  pour les ions  $\text{H}^+$ , soit un total de  $2\ominus + 14\oplus = 12\oplus$ .
- À droite :  $3\oplus$  pour chaque ion chrome (III), seule espèce chargée, mais présente deux fois, d'où un total de  $2 \times 3\oplus = 6\oplus$ .

Le membre de gauche présente donc un excès de  $12 - 6 = 6$  charges  $\oplus$ , que l'on équilibre en y ajoutant 6 électrons. La demi-équation devient alors :



**Remarque :** comme pour les demi-équations acido-basiques, il ne s'agit que d'une écriture formelle. En effet les électrons, à l'instar des protons, ne peuvent exister à l'état solvaté dans l'eau. Il n'est donc pas question d'envisager l'existence d'électrons isolés dans le milieu réactionnel : ceux-ci ne peuvent provenir que d'un contact direct de l'ion dichromate avec un réducteur.

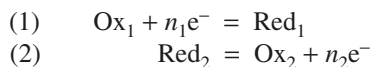
## — Réaction entre un oxydant et un réducteur

Étant désormais en mesure de connaître la demi-équation électronique de n'importe quel couple oxydant/réducteur, voyons ce qu'il se passe concrètement lorsque l'on met en présence l'oxydant d'un couple et le réducteur d'un autre.

On commence par écrire les demi-équations électroniques respectives des couples en plaçant à chaque fois à gauche, celui des membres du couple qui intervient en tant que réactif :

- Celle du couple dont l'oxydant est engagé en tant que réactif est écrite avec l'oxydant à gauche (donc les électrons à gauche).
- Celle du couple dont le réducteur est engagé en tant que réactif est écrite avec le réducteur à gauche (donc les électrons à droite).

On se retrouve ainsi, en faisant abstraction des éventuels termes  $\text{H}_2\text{O}$  et  $\text{H}^+$ , avec une représentation du type suivant :

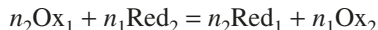


Si chacune de ces demi-équations engage le même nombre d'électrons (c'est-à-dire si  $n_1 = n_2$ ), il suffit de les sommer membre à membre. Dans le cas contraire, nous retrouvons un problème déjà rencontré avec les couples acide/base, à savoir que la/les particule(s) échangée(s) (ici, les électrons) ne peu(ven)t exister à l'état solvaté dans l'eau. En conséquence, la réaction pourra se faire uniquement si les deux demi-équations électroniques sont engagées dans des proportions telles, que le nombre total d'électrons cédés par le réducteur soit égal au nombre total d'électrons acquis par l'oxydant.

Pour que ce soit le cas, il suffit de considérer que la première demi-équation se déroule  $n_2$  fois, quand la deuxième se déroule  $n_1$  fois. Nous avons alors en effet :

$$\begin{aligned}n_2 \times (1) \Rightarrow & n_2 \text{Ox}_1 + n_2 n_1 e^- = n_2 \text{Red}_1 \\ n_1 \times (2) \Rightarrow & n_1 \text{Red}_2 = n_1 \text{Ox}_2 + n_1 n_2 e^-\end{aligned}$$

Chaque couple engage ainsi le même nombre  $n_1 n_2$  d'électrons. On peut alors sommer membre à membre, et l'on obtient l'équation :



**Remarque :** il n'est pas toujours indispensable de multiplier chaque demi-équation par le nombre d'électrons engagés par l'autre. Si par exemple l'une engage 6 électrons et l'autre, 3, il suffit de multiplier la deuxième par 2.

Une fois obtenue cette équation, on peut encore lui apporter quelques finitions :

1. Simplifier les espèces redondantes à gauche et à droite (notamment  $\text{H}_2\text{O}$  et  $\text{H}^+$ ).
2. Réduire les coefficients portant sur les différentes espèces (rappelons que les nombres stœchiométriques doivent être les **plus petits entiers** permettant d'équilibrer l'équation ; s'ils valent par exemple 4, 2, 6 et 8, on simplifiera en 2, 1, 3 et 4).
3. Quoique l'écriture  $\text{H}^+$  soit en général tolérée dans le cadre de l'oxydoréduction, on demande parfois de transformer les ions  $\text{H}^+$  en ions  $\text{H}_3\text{O}^+$  (rappelons qu'en toute rigueur les protons ne peuvent exister seuls dans l'eau). Pour ce faire, on ajoute de part et d'autre de l'équation autant de molécules d'eau qu'il reste d'ions  $\text{H}^+$ .
4. Cette équation reflète une réalité ; on doit donc préciser les états des espèces qu'elle engage.
5. Vérifier, à l'issue de toutes manipulations, que l'équation est toujours bien équilibrée.

Si l'équation fait figurer des ions oxonium en tant que réactif, le déroulement de la réaction nécessite un apport en ions  $\text{H}^+$  et l'expérience devra donc être menée en milieu acide. Cette acidification est souvent obtenue au moyen d'acide sulfurique concentré.

Une équation d'oxydoréduction peut également être équilibrée en milieu basique. Pour ce faire, on commence par l'établir en milieu acide, puis l'on ajoute de part et d'autre autant d'ions hydroxyde  $\text{HO}^-$  qu'il reste d'ions  $\text{H}^+$ . Du côté où ne figure aucun ion  $\text{H}^+$ , les ions hydroxyde restent en l'état. Du côté où figurent les ions  $\text{H}^+$ , ceux-ci réagissent avec les ions hydroxyde pour former autant de molécules d'eau. On aura alors soin de vérifier une possible nouvelle redondance des molécules d'eau entre gauche et droite de l'équation.

## — Les réactions d'oxydoréduction aboutissent-elles à des équilibres chimiques ?

Nous avons vu que, dans le cas des couples acide/base, la réaction aboutit souvent à une situation d'équilibre où l'acide et la base de chaque couple coexistent en proportions comparables. Cependant pour des différences de  $pK_A$  suffisamment importantes (supérieure à 4 unités de pH, typiquement), la réaction peut être considérée comme totale ou bien au contraire comme ne se produisant pratiquement pas.

Les réactions d'oxydoréduction aboutissent elles aussi, en toute rigueur, à des situations d'équilibre. Mais les valeurs des constantes  $K$  associées à ces équilibres sont la plupart du temps soit très importantes (entre  $10^4$  et  $10^{50}$ ), soit au contraire très faibles (réactions inverses des précédentes, donc  $K$  compris entre  $10^{-50}$  et  $10^{-4}$ ). Les équilibres en question sont donc soit :

- extrêmement déplacés vers les produits :  $K$  très élevée, réaction quasi totale ;
- extrêmement déplacés vers les réactifs :  $K$  très faible, réaction quasi ineffective.

On n'aura donc pas, dans un premier temps, à se poser la question de l'existence d'un éventuel équilibre : les réactions d'oxydoréduction pourront être considérées soit comme totales, soit comme inexistantes. Dans le premier cas, la méthodologie du réactif limitant pourra donc être appliquée et on pourra considérer que  $x_f \approx x_{\max}$ .

Si une réaction d'oxydoréduction se fait bien entre le réducteur d'un couple 1 et l'oxydant d'un couple 2, alors c'est que la constante  $K$  associée a une valeur très élevée. Donc la constante de la réaction inverse (entre l'oxydant du couple 1 et le réducteur du couple 2)

a pour valeur  $K' = \frac{1}{K}$ , a autrement dit cette réaction inverse ne se fera pratiquement pas.

Ainsi, l'immersion d'un morceau de cuivre métallique dans une solution contenant des ions argent (I) donne naissance à des dendrites d'argent métallique (arbre de Diane) et provoque le bleuissement de la solution (formation d'ions cuivre (II)). Mais l'immersion d'un fil d'argent dans une solution contenant des ions cuivre (II) ne laisse apparaître aucune évolution significative.

Notons qu'il existe un classement des couples d'oxydoréduction, par une grandeur appelée **potentiel standard** et notée  $E^0$ . Cette grandeur, qui se mesure en volts (V), est en quelque sorte le pendant du  $pK_A$  des couples acide/base pour les couples oxydant/réducteur. Sa valeur, en général de quelques volts, est d'autant plus élevée que l'oxydant du couple est plus fort.

**Remarque :** de même qu'il existe des indicateurs colorés acido-basiques, il existe des indicateurs colorés d'oxydoréduction (ou « *redox* »). Il s'agit d'oxydants ou de réducteurs dont l'espèce conjuguée possède une couleur différente de la leur en solution aqueuse. Selon la valeur du potentiel électrique de la solution dans laquelle on les introduit, c'est l'une ou l'autre de ces deux formes qui prédomine. Les dosages potentiométriques, qui fonctionnent sur le même principe que les dosages acido-basiques (et donnent d'ailleurs à cette occasion des courbes de dosage  $E = f(V_{\text{titrant}})$  très similaires), les utilisent pour repérer l'équivalence, au niveau de laquelle le potentiel varie brusquement. Peut-être aviez-vous déjà remarqué, sur les pH-mètres utilisés en terminale, un bouton que l'on pouvait commuter au choix sur « pH » ou sur « mV » ; en voici, quoique sommaire, l'explication.

## — Pourquoi les réactions d'oxydoréduction présentent-elles un intérêt particulier dans le domaine électrique ?

Outre leur intérêt sur le seul plan de la connaissance, les réactions d'oxydoréduction présentent de nombreuses applications dans le domaine électrique. Elles mettent en effet en

jeu un transfert d'électrons, autrement dit un **courant électrique**. Or qui dit circulation spontanée de courant électrique dit possibilité de concevoir un générateur de courant et, à partir de là, de faire fonctionner un appareil électrique.

Le problème étant que ce transfert, comme nous l'avons vu, se fait par contact direct entre le réducteur et l'oxydant, sans possibilité de canaliser le courant engagé.

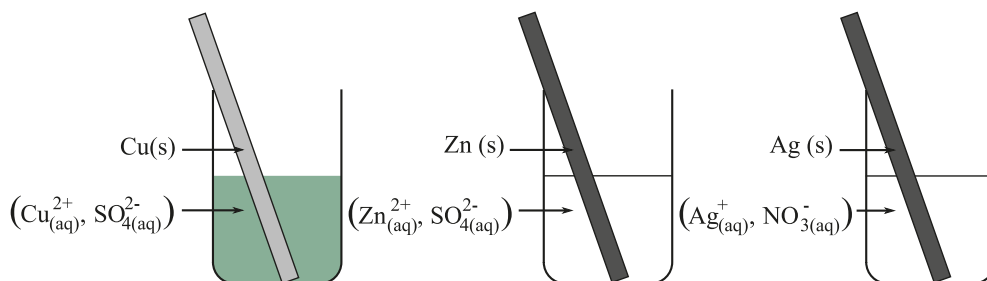
## — Comment peut-on récupérer l'énergie mise en jeu au cours d'une réaction d'oxydoréduction ?

Posons-nous la question suivante : que se passerait-il si nous mettions le réducteur d'un côté d'un fil conducteur, et l'oxydant à l'autre extrémité de ce même fil ? Le réducteur pourrait-il spontanément libérer des électrons dans ce fil, qui circuleraient alors jusqu'à l'oxydant ? On observerait alors l'oxydation du réducteur et la réduction de l'oxydant, chacun demeurant à une extrémité du fil, tandis que le fil deviendrait le siège d'un courant d'électrons ; une sorte de « relation longue distance » d'oxydoréduction.

C'est en effet ce que l'on observe et, en plaçant un dipôle électrique sur le chemin emprunté par le courant, nous allons être en mesure de profiter de l'énergie qu'il véhicule.

## — Qu'est-ce qu'une pile ?

Avant de définir ce qu'est une **pile**, nous devons d'abord définir ce qu'est une **demi-pile**. On appelle ainsi le système résultant de la mise en contact d'un réducteur, avec son oxydant conjugué. Un morceau de cuivre trempant dans une solution aqueuse de sulfate de cuivre, par exemple, est une demi-pile mettant en jeu le couple  $\text{Cu}^{2+}/\text{Cu}$ . De même pour un morceau de zinc trempant dans une solution aqueuse de sulfate de zinc, ou un fil d'argent trempant dans une solution aqueuse de nitrate d'argent.



Ces exemples sont relativement restreints, en ce qu'ils présentent uniquement des couples dont l'un des membres est un solide, ce qui permet de le connecter au circuit extérieur qui doit canaliser le courant électrique. Dans les faits, on peut contourner ce problème : on met l'oxydant et son réducteur conjugué en contact dans la solution au niveau d'une électrode de platine, chimiquement très difficile à attaquer, mais conduisant très bien le courant électrique. On peut alors observer un gain d'électrons par l'oxydant aux dépens de l'électrode de platine lorsqu'il se trouve au contact de la surface de celle-ci (cas d'une demi-pile siège d'une réduction), ou une cession d'électrons au profit de cette même électrode (cas d'une demi-pile siège d'une oxydation). La demi-pile sera le siège d'une oxydation (respectivement d'une réduction) uniquement si elle est mise en contact avec une autre demi-pile, qui sera alors siège d'une réduction (respectivement d'une oxydation).

**Remarque :** vous apprendrez en fait, en CPGE, à définir 3 types d'électrodes (= demi-piles), dites de 1<sup>re</sup>, 2<sup>e</sup> et 3<sup>e</sup> espèces. Celles que vous venez de voir constituent respectivement les électrodes de 1<sup>re</sup>

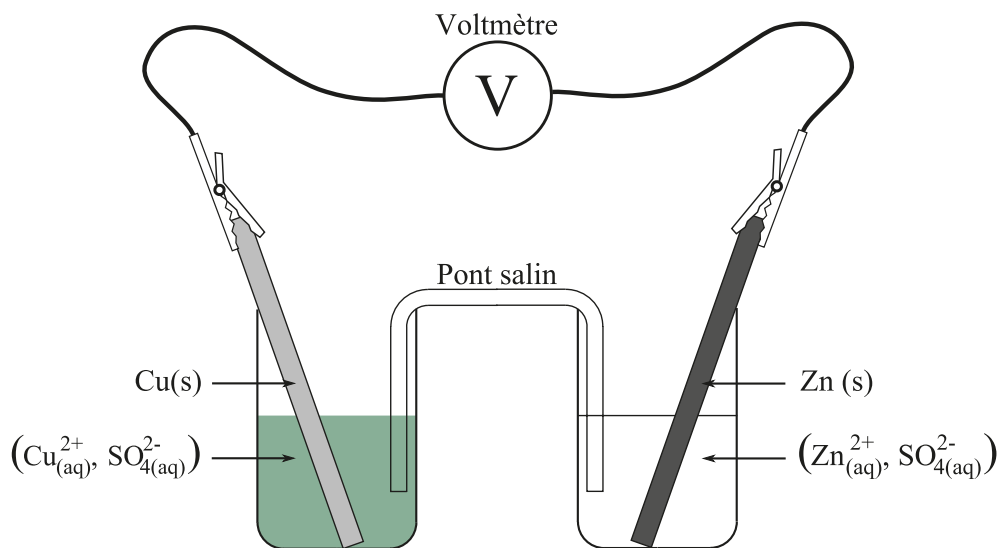
(métal plongeant dans une solution de ses propres ions) et 3<sup>e</sup> espèces (oxydant et réducteur conjugués au contact d'une électrode de platine).

La fabrication d'une pile est alors réalisée par l'association de deux demi-piles à l'aide d'un fil métallique. Ce dispositif doit cependant encore être amélioré sur deux points :

- Un simple fil n'opposant pratiquement aucune résistance au passage du courant électrique, il court-circuiterait la pile. Dans les faits, on place donc un dipôle électrique (par exemple un conducteur ohmique, dont la résistance modère l'importance du courant débité) ou un voltmètre, qui permettra de mesurer la différence de potentiel entre les deux piles (pratiquement sans aucun passage de courant cette fois).
- Pour que le courant électrique puisse circuler, le circuit doit être fermé. En effet, dans le cas contraire, des électrons quitteraient l'une des demi-piles pour aller dans l'autre, ce qui romprait la neutralité électrique dans chacune. On supplée à ce problème en installant entre les deux solutions un **pont salin**. Il peut être constitué par une bandelette de papier imbibé d'une solution ionique, qui permet la circulation des ions d'une solution à l'autre et ferme ainsi le circuit.

**Remarque :** on doit veiller à ce que les ions contenus dans le pont salin soient indifférents aux espèces engagées par les couples des deux demi-piles. En effet, les cations du pont salin vont véritablement migrer dans la demi-pile où se déroule la réduction (pour compenser l'afflux d'électrons par le circuit extérieur), tandis que les anions vont migrer dans la demi-pile où a lieu l'oxydation (pour compenser le départ des électrons). On évitera par exemple un pont salin comportant des ions chlorure dans une pile engageant des ions argent. Par ailleurs, dans le cas de ponts salins réutilisables, il est nécessaire de les recharger périodiquement en ions (par immersion prolongée dans une solution ionique dont ils vont absorber les ions par diffusion).

Un exemple classique est la **pile Daniell**, qui repose sur les couples  $\text{Cu}^{2+}/\text{Cu}$  et  $\text{Zn}^{2+}/\text{Zn}$  :

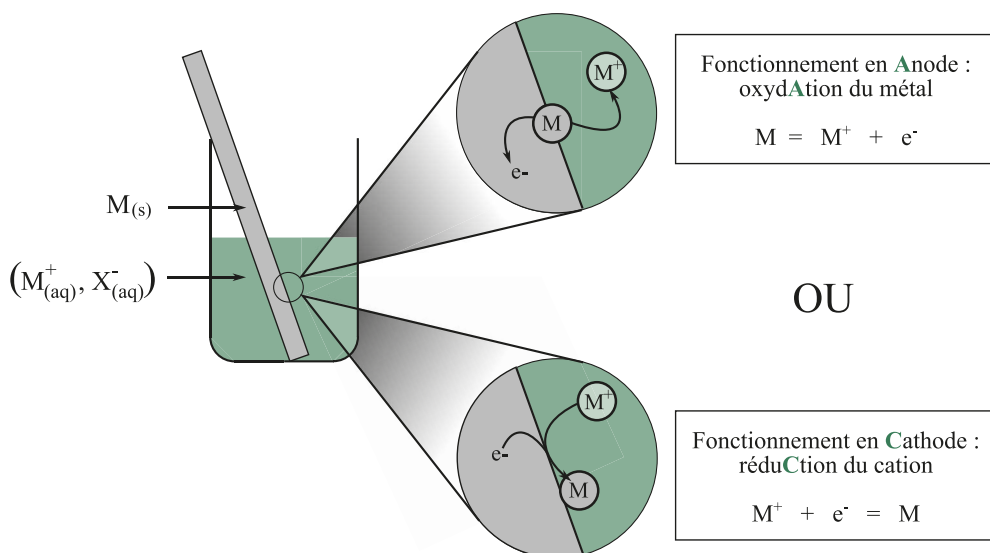


Les pièces métalliques constituant les demi-piles sont appelées **électrodes**. On les différencie par le type de réaction qui se déroule à leur niveau. On appelle ainsi :

- **anode**, l'électrode où se produit l'**oxydation** ;
- **cathode**, l'électrode où se produit la **réduction**.

**Remarque :** on peut trouver d'autres définitions, mais celle-ci présente le mérite d'être systématiquement exacte, que la pile soit en train de débiter, ou bien au contraire en situation de charge

(cf. la question de l'électrolyse). On peut retenir la différence entre anode et cathode à l'aide du moyen mnémotechnique suivant : oxyd**A**tion à l'**A**node, et rédu**C**tion à la **C**athode.



## — Comment une pile fonctionne-t-elle ?

Une pile est donc constituée de deux compartiments (ou demi-piles) :

- Un dans lequel les cations métalliques sont réduits en atomes de métal solide, qui se déposent sur l'électrode : c'est la cathode. Il s'agit donc de la demi-pile vers laquelle les électrons affluent par le circuit extérieur pour alimenter cette réduction.
- Un dans lequel l'électrode de métal se désagrège progressivement par oxydation, pour former des cations métalliques : c'est l'anode. Il s'agit donc de la demi-pile depuis laquelle les électrons libérés affluent et partent dans le circuit extérieur.

Sachant que le courant électrique d'intensité positive se déplace en sens inverse des électrons dans le circuit extérieur, on pourra donc retenir que **lorsqu'une pile débite**, la cathode constitue sa borne  $\oplus$ , et l'anode constitue sa borne  $\ominus$ .

**Remarque :** lorsque la pile débite, le courant électrique d'intensité positive se déplace dans le circuit extérieur de la cathode vers l'anode, et dans le pont salin de l'anode vers la cathode. Il importe bien de garder en mémoire que lorsque l'orientation du courant électrique entre deux demi-piles est évoquée, c'est par défaut celle du courant circulant **dans le circuit extérieur** dont il est question.

La réaction peut être considérée comme totale (réaction d'oxydoréduction), et elle se déroulera donc jusqu'à ce qu'il ne reste pratiquement plus de cations dans la cathode (c'est en général ce cas de figure qui se présente), ou pratiquement plus de métal solide dans l'anode.

**Remarque :** nous avons vu qu'une réaction d'oxydoréduction ne pouvait pas toujours se faire : si elle se fait entre l'oxydant d'un premier couple et le réducteur d'un deuxième, elle ne se fera pas entre leurs réducteur et oxydant conjugués respectifs. Dans le cas où nous associons deux demi-piles, le problème ne se pose pas : chacune des demi-piles contenant à la fois l'oxydant et le réducteur d'un couple, elle peut fonctionner :

- En anode, si son oxydant est le plus fort des deux oxydants engagés par chacune des demi-piles.
- En cathode, si son réducteur est le plus fort des deux réducteurs engagés par chacune des demi-piles.

Chaque demi-pile est donc, formellement, le siège d'une demi-équation électronique. La pile résultant de leur association réalise alors, quoique délocalisée sur deux compartiments, la réaction reposant sur les deux demi-équations en jeu.

## Quelles sont les grandeurs caractéristiques d'une pile, et comment les déterminer expérimentalement ?

Comme tout générateur, une pile peut être modélisée par la mise en série d'un générateur de tension idéal de f.é.m.  $e$ , avec un conducteur ohmique de résistance  $r$  (typiquement, quelques ohms). La tension à ses bornes est donc  $u_{\text{pile}} = e - ri$ , où  $i$  est l'intensité du courant débité. Pour mesurer  $e$ , il suffit de placer un voltmètre aux bornes de la pile. En notant  $R$  la résistance interne du voltmètre (typiquement de l'ordre du  $M\Omega$ ), on peut aisément

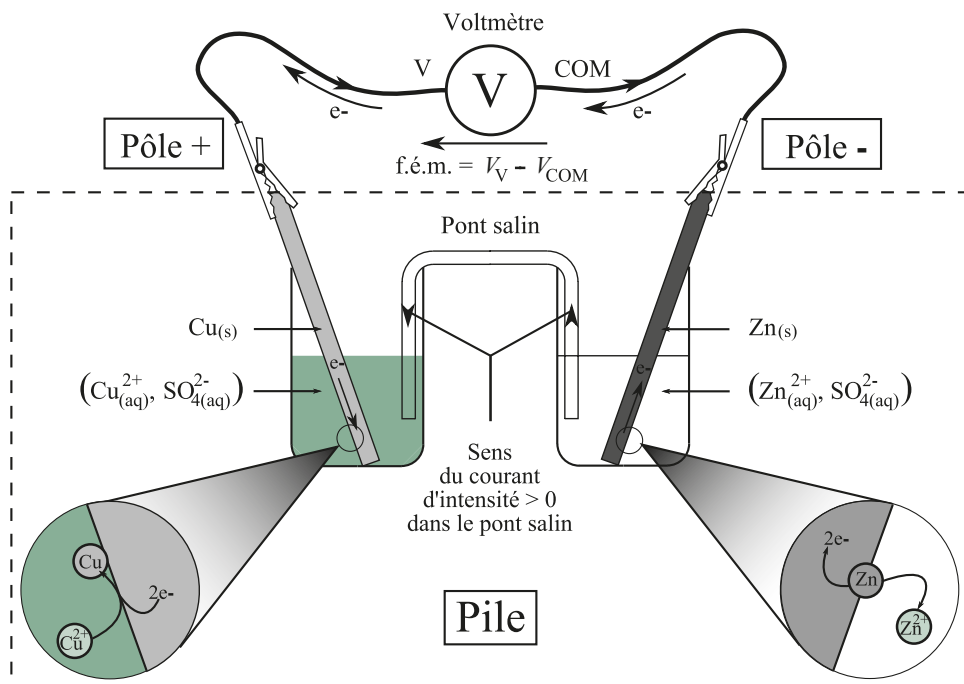
montrer que  $i = \frac{e}{r+R}$ , d'où  $u_{\text{pile}} = e - \frac{r}{r+R}e = \frac{R}{r+R}e$ . Comme  $R \gg r$ , on trouve alors que  $u_{\text{pile}} \approx e$  : on dit que l'on mesure la **tension à vide**.

Le voltmètre affiche toujours la différence de potentiel  $u_{\text{aff}} = V_V - V_{\text{COM}}$ , où les indices V et COM représentent les deux bornes du voltmètre. Ainsi, si la tension  $u_{\text{aff}}$  affichée par le voltmètre est positive, c'est que  $V_V > V_{\text{COM}}$ . Dans le circuit extérieur, le courant circule donc de la demi-pile branchée sur la borne V, vers celle branchée sur la borne COM. Les électrons circulant dans l'autre sens, on en déduit que des électrons sont libérés dans le circuit extérieur (donc que l'oxydation se produit) au niveau de la demi-pile côté COM, et captés au niveau de l'autre demi-pile (où se produit donc la réduction).

Ainsi, dans une pile qui débite, la demi-pile dont le potentiel est le plus bas est l'anode, tandis que celle dont le potentiel est le plus élevé est la cathode. Dans le cas de la pile Daniell, proposée plus haut, on observe :

- l'oxydation du zinc solide (la plaque se désagrège progressivement) dans la demi-pile au zinc ;
- la réduction des ions cuivre (II) (la solution se décolore) dans la demi-pile au cuivre.

La demi-pile au zinc constitue donc l'anode de ce montage, c'est-à-dire le pôle - de la pile, tandis que la demi-pile au cuivre constitue la cathode, c'est-à-dire le pôle +. La circulation du courant électrique suit donc le schéma suivant :



Une autre grandeur caractéristique d'une pile est l'intensité maximale du courant électrique qu'elle peut débiter. Nous avons vu plus haut que si un dipôle de résistance  $R$  est placé en série avec cette pile, l'intensité du courant circulant avait pour expression  $i = \frac{e}{r + R}$ . Elle est donc

maximale si  $R = 0$ , c'est-à-dire si le dipôle se résume à un fil (ou bien à un ampèremètre, de résistance interne pratiquement nulle). Pour cette raison, cette intensité est appelée **intensité de court-circuit**. Elle se mesure en plaçant un ampèremètre entre les bornes de la pile.

**Remarque :** la mesure de l'intensité de court-circuit se fait donc en court-circuitant la pile (logique). Il s'agit cependant d'une mesure qui doit appeler à la vigilance et n'être menée que sur une durée très courte, la mise en court-circuit d'un générateur entraînant une décharge très rapide de la pile, et pouvant mener à la détérioration du générateur, de l'ampèremètre voire de l'expérimentateur.

De la tension à vide (f.é.m.  $e$ ) et de l'intensité de court-circuit ( $i_{cc}$ ), on peut déduire la résistance interne de la pile. En effet, comme  $i_{cc} = \frac{e}{r}$ , on trouve ainsi que :

|                        |                 |
|------------------------|-----------------|
| $r = \frac{e}{i_{cc}}$ | $r$ en $\Omega$ |
|                        | $e$ en V        |
|                        | $i_{cc}$ en A   |

Une dernière grandeur que l'on peut déterminer concernant cette pile est sa **quantité d'électricité**  $Q$ , qui représente la charge totale qu'elle peut débiter. En effet, la pile ne peut débiter que tant qu'il reste du réducteur à oxyder et de l'oxydant à réduire. Les quantités de matière initiales étant finies, c'est un nombre fini d'électrons qui sera échangé entre les deux demi-piles, auquel on pourra associer une charge électrique également finie.

Pour déterminer la quantité d'électricité d'une pile, on procède de la manière suivante :

1. Connaissant la réaction qui se déroule entre les deux demi-piles ainsi que les quantités initiales des réactifs, on détermine le réactif limitant et la valeur  $x_{\max}$ .
2. On détermine le nombre  $n$  d'électrons échangés entre les réactifs à chaque fois que cette réaction a lieu.
3. Pour un avancement  $x$ , la quantité de matière d'électrons échangés entre l'oxydant et le réducteur a pour valeur  $nx$ , et le nombre total d'électrons échangés vaut donc  $n \times \mathcal{N}_A \times x$ , où  $\mathcal{N}_A$  est la constante d'Avogadro.
4. La valeur absolue de la charge véhiculée par ces électrons vaut donc :  $n \times \mathcal{N}_A \times x \times e$ .

Pour alléger les écritures, on introduit la **constante de Faraday**,  $\mathcal{F}$ , qui représente la valeur absolue de la charge électrique véhiculée par 1 mol d'électrons, rapportée à cette mole :

|                                 |                                                   |
|---------------------------------|---------------------------------------------------|
| $\mathcal{F} = \mathcal{N}_A e$ | $\mathcal{F}$ en $\text{C} \cdot \text{mol}^{-1}$ |
|                                 | $\mathcal{N}_A$ en $\text{mol}^{-1}$              |
|                                 | $e$ en C                                          |

Dans la pratique, la constante de Faraday a pour valeur  $\mathcal{F} = 9,65 \cdot 10^4 \text{ C} \cdot \text{mol}^{-1}$ .

La charge électrique  $Q$  débitée par une pile siége d'une réaction d'avancement  $x$  prend alors pour expression :

|                     |                                                   |
|---------------------|---------------------------------------------------|
| $Q = n\mathcal{F}x$ | $Q$ en C                                          |
|                     | $n$ sans unité                                    |
|                     | $\mathcal{F}$ en $\text{C} \cdot \text{mol}^{-1}$ |
|                     | $x$ en mol                                        |

**Remarque :** dans une installation domestique, les courants électriques sont souvent de l'ordre de l'ampère et les durées d'utilisation de l'ordre de l'heure. Il n'est donc pas rare, dans un exercice sur ce thème, de voir une quantité d'électricité exprimée non pas en coulomb mais en ampère-heure (Ah), sachant que, comme  $1 \text{ A.s} = 1 \text{ C}$ , alors  $1 \text{ Ah} = 1 \text{ A} \times 1 \text{ h} = 1 \text{ A} \times 3\,600 \text{ s} = 3\,600 \text{ A.s} = 3\,600 \text{ C}$ . C'est par exemple fréquemment le cas pour les piles et les batteries d'appareils électroportatifs.

## — Problème : que peut-on faire d'une pile vide ?

Les espèces formées au sein d'une pile pendant qu'elle débite ne peuvent réagir entre elles. Que faire alors de ces produits ? Par ailleurs, ne risque-t-on pas d'épuiser, à termes, les ressources en oxydant et en réducteur ?

Une solution intéressante serait d'inverser le processus, de manière à régénérer l'oxydant et le réducteur d'origine qui pourraient ainsi de nouveau échanger spontanément leurs électrons.

Si l'on branche, aux bornes d'une pile de f.é.m.  $e$ , un générateur de tension électrique  $u_G$  susceptible de débiter un courant électrique d'intensité positive, dans le sens opposé à celui débité par la pile (bornes  $\oplus$  en coïncidence et bornes  $\ominus$  également, donc), on constate que :

- Tant que  $u_G < e$ , le courant total d'intensité positive circule dans le sens imposé par la pile mais son intensité est d'autant plus réduite que la valeur de  $u_G$  est plus élevée.
- Si  $u_G = e$ , aucun courant ne circule.
- Si  $u_G > e$ , le courant d'intensité positive se renverse et circule donc cette fois dans le sens imposé par le générateur, contraire à celui qu'imposait la pile.

Expérimentalement, on constate alors que :

- la demi-pile où se produisait spontanément l'oxydation devient le siège d'une réduction ;
- l'autre demi-pile, où se produisait la réduction, devient le siège d'une oxydation.

L'imposition d'un générateur de tension  $u_G > e$  permet donc d'invertir l'anode et la cathode de la pile, et de forcer la réaction en sens inverse de son sens spontané.

**Remarque :** ce processus ne peut être convenablement réalisé qu'avec des piles spécialement conçues à cette fin, et que l'on appelle accumulateurs ou encore « accus » tout court.

Parfois, l'intensité débitée est nulle sur toute une plage de valeurs de  $u_G$  : le courant cesse de circuler avant que  $u_G$  n'atteigne  $e$ , et ne reprend en sens inverse que pour une valeur significativement supérieure à  $e$ . Ceci provient du fait que certains couples (dits **couples lents**) réclament un écart de potentiel particulièrement marqué pour que la réaction s'amorce. L'écart en question est appelé **surtension**.

## — Qu'est-ce que le phénomène d'électrolyse ?

Le phénomène mis en lumière à la section ci-dessus porte le nom d'**électrolyse**. Ses applications les plus courantes sont :

- Il permet de recharger des batteries conçues à cet effet.
- La cathode n'a pas besoin d'être constituée du même métal que les cations qui vont s'y trouver réduits. En plaçant en guise de cathode un objet métallique quelconque, celui-ci va se retrouver progressivement recouvert d'un dépôt de métal solide issu de la réduction des cations de la solution dans laquelle il baigne. Ce procédé permet le plaquage de certains bijoux, éléments de véhicules automobiles (chromes notamment)...

## Halte aux idées reçues

Expliquer, pour chacune des affirmations suivantes, pourquoi elle est fautive :

1. Le volume molaire d'une espèce chimique ne peut être défini que si cette espèce est gazeuse.
2. Masse volumique et concentration massique sont deux appellations d'une même grandeur.
3. Dans le tableau d'avancement d'une équation de réaction où figure une espèce spectatrice, les deux colonnes

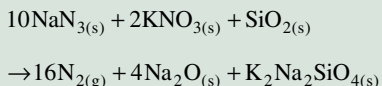
associées à cette espèce (une du côté des réactifs et l'autre du côté des produits) présentent des contenus différents.

4. Une solution est d'autant plus acide que le  $pK_A$  de l'acide s'y trouvant dissout est plus faible.
5. Lors d'une réaction d'oxydoréduction, une molécule réductrice cède forcément autant d'électrons qu'en acquiert une molécule d'oxydant.

## Du Tac au Tac

(Exercices sans calculatrice)

1. Le fonctionnement des airbags utilisés dans les automobiles repose sur l'explosion d'une cartouche contenant de l'azote de sodium  $\text{NaN}_{3(s)}$  ( $M_1 = 65,0 \text{ g.mol}^{-1}$ ), du nitrate de potassium  $\text{KNO}_{3(s)}$  ( $M_2 = 101,1 \text{ g.mol}^{-1}$ ) et de la silice  $\text{SiO}_{2(s)}$  ( $M_3 = 60,1 \text{ g.mol}^{-1}$ ). Les réactifs réagissent selon la réaction d'équation (supposée totale) :



Le violent dégagement de diazote qui s'ensuit provoque le gonflage du ballon.

Déterminer la masse de réactifs à prévoir pour gonfler un airbag de volume  $V_A = 160 \text{ L}$  (la pression et la température dans l'airbag sont telles, que le volume molaire des gaz a pour valeur  $V_m = 50,0 \text{ L.mol}^{-1}$  et l'on négligera le volume des produits solides devant celui du diazote formé).

2. On considère la réaction d'oxydation des ions iode  $\text{I}^-$  (couple  $\text{I}_2/\text{I}^-$ ) par les ions permanganate  $\text{MnO}_4^-$  (couple  $\text{MnO}_4^-/\text{Mn}^{2+}$ ). On mélange un volume  $V_1 = 20,0 \text{ mL}$  de solution aqueuse d'iodure de potassium ( $\text{K}_{(aq)}^+, \text{I}_{(aq)}^-$ ) à la concentration  $c_1 = 1,00.10^{-2} \text{ mol.L}^{-1}$ , et un volume  $V_2 = 50,0 \text{ mL}$  de solution aqueuse de permanganate de potassium ( $\text{K}_{(aq)}^+, \text{MnO}_{4(aq)}^-$ ) acidifiée, à la concentration  $c_2 = 4,00.10^{-3} \text{ mol.L}^{-1}$ .

On admettra que les ions  $\text{H}^+$  sont apportés en excès.

Déterminer le réactif limitant.

3. On mélange un volume  $V_1 = 100,0 \text{ mL}$  de solution aqueuse d'acide chlorhydrique à la concentration

$c_1 = 5,00.10^{-3} \text{ mol.L}^{-1}$ , et un volume  $V_2 = 50,0 \text{ mL}$  de solution aqueuse d'acide nitrique, à la concentration  $c_2 = 2,00.10^{-2} \text{ mol.L}^{-1}$ .

Déterminer le pH de la solution ainsi obtenue.

4. On considère un volume  $V_0 = 100,0 \text{ mL}$  de solution aqueuse d'acide éthanóïque  $\text{H}_3\text{CCOOH}$ , de concentration  $c_0 = 3,30.10^{-2} \text{ mol.L}^{-1}$ , à laquelle on ajoute une solution aqueuse de soude. Le pH augmente progressivement et, pour un volume total de solution aqueuse de soude versé  $V_1 = 50,0 \text{ mL}$ , atteint la valeur  $\text{pH} = 5,8$ .

$pK_A$  du couple acide éthanóïque/ion éthanóate : 4,8.

Déterminer les concentrations effectives à ce stade, respectivement en acide éthanóïque et en ions éthanóate.

5. On considère une pile, formée de :

- une anode constituée d'une pièce de cuivre métallique de masse  $m_1 = 12,7 \text{ g}$  baignant dans un volume  $V_1 = 100 \text{ mL}$  de solution aqueuse de sulfate de cuivre ( $\text{Cu}_{(aq)}^{2+}, \text{SO}_{4(aq)}^{2-}$ ) de concentration  $c_1 = 1,00.10^{-2} \text{ mol.L}^{-1}$  ;
- une cathode constituée d'une pièce d'argent métallique de masse  $m_2 = 10,8 \text{ g}$  baignant dans un volume  $V_2 = 100 \text{ mL}$  de solution aqueuse de nitrate d'argent ( $\text{Ag}_{(aq)}^+, \text{NO}_{3(aq)}^-$ ) de concentration  $c_2 = 5,00.10^{-2} \text{ mol.L}^{-1}$ .

Couples engagés respectivement dans les demi-piles 1 et 2 :  $\text{Cu}^{2+}/\text{Cu}$  et  $\text{Ag}^+/\text{Ag}$ , masses molaires atomiques du cuivre  $M_{\text{Cu}} = 63,5 \text{ g.mol}^{-1}$  et de l'argent  $M_{\text{Ag}} = 107,9 \text{ g.mol}^{-1}$ , la constante d'Avogadro  $N_A = 6,0.10^{23} \text{ mol}^{-1}$ , et la charge électrique élémentaire  $e = 1,6.10^{-19} \text{ C}$ .

Déterminer l'équation de la réaction dont cette pile est le siège, le réactif limitant, ainsi que sa quantité d'électricité.

 Vers la prépa

Une solution tampon est une solution contenant des quantités égales d'un acide et de sa base conjuguée.

Le pH du sang est régulé entre autres par le couple  $\text{H}_2\text{PO}_4^-/\text{HPO}_4^{2-}$  ( $\text{p}K_{\text{A},2} = 7,2$ ). Déterminer comment fabriquer un volume  $V = 1,0 \text{ L}$  de solution tampon de concentration totale en espèce phosphatée  $c_P = 2,00 \cdot 10^{-2} \text{ mol.L}^{-1}$ , à partir des réactifs proposés respectivement dans chacun des trois cas suivants :

**1.** Dihydrogénophosphate de potassium  $\text{KH}_2\text{PO}_4$  et hydrogénophosphate de potassium  $\text{K}_2\text{HPO}_4$ .

**2.** Dihydrogénophosphate de potassium  $\text{KH}_2\text{PO}_4$  et solution aqueuse de soude à  $c_0 = 5,0 \cdot 10^{-2} \text{ mol.L}^{-1}$  (on montrera que la réaction qui s'ensuit est quantitative).

**3.** Solution aqueuse d'acide orthophosphorique  $\text{H}_3\text{PO}_4$  (couple  $\text{H}_3\text{PO}_4/\text{H}_2\text{PO}_4^-$ ,  $\text{p}K_{\text{A},1} = 2,1$ ) à  $c_1 = 4,0 \cdot 10^{-2} \text{ mol.L}^{-1}$ , et solution aqueuse de soude à  $c_2 = 1,0 \text{ mol.L}^{-1}$  (on montrera que la réaction qui s'ensuit est quantitative).

## Halte aux idées reçues

**1.** Le volume molaire d'un gaz est défini comme le rapport du volume occupé par un gaz, à la quantité de matière d'entités gazeuses (qu'il s'agisse d'atomes ou de molécules) contenue dans ce volume. On peut parfaitement appliquer cette même définition pour un liquide ou un solide.

Pour l'eau liquide, par exemple, on a alors :

$$V_{m,H_2O} = \frac{V_{H_2O}}{n_{H_2O}} = \frac{V_{H_2O}}{\frac{m_{H_2O}}{M_{H_2O}}} = \frac{M_{H_2O}}{\rho_{H_2O}} = 18,0 \text{ mL} \cdot \text{mol}^{-1}$$

Ceci a quelque chose de logique : l'eau a une masse de 18,0 g par mole de molécules d'eau considérée. Or, on sait que l'eau liquide occupe un volume de 1 mL par gramme d'eau considérée. De ce fait, une mole d'eau pèse 18,0 g, et occupe donc 18,0 mL, d'où ce volume molaire de 18,0 mL·mol<sup>-1</sup> (notons au passage que ce volume molaire est très inférieur à celui des gaz, puisqu'en phase liquide la matière se trouve beaucoup plus condensée).

Pourquoi dans ce cas réserver l'usage du volume molaire aux gaz ? Tout simplement parce que le volume molaire tel qu'il est calculé ci-dessus repose sur deux grandeurs dépendant de l'espèce considérée :

- Sa masse molaire, qui se calcule très simplement.
- Sa masse volumique qui, elle, ne peut se calculer simplement à partir de grandeurs relatives aux atomes la constituant. On a donc alors besoin de connaître individuellement la masse volumique de chaque espèce, sachant que s'il n'existe qu'une centaine d'éléments chimiques, leurs combinaisons donnent en revanche des millions d'espèces chimiques différentes.

Ce problème lié à la masse volumique rend donc le volume molaire sans grand intérêt dans le cas général. Mais dans le cas particulier des gaz, on constate le petit miracle que résume la loi d'Avogadro-Ampère : le volume molaire ne dépend pas de la nature du gaz considéré : qu'il s'agisse d'hélium, de dihydrogène, de néon, de dioxygène, d'argon, de dioxyde de carbone, de vapeur d'eau, de diazote... Tous ont le même volume molaire, qui ne dépend que la température et de la pression.

Dans le cas d'un gaz, on n'a donc plus besoin d'écrire  $V_{m,X} = \frac{V_X}{n_X}$ , mais seulement  $V_m = \frac{V_X}{n_X}$ , où  $V_m$  a, pour une pression et une température données, la même valeur quel que soit X.

Cette simplification considérable fait que le volume molaire, est (mais dans le cas des gaz uniquement) une façon commode de déterminer une quantité de matière.

On peut ajouter que le volume molaire d'un gaz a pour expression, si celui-ci satisfait à la loi du gaz parfait :  $V_m = \frac{V_X}{n_X} = \frac{RT}{P}$ . Ceci permet de souligner deux points importants :

- Cette expression permet de calculer le volume molaire du gaz (supposé parfait), à partir de sa pression et de sa température. Attention toutefois : si  $R$  est fournie en USI, alors  $P$  doit être exprimée en pascals et  $T$  en kelvins, et le volume molaire sera obtenu en m<sup>3</sup>·mol<sup>-1</sup>.
- Si un gaz répond mal au modèle du gaz parfait (situation assez rare), son volume molaire diffère également de l'expression fournie ci-dessus, et l'on commence à perdre le bénéfice du volume molaire identique pour tous les gaz.

**2.** Si l'on utilise les formules sans réfléchir, par exemple dans le cas présent  $\frac{m}{V}$ , on peut effectivement être tenté d'assimiler

masse volumique et concentration massique. Les deux sont définies par le rapport d'une masse à un volume, donc ont la même unité.

Si l'on regarde d'un peu plus près, on constate que concentration massique et masse volumique sont deux grandeurs quantifiant deux relations de proportionnalités différentes :

- La masse d'un échantillon de corps pur est proportionnelle au volume qu'il occupe. La masse volumique de ce corps est le coefficient de proportionnalité entre le volume et la masse d'un échantillon de ce corps. On peut sous-titrer son unité comme kg (de corps pur), par mètre cube (de ce même corps pur) :

$$\rho_X = \frac{m_X}{V_X}$$

- La masse de soluté apporté dans un échantillon **de solution**, est proportionnelle au volume de cet échantillon : si l'on prélève un échantillon de cette solution, dont le volume est deux fois plus important qu'un autre, alors cet échantillon contient une masse de soluté apporté deux fois plus importante. La concentration massique de cette solution en ce soluté est le coefficient de proportionnalité entre le volume de solution prélevé et la masse de soluté apporté qui vient avec. On peut sous-titrer son unité comme kg (de soluté X apporté) par mètre cube (de solution) :

$$c_X = \frac{m_X}{V_{\text{solution}}}$$

3. Il est en règle général chaudement recommandé de ne jamais faire figurer les espèces spectatrices dans une équation de réaction. Outre le fait qu'elles n'y servent à rien puisque par nature elles n'y participent pas, elles font courir le risque d'écrire de graves incohérences.

En effet, une cellule du tableau d'avancement affiche la quantité de matière de l'espèce figurant en tête de colonne, au stade de la réaction figurant en tête de ligne. Si une espèce figure deux fois dans l'équation, on lui affectera deux colonnes, et les deux devront obligatoirement contenir les mêmes informations, à défaut de quoi cela signifierait qu'une même espèce peut être présente dans le milieu réactionnel en deux quantités différentes à un même instant.

On peut ajouter que si une espèce est spectatrice, alors elle n'est par définition ni consommée, ni formée au cours de la réaction. Il s'ensuit que sa quantité de matière ne doit pas varier au cours de la réaction. Ainsi, sa colonne dans un tableau d'avancement devrait non seulement présenter les mêmes valeurs partout où elle figure, mais aussi présenter la même valeur à chaque ligne.

Un autre point de vue est de considérer, dans chaque colonne, que l'espèce est à la fois consommée et formée. La soustraction de la quantité consommée compense la quantité formée, assurant ainsi le maintien de la quantité de matière à sa valeur initiale.

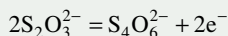
On le voit, la gestion rigoureuse d'une espèce spectatrice n'est pas seulement inutile : elle est également difficile. La meilleure option est donc de **ne pas faire figurer une espèce spectatrice dans l'équation de réaction**, et de supprimer d'éventuelles redondances d'espèces lorsqu'elle figurent à la fois en tant que produit et en tant que réactif (dans les équations d'oxydoréduction notamment).

4. Un  $pK_A$  faible atteste d'un couple dont l'acide se dissocie avec une aisance particulière, et corrélativement la base se protone difficilement. De ce fait, lorsque l'on introduit un acide dans l'eau, plus son  $pK_A$  est faible, mieux il se dissocie, plus il libère d'ions  $H^+$ , plus cela forme d'ions oxonium avec l'eau, et plus le pH baisse.

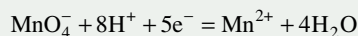
Donc, si l'on confectionne deux solutions de même concentration apportée en acide, celle dont l'acide aura le  $pK_A$  le plus faible sera celle de pH le plus faible. Mais si l'on envisage deux concentrations différentes, une concentration plus élevée de l'acide le moins dissocié peut tout à fait compenser cette moindre dissociation.

5. On sait que les électrons n'existent pas à l'état solvaté en solution aqueuse : ils ne peuvent être libérés dans l'eau par un réducteur et un oxydant ne les trouvera jamais se promenant dans l'eau tout seuls. Lorsqu'une réaction d'oxydoréduction se produit, il est donc vrai que les molécules d'oxydant doivent recevoir exactement autant d'électrons, qu'en cèdent les molécules de réducteur. Mais cette égalité collective n'est pas nécessairement individuelle.

On peut par exemple mettre en présence des ions thiosulfate  $S_2O_3^{2-}$ , intervenant en tant que réducteurs dans la demi-équation :



et des ions permanganate  $MnO_4^{2-}$  intervenant en tant qu'oxydant dans la demi-équation :



Les premiers cèdent les électrons par paquet de 2, tandis que les autres les captent par paquet de 5. Mais pour que le nombre total d'électrons cédé par les ions thiosulfate  $S_2O_3^{2-}$  soit égal à celui capté par les ions permanganate  $MnO_4^{2-}$ , chacune de ces demi-équations ne va pas être engagée autant de fois que l'autre. Ainsi, si la première se produit une fois, elle libère deux électrons, ce qui ne suffit pas à réduire un ion permanganate. Idem si elle se produit deux fois. Si elle se produit trois fois, 6 électrons sont libérés : on peut réduire un ion permanganate, mais il reste un électron inutilisé, qui ne peut rester sans preneur. En envisageant encore 2 cessions de 2 électrons, on passe à 10 électrons en tout, ce qui permet cette fois de réduire deux ions permanganate supplémentaires.

Le bilan de l'histoire est que, pour que les ions thiosulfate cèdent un nombre d'électrons à même d'être reçus par les ions permanganate, la première demi-équation doit être engagée 5 fois, lorsque la seconde l'est 2 fois. Ainsi les molécules réductrices ont cédé 10 électrons, autant qu'en ont reçu les molécules oxydantes, malgré le fait qu'individuellement, les molécules d'oxydant et de réducteur reçoivent et cèdent respectivement des nombres d'électrons différents.

## Du Tac au Tac

1. Connaissant le volume de l'airbag à gonfler de diazote, nous pouvons déterminer la quantité de matière  $n_{N_2,f}$  de diazote correspondante, à l'aide du volume molaire des gaz :

$$n_{N_2,f} = \frac{V_A}{V_m} = \frac{160 \text{ L}}{50,0 \text{ L.mol}^{-1}} = 3,20 \text{ mol}$$

Exprimons cette quantité de matière en fonction de la valeur finale  $x_f$  de l'avancement, sachant que le diazote est initialement absent (airbag vide :  $n_{N_2,i} = 0$ ), et qu'il est formé avec un nombre stœchiométrique égal à 16 :

$$n_{N_2,f} = 16x_f \Leftrightarrow x_f = \frac{n_{N_2,f}}{16} = 200 \text{ mmol}$$

La réaction doit donc avoir lieu  $x_f = 200 \text{ mmol}$  de fois pour former les  $n_{N_2,f} = 3,20 \text{ mol}$  de diazote nécessaires à gonfler les  $V_A = 160 \text{ L}$  de l'airbag. Calculons alors les quantités de matière de réactifs nécessaires et suffisantes pour atteindre cet objectif : on doit faire en sorte que :

$n_{1,i} \geq 0$ , donc la quantité juste suffisante en azoture de sodium  $NaN_{3(s)}$  vérifie :

$$n_{1,i} - 10 \times x_f = 0 \Leftrightarrow n_{1,i} = 10x_f = 2,00 \text{ mol}$$

$n_{2,f} \geq 0$ , donc la quantité juste suffisante en nitrate de potassium  $KNO_{3(s)}$  vérifie :

$$n_{2,i} - 2 \times x_f = 0 \Leftrightarrow n_{2,i} = 2x_f = 0,400 \text{ mol}$$

$n_{3,f} \geq 0$ , donc la quantité juste suffisante en silice  $\text{SiO}_{2(s)}$  vérifie :

$$n_{3,i} - 1 \times x_f = 0 \Leftrightarrow n_{3,i} = x_f = 0,200 \text{ mol}$$

Il ne nous reste plus, pour connaître la masse totale  $m_R$  de réactifs, qu'à calculer les masses correspondantes de réactifs :

$$m_{1,i} = n_{1,i} \times M_{\text{NaN}_3} = 130 \text{ g} \quad m_{2,i} = n_{2,i} \times M_{\text{KNO}_3} = 40,4 \text{ g}$$

$$m_{3,i} = n_{3,i} \times M_{\text{SiO}_2} = 12,0 \text{ g}$$

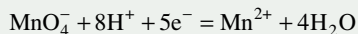
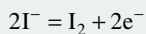
$$m_R = m_{1,i} + m_{2,i} + m_{3,i} = 182 \text{ g}$$

La cartouche doit donc contenir  $m_R = 182 \text{ g}$  de mélange stœchiométrique de réactif.

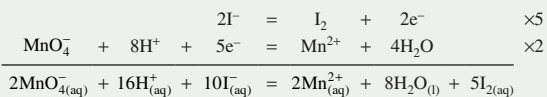
2. On commence par équilibrer la demi-équation d'oxydo-réduction de chacun des couples proposés.

Il est recommandé d'écrire les demi-équations en plaçant d'entrée de jeu le réactif à gauche, qu'il soit oxydant ou réducteur. Ceci d'une part évite d'avoir à inverser gauche et droite d'une demi-équation lors de l'écriture de l'équation finale. D'autre part on voit ainsi tout de suite, en observant les deux demi-équations, leur cohérence mutuelle, puisque l'une doit avoir les électrons à gauche (celle du réactif oxydant) et l'autre à droite (celle du réactif réducteur).

Dans le cas présent, on obtient donc :



La première demi-équation est donc engagée 5 fois lorsque la seconde est engagée 2 fois. On obtient alors :



Pour déterminer le réactif limitant, on commence par déterminer les quantités de matière de chaque réactif initialement introduites. Dans le cas présent, on a apporté :

- une quantité  $n_{1,\text{app}} = c_1 V_1 = 2,00 \cdot 10^{-4} \text{ mol}$  d'iodure de potassium, dont la dissociation fournit une quantité  $n_{\text{I}^-,i} = n_{1,\text{app}} = 2,00 \cdot 10^{-4} \text{ mol}$  d'ions iodure ;
- une quantité  $n_{2,\text{app}} = c_2 V_2 = 2,00 \cdot 10^{-4} \text{ mol}$  de permanganate de potassium, dont la dissociation fournit une quantité  $n_{\text{MnO}_4^-,i} = n_{2,\text{app}} = 2,00 \cdot 10^{-4} \text{ mol}$  d'ions permanganate.

Les deux réactifs ont donc été introduits en quantités égales, ce qui pourrait laisser croire qu'ils vont limiter la réaction au même stade. À ceci près qu'ils ne sont pas consommés au même rythme : pour 2 ions permanganate consommés, ce sont 10 ions iodure qui le sont également. Cette plus grande proportion d'ions consommée permet, à quantités de matière initiales égales, d'affirmer que les ions iodure constituent le réactif limitant. On peut s'arrêter à ce stade, mais on peut également préciser les valeurs maximales autorisées à l'avancement par chacun des réactifs :

- Pour les ions iodure : la réaction prend fin si et seulement si  $n_{\text{I}^-,f} = 0 = n_{\text{I}^-,i} - 10x_{\text{max,I}^-}$ ,

$$\text{d'où } x_{\text{max,I}^-} = \frac{n_{\text{I}^-,i}}{10} = 2,00 \cdot 10^{-5} \text{ mol.}$$

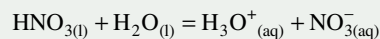
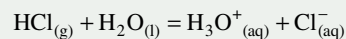
- Pour les ions permanganate : la réaction prend fin si et seulement si  $n_{\text{MnO}_4^-,f} = 0 = n_{\text{MnO}_4^-,i} - 2x_{\text{max,MnO}_4^-}$ ,

$$\text{d'où } x_{\text{max,MnO}_4^-} = \frac{n_{\text{MnO}_4^-,i}}{2} = 1,00 \cdot 10^{-4} \text{ mol.}$$

On constate donc, comme annoncé précédemment, que les ions iodure s'éteignent pour un avancement inférieur à celui provoquant l'extinction des ions permanganate. Ainsi constituent-ils le réactif limitant et imposent-ils à la réaction la valeur finale de son avancement  $x_f = x_{\text{max,I}^-} = 2,00 \cdot 10^{-5} \text{ mol}$ .

On peut ajouter que lorsque la réaction s'arrête, il reste en solution  $n_{\text{MnO}_4^-,f} = n_{\text{MnO}_4^-,i} - 2x_f = 1,60 \cdot 10^{-4} \text{ mol}$ .

3. Les deux acides considérés sont des acides forts, c'est-à-dire se dissocient totalement en solution aqueuse, selon les équations de réaction respectives :



Connaissant leurs concentrations respectives dans leurs solutions d'origine, nous pouvons en déduire les quantités de matière d'ions oxonium formées respectivement par chacune :

- Pour la solution aqueuse d'acide chlorhydrique :

$$n_{\text{H}_3\text{O}^+,1} = n_{\text{HCl,app}} = c_1 V_1 = 5,00 \cdot 10^{-4} \text{ mol.}$$

- Pour la solution aqueuse d'acide nitrique :

$$n_{\text{H}_3\text{O}^+,2} = n_{\text{HNO}_3,\text{app}} = c_2 V_2 = 1,00 \cdot 10^{-3} \text{ mol.}$$

Nous en déduisons une quantité de matière totale d'ions oxonium :  $n_{\text{H}_3\text{O}^+} = n_{\text{H}_3\text{O}^+,1} + n_{\text{H}_3\text{O}^+,2} = 1,50 \cdot 10^{-3} \text{ mol}$ .

Il ne nous reste plus, pour accéder à la concentration de ces ions dans le mélange, qu'à rapporter cette quantité de matière au volume du mélange en question, soit  $V = V_1 + V_2 = 150 \text{ mL}$  :

$$[\text{H}_3\text{O}^+] = \frac{n_{\text{H}_3\text{O}^+}}{V} = 1,00 \cdot 10^{-2} \text{ mol.L}^{-1}, \text{ d'où nous déduisons un pH égal à } 2,0.$$

4. On sait que les concentrations effectives de l'acide et de sa base conjuguée sont liées par :

$$\text{pH}_f = \text{pK}_A + \log \left( \frac{[\text{H}_3\text{CCOO}^-]_f}{[\text{H}_3\text{CCOOH}]_f} \right)$$

Soit encore :

$$\frac{[\text{H}_3\text{CCOO}^-]_f}{[\text{H}_3\text{CCOOH}]_f} = 10^{\text{pH}_f - \text{pK}_A} = 10^{-1,0} = 0,10$$

Nous en déduisons que, dans la solution finale, la forme acide du couple (l'acide éthanoïque  $\text{H}_3\text{CCOOH}$ ) est dix fois plus

présente que la forme basique, ce que nous pouvons réécrire sous la forme :

$$[\text{H}_3\text{CCOOH}]_f - 10[\text{H}_3\text{CCOO}^-]_f = 0$$

Nous disposons ici d'une équation à deux inconnues. Pour résoudre le problème, nous avons besoin d'une seconde équation engageant les mêmes inconnues. Celle-ci est fournie par la conservation de la matière : l'énoncé indique en effet que l'espèce apportée est l'acide éthanóique, mais si celui-ci réagit avec l'eau, il donne sa base conjuguée, l'ion éthanóate. Ainsi la population totale de ces deux espèces dans la solution demeure-t-elle constante.

Nous pouvons ainsi écrire :

$$n_{\text{H}_3\text{CCOOH},f} + n_{\text{H}_3\text{CCOO}^-,f} = n_{\text{H}_3\text{CCOOH},i} + n_{\text{H}_3\text{CCOO}^-,i}$$

Soit donc, en l'absence d'apport initial en ions éthanóate :

$$n_{\text{H}_3\text{CCOOH},f} + n_{\text{H}_3\text{CCOO}^-,f} = n_{\text{H}_3\text{CCOOH},i} = c_0 V_0 = 3,30 \cdot 10^{-3} \text{ mol}$$

En ramenant cette égalité au volume offert aux espèces solvatées, soit  $V'_0 = V_0 + V_1 = 150,0 \text{ mL}$ , nous obtenons alors la deuxième équation :

$$[\text{H}_3\text{CCOOH}]_f + [\text{H}_3\text{CCOO}^-]_f = \frac{c_0 V_0}{V'_0} = 2,20 \cdot 10^{-2} \text{ mol.L}^{-1}$$

En injectant dans la première relation, nous obtenons :

$$11[\text{H}_3\text{CCOO}^-]_f = 2,20 \cdot 10^{-2} \text{ mol.L}^{-1} \\ \Leftrightarrow [\text{H}_3\text{CCOO}^-]_f = 2,0 \cdot 10^{-3} \text{ mol.L}^{-1}$$

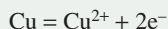
Nous trouvons enfin la concentration effective en acide :

$$[\text{H}_3\text{CCOOH}]_f = 10[\text{H}_3\text{CCOO}^-]_f = 2,0 \cdot 10^{-2} \text{ mol.L}^{-1}$$

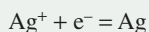
5. La pile présentée contient des quantités de matière significatives de l'oxydant aussi bien que du réducteur de chaque couple (ion et forme solide de chacun des deux métaux). Elle se prête donc *a priori* à deux réactions d'oxydoréduction : ion  $\text{Cu}^{2+}_{(\text{aq})}$  avec  $\text{Ag}_{(\text{s})}$ , ou  $\text{Cu}_{(\text{s})}$  avec  $\text{Ag}^{+}_{(\text{aq})}$ . Pour savoir laquelle est la bonne, nous avons besoin d'une information sur la polarité de la pile. Celle-ci nous est fournie par l'énoncé, qui précise que la demi-pile au cuivre est l'anode de la pile, tandis que celle à l'argent est la cathode.

Nous pouvons ainsi déduire que :

- La demi-pile au cuivre est le siège d'une oxydation. La demi-équation traduisant ce qu'il se passe dans cette demi-pile est donc :



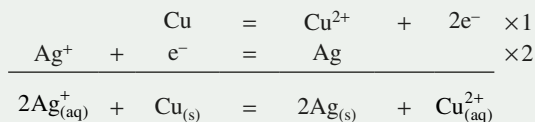
- De même, la demi-pile à l'argent est le siège d'une réduction. La demi-équation traduisant ce qu'il se passe dans cette demi-pile est donc cette fois :



Pour que le nombre total d'électrons cédé par l'espèce réductrice soit égal au nombre total d'électrons capté par l'espèce

oxydante, la seconde demi-équation doit se dérouler 2 fois lorsque la première se déroule 1 fois.

La pile est donc le siège de la réaction d'équation :



Pour déterminer le réactif limitant, il nous suffit dans un premier temps de déterminer les quantités de matière initiales de cuivre métallique et d'ions argent, soit respectivement :

- $n_{\text{Cu},i} = \frac{m_1}{M_{\text{Cu}}} = 2,00 \cdot 10^{-1} \text{ mol}$  ;
- $n_{\text{Ag}^+,i} = [\text{Ag}^+]_i V_2 = c_2 V_2 = 5,00 \cdot 10^{-3} \text{ mol}$ .
- Si le cuivre métallique est limitant, alors :

$$n_{\text{Cu},f} = 0 = n_{\text{Cu},i} - 1 \times x_{\text{max,Cu}} \\ \Leftrightarrow x_{\text{max,Cu}} = \frac{n_{\text{Cu},i}}{1} = 2,00 \cdot 10^{-1} \text{ mol}$$

- Si les ions argent sont limitants, alors :

$$n_{\text{Ag}^+,f} = 0 = n_{\text{Ag}^+,i} - 2 \times x_{\text{max,Ag}^+} \\ \Leftrightarrow x_{\text{max,Ag}^+} = \frac{n_{\text{Ag}^+,i}}{2} = 2,50 \cdot 10^{-3} \text{ mol}$$

On constate que les ions argent s'épuisent avant le cuivre métallique et constituent de ce fait le réactif limitant. Ils imposent donc à l'ensemble du système une valeur finale de l'avancement  $x_f = x_{\text{max,Ag}^+} = 2,50 \cdot 10^{-3} \text{ mol}$ .

Concernant la charge électrique qui aura transité à travers le circuit une fois la pile vidée (c'est-à-dire une fois les ions argent entièrement consommés), il nous suffit de déterminer le nombre  $N_{e^-}$  d'électrons qui a été échangé et de le multiplier par la charge électrique d'un seul électron.

Or nous constatons qu'à chaque fois qu'un ion argent est réduit, un électron est échangé (il y en a 2 de libérés lors de l'oxydation d'un atome de cuivre, mais vu que cette oxydation s'accompagne de la réduction de 2 ions argent, le résultat reste le même quel que soit le réactif sur lequel on fonde le raisonnement). Nous pouvons donc affirmer que la quantité de matière d'électrons échangée est égale à la quantité de matière d'ions argent consommée au cours de la réaction, soit donc  $n_{e^-} = n_{\text{Ag}^+,i} = 5,00 \cdot 10^{-3} \text{ mol}$  si l'on considère que la réaction est totale.

Nous en déduisons le nombre d'électrons correspondant :  $N_{e^-} = n_{e^-} N_A = 3,0 \cdot 10^{21}$  électrons, et enfin la charge électrique correspondante, qui en valeur absolue nous donne :

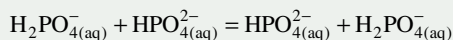
$$Q = N_{e^-} \times e = 4,8 \cdot 10^2 \text{ C}$$

En supposant que la pile puisse débiter par exemple un courant de 1 mA d'intensité (soit donc  $10^{-3} \text{ C}$  par seconde), nous constatons qu'elle pourra débiter durant  $4,8 \cdot 10^5 \text{ s}$ , c'est-à-dire 133 h. On dit de cette pile qu'elle contient une quantité d'électricité de 133 mAh.

Notons au passage que si elle débite 10 fois plus, elle durera 10 fois moins longtemps. L'intensité qu'elle débite dépend du circuit qu'elle alimente, et est limitée intrinsèquement par la nature et la qualité des réactifs engagés.

## Vers la prépa

1. Les espèces proposées se dissolvent dans l'eau en libérant respectivement les ions dihydrogénophosphate  $\text{H}_2\text{PO}_4^-$  et hydrogénophosphate  $\text{HPO}_4^{2-}$ . Ces deux espèces sont amphotères mais constituent respectivement l'acide le plus fort et la base la plus forte introduits dans le milieu. La réaction qui fixe prioritairement les concentrations des espèces qu'elle engage est donc celle du premier sur le second, d'équation :



Nous constatons que le quotient de réaction de cet équilibre prend systématiquement la valeur 1, ce qui semble logique : si par exemple une molécule d'acide est consommée, c'est qu'elle réagit avec une molécule de base qui se transforme du même coup en molécule d'acide. La molécule consommée est donc aussitôt remplacée par une autre. Les concentrations restent donc égales à leurs valeurs initiales et il suffit d'introduire les deux espèces proposées en quantités de matière égales et telles que la somme de leurs concentrations soit égale à la concentration requise, soit ici  $c_p$ .

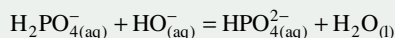
Nous souhaitons donc que la somme des quantités de matière en espèces phosphatées introduites vérifie :  $n_A + n_B = c_p V = 2,0 \cdot 10^{-2} \text{ mol}$ .

Comme, par ailleurs, nous souhaitons des quantités de matière égales de chacune de ces deux espèces, nous pouvons écrire que  $n_A = n_B$ . Nous disposons d'un système de deux équations à deux inconnues, dont la résolution donne immédiatement  $n_A = n_B = 1,0 \cdot 10^{-2} \text{ mol}$ .

Les espèces proposées étant des solides ioniques, nous devons mesurer leurs quantités en termes de masse, soit ici :

- $m_A = n_A M_A = 1,36 \text{ g}$  ;
- $m_B = n_B M_B = 1,74 \text{ g}$ .

2. Les espèces introduites sont cette fois  $\text{H}_2\text{PO}_4^-$  et  $\text{HO}^-$ . Elles constituent respectivement l'acide le plus fort et la base la plus forte introduits dans le milieu. C'est donc leur réaction qui va fixer prioritairement les concentrations des espèces qu'elle engage, selon l'équation :



dont la constante d'équilibre peut s'exprimer :

$$K = \frac{[\text{HPO}_4^{2-}]_f}{[\text{H}_2\text{PO}_4^-]_f [\text{HO}^-]_f} = \frac{K_{A,2}}{K_e} = 10^{6,8}$$

Cette réaction peut donc être considérée comme quantitative et nous pouvons en déduire qu'il suffit de verser une quantité d'ions  $\text{HO}^-$  égale à la moitié de la quantité d'ions  $\text{H}_2\text{PO}_4^-$

initialement présente. Ainsi, partant d'une quantité  $n_{A,i}$  de  $\text{H}_2\text{PO}_4^-$ , nous devons ajouter  $\frac{n_{A,i}}{2}$  d'ions  $\text{HO}^-$  pour transformer la moitié de ces ions  $\text{H}_2\text{PO}_4^-$  en ions  $\text{HPO}_4^{2-}$ .

Au final, nous obtiendrons donc  $\frac{n_{A,i}}{2}$  d'ions  $\text{HPO}_4^{2-}$  nouvellement formés, et  $\frac{n_{A,i}}{2}$  d'ions  $\text{H}_2\text{PO}_4^-$  restant.

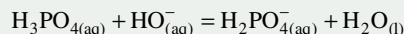
Comme nous avons vu à la question précédente que les quantités de ces deux ions étaient de  $n_A = n_B = 1,0 \cdot 10^{-2} \text{ mol}$ , nous en déduisons que  $\frac{n_{A,i}}{2} = 1,0 \cdot 10^{-2} \text{ mol}$ , d'où  $n_{A,i} = 2,0 \cdot 10^{-2} \text{ mol}$ . Nous en déduisons que la masse de dihydrogénophosphate de potassium à introduire est :  $m_{A,i} = n_{A,i} M_A = 2,72 \text{ g}$ .

Pour le reste, nous devons ajouter  $\frac{n_{A,i}}{2} = 1,0 \cdot 10^{-2} \text{ mol}$  d'ions  $\text{HO}^-$ . La solution de soude proposée étant concentrée à  $c_0 = 5,0 \cdot 10^{-2} \text{ mol} \cdot \text{L}^{-1}$ , nous en déduisons qu'il faut ajouter un volume  $V_2 = \frac{n_{A,i}}{2c_0} = 2,0 \cdot 10^{-1} \text{ L}$ .

Il ne reste plus qu'à verser cette masse et ce volume dans une fiole jaugée de 1,0 L et à compléter jusqu'au trait de jauge avec de l'eau distillée.

3. Les espèces introduites cette fois sont  $\text{H}_3\text{PO}_4$  et  $\text{HO}^-$ . Elles constituent respectivement l'acide le plus fort et la base la plus forte introduits dans le milieu.

C'est donc leur réaction qui va fixer prioritairement les concentrations des espèces qu'elle engage, selon l'équation :



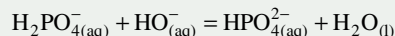
La constante d'équilibre peut s'exprimer :

$$K = \frac{[\text{H}_2\text{PO}_4^-]_f}{[\text{H}_3\text{PO}_4]_f [\text{HO}^-]_f} = \frac{K_{A,1}}{K_e} = 10^{11,9}$$

avec  $K_{A,1} = 10^{-pK_{A,1}} = 10^{-2,1}$ .

La valeur ainsi trouvée indique le caractère quantitatif de cette réaction et nous pouvons donc considérer qu'elle est pratiquement totale. Ainsi, en apportant  $\text{H}_3\text{PO}_4$  et  $\text{HO}^-$  en quantités égales, transformera-t-on tout l'acide  $\text{H}_3\text{PO}_4$  en sa base conjuguée,  $\text{H}_2\text{PO}_4^-$ . Mais nous souhaitons réaliser une solution contenant  $\text{H}_2\text{PO}_4^-$  et  $\text{HPO}_4^{2-}$  en quantités égales. Il est donc nécessaire de transformer la moitié des  $\text{H}_2\text{PO}_4^-$  en  $\text{HPO}_4^{2-}$ , par réaction avec les ions  $\text{HO}^-$  par exemple.

Intéressons-nous donc à la réaction qui se produit si nous ajoutons encore des ions  $\text{HO}^-$  : tout l'acide  $\text{H}_3\text{PO}_4$  ayant été consommé, l'acide le plus fort qui existe encore en solution est  $\text{H}_2\text{PO}_4^-$ . Celui-ci réagit avec la base la plus forte, soit ici  $\text{HO}^-$ , selon la réaction d'équation :



dont la constante d'équilibre peut s'exprimer :

$$K = \frac{[\text{HPO}_4^{2-}]_f}{[\text{H}_2\text{PO}_4^-]_f [\text{HO}^-]_f} = \frac{K_{A,2}}{K_e} = 10^{6,8}$$

Cette réaction est également quantitative, et nous pouvons en déduire qu'il suffit de verser une quantité supplémentaire en ions  $\text{HO}^-$ , égale à la moitié de la quantité d'ions  $\text{H}_2\text{PO}_4^-$  initialement présente dans cette deuxième réaction.

Ainsi, si nous partons d'une quantité  $n_{\text{AP}}$  de  $\text{H}_3\text{PO}_4$ , nous devons ajouter :

- Dans un premier temps,  $n_{\text{AP}}$  d'ions  $\text{HO}^-$  pour transformer toutes les molécules  $\text{H}_3\text{PO}_4$  en ions  $\text{H}_2\text{PO}_4^-$ .
- Dans un deuxième temps,  $\frac{n_{\text{AP}}}{2}$  d'ions  $\text{HO}^-$  pour transformer la moitié de ces ions  $\text{H}_2\text{PO}_4^-$  en ions  $\text{HPO}_4^{2-}$ .

Au final, nous aurons donc introduit au total  $n_{\text{AP}}$  d'acide  $\text{H}_3\text{PO}_4$ ,  $3\frac{n_{\text{AP}}}{2}$  d'ions  $\text{HO}^-$ , et nous aurons en solution  $\frac{n_{\text{AP}}}{2}$  d'ions  $\text{HPO}_4^{2-}$  nouvellement formés, et  $\frac{n_{\text{AP}}}{2}$  d'ions  $\text{H}_2\text{PO}_4^-$  restant.

Comme nous avons vu que les quantités de ces deux ions étaient de  $n_{\text{A}} = n_{\text{B}} = 1,0 \cdot 10^{-2}$  mol, nous en déduisons que  $\frac{n_{\text{AP}}}{2} = 1,0 \cdot 10^{-2}$  mol, d'où  $n_{\text{AP}} = 2,0 \cdot 10^{-2}$  mol. La solution d'acide proposée étant concentrée à  $c_1 = 4,0 \cdot 10^{-2}$  mol.L<sup>-1</sup>, nous en déduisons que le volume à prélever est  $V_1 = \frac{n_{\text{AP}}}{c_1} = 5,0 \cdot 10^{-1}$  L.

Pour le reste, nous devons ajouter  $3\frac{n_{\text{AP}}}{2} = 3,0 \cdot 10^{-2}$  mol d'ions  $\text{HO}^-$ . La solution de soude proposée étant concentrée à  $c_2 = 1,0$  mol.L<sup>-1</sup>, nous en déduisons qu'il suffit d'ajouter un volume  $V_2 = \frac{3n_{\text{AP}}}{2c_2} = 30$  mL de cette solution.

Il ne reste plus qu'à verser ces deux volumes dans une fiole jaugée de 1,0 L et à compléter jusqu'au trait de jauge avec l'eau distillée.

### 4.1 Principe d'un titrage

#### — Qu'est-ce qu'un dosage ?

L'analyse quantitative des espèces contenues dans une solution est indispensable dans de très nombreux domaines : diététique, police scientifique, analyses médicales... Dans ce but, on est souvent amené à déterminer la quantité de matière d'une certaine espèce, effectivement présente dans un milieu.

On appelle ainsi **dosage** d'une espèce dans un milieu la détermination de la quantité de matière de cette espèce, dans ce milieu.

**Remarque :** dans les faits, l'information importante est en général la **concentration** de cette espèce dans ce milieu. Une quantité de globules rouges dans une analyse sanguine, par exemple, ne prendra de sens que rapportée au volume dans lequel ces globules étaient contenus.

On peut distinguer deux méthodes de dosage : physiques et chimiques (titrage).

Les **méthodes physiques** : on mesure, sur un système chimique, une grandeur physique qui dépend de la concentration dans ce système de l'espèce à doser, selon une loi connue. Ainsi, par exemple, comme nous l'avons détaillé au chapitre 2 :

- Une mesure de pH permet de remonter à la concentration en ions oxonium.
- Une mesure de conductivité permet de remonter à la concentration apportée en un soluté ionique.
- Une mesure d'absorbance permet de remonter à la concentration d'une espèce colorée.

Le principal avantage que l'on peut tirer d'un usage direct de ces méthodes, est le fait qu'elles sont non destructives (l'espèce dont on souhaite déterminer la concentration est toujours présente en fin d'analyse). Ceci importe dans le cas où cette espèce est chère (sels d'argent, par exemple).

L'utilisation directe de ces mesures présente cependant le plus souvent une précision discutable : les conditions pour la validité de la loi liant la grandeur mesurée et la concentration recherchée sont souvent difficiles à respecter et chaque résultat découle d'une mesure unique.

Pour contourner ce problème, on préfère en général utiliser des **méthodes chimiques** ; on parle alors plutôt de **titrage** que de dosage. Elles consistent à introduire, en quantité connue, une espèce (dite titrante), avec laquelle l'espèce à titrer réagit selon une **réaction de stœchiométrie connue** (réaction dite de dosage) :

1. On introduit ainsi l'espèce titrante petit à petit, jusqu'à ce que la réaction de dosage ne se fasse plus malgré l'apport continu de cette espèce.
2. On sait alors que toute l'espèce titrée a réagi.
3. On détermine quelle quantité d'espèce titrante il a fallu introduire pour faire réagir toute l'espèce titrée.
4. Sachant dans quelles proportions les espèces titrée et titrante réagissent (puisque l'équation de la réaction, donc sa stœchiométrie, est connue), on peut en déduire la quantité d'espèce titrée qui était présente au début du titrage, objectif avoué de la démarche.

La précision des résultats dans ce nouveau cas repose sur la qualité de deux évaluations :

- celle du stade où la réaction de dosage cesse de se faire, faute d'espèce titrée ;
- celle de la quantité de matière d'espèce titrante versée lorsque ce stade est atteint.

Si la mesure de la quantité de réactif titrant versée au stade où la réaction de dosage cesse de se faire doit être réalisée avec la meilleure précision possible, la détection de cette fin de réaction, elle, ne réclame pas une analyse fine. Elle repose en général sur une simple appréciation qualitative telle que :

- la disparition d'une couleur (cas d'une espèce titrée colorée) ;
- la persistance d'une couleur (cas d'une espèce titrante colorée) ;
- le changement brutal dans l'évolution d'une grandeur physique dont on suivait l'évolution (pH, conductivité, et indirectement absorbance dans les deux cas invoqués ci-dessus) ;
- le changement de couleur d'un éventuel indicateur coloré, sous l'effet du changement de cette grandeur.

Tout l'intérêt du titrage repose donc sur le fait que contrairement aux méthodes physiques vues au chapitre 2, il n'est cette fois pas nécessaire de mesurer précisément les grandeurs qui varient lors de ce changement : il nous suffit de mesurer **quand** il se produit. On déplace ainsi la question de la précision des instruments, depuis ceux mesurant les grandeurs qui varient, vers ceux mesurant les volumes prélevés et versés. Or ceux-ci bénéficient généralement d'une fiabilité nettement meilleure.

La précision sur cette seconde catégorie de mesures devra donc être irréprochable, d'où le choix systématique d'éléments de verrerie jaugés ou gradués pour :

- confectionner les solutions ( fioles jaugées ) ;
- les prélever (pipettes jaugées ou graduées) ;
- les verser (burettes graduées).

Ces instruments sont disponibles en deux classes : B (précision de 0,5 % sur les volumes mesurés) et A (0,2 %).

## — Que se passe-t-il lorsque l'espèce titrée achève d'être consommée ?

Le stade d'un titrage où la réaction de dosage ne se fait plus malgré l'ajout d'espèce titrante constitue ce que l'on appelle l'**équivalence** de ce titrage.

Ce stade est atteint lorsque la quantité d'espèce titrante introduite depuis le début est juste suffisante pour avoir fait réagir la totalité de l'espèce titrée initialement présente dans le milieu réactionnel, c'est-à-dire lorsque les deux espèces ont été introduites en **proportions stœchiométriques**.

Les espèces sont alors toutes les deux consommées en totalité et il ne reste aucune trace de l'une ni de l'autre.

**Remarque :** il importe de bien retenir que les quantités de matière effectives des espèces titrée et titrante dans le milieu réactionnel sont **nulles à l'équivalence**. Les quantités non nulles que nous allons invoquer dans la suite sont les quantités **introduites**.

On retiendra donc que l'équivalence d'un titrage est le stade de ce titrage, où l'espèce titrante a été introduite dans le milieu réactionnel en proportions stœchiométriques, avec la quantité initialement présente d'espèce titrée. Notons alors :

- $X$  l'espèce titrée,  $\alpha_X$  son nombre stœchiométrique dans l'équation de la réaction de titrage,  $c_X$  sa concentration apportée (*a priori* inconnue) dans la solution titrée, et  $V_{X,i}$  le volume de cette solution sur lequel on réalise le titrage.

- Y l'espèce titrante,  $\alpha_Y$  son nombre stœchiométrique dans l'équation de la réaction de titrage,  $c_Y$  sa concentration (connue) dans la solution titrante, et  $V_{Y,E}$  le volume de cette solution versé à l'équivalence.

Nous pouvons alors écrire qu'à l'équivalence :

- La totalité de la quantité initiale d'espèce titrée a été consommée :  $n_{X(x_{f,E})} = 0 = n_{X,i} - \alpha_X x_{f,E}$ , d'où  $x_{f,E} = \frac{n_{X,i}}{\alpha_X}$ .
- La totalité de la quantité versée d'espèce titrante a été consommée :  $n_{Y(x_{f,E})} = 0 = n_{Y,E} - \alpha_Y x_{f,E}$ , d'où  $x_{f,E} = \frac{n_{Y,v}}{\alpha_Y}$ .

En explicitant alors  $n_{X,i} = c_X V_{X,i}$  et  $n_{Y,E} = c_Y V_{Y,E}$ , nous obtenons la relation :

|                                                               |                                         |
|---------------------------------------------------------------|-----------------------------------------|
| $\frac{c_X V_{X,i}}{\alpha_X} = \frac{c_Y V_{Y,E}}{\alpha_Y}$ | $c_X, c_Y$ en $\text{mol.L}^{-1}$       |
|                                                               | $V_{X,i}, V_{Y,E}$ en L                 |
|                                                               | $\alpha_X, \alpha_Y$ entiers sans unité |

Notons qu'à l'équivalence se produit un **changement de réactif limitant** :

- **Avant l'équivalence**, chaque ajout d'espèce titrante est entièrement consommé tandis qu'il reste de l'espèce titrée à faire réagir : **l'espèce titrante constitue le réactif limitant**.
- **Après l'équivalence**, on a beau rajouter de l'espèce titrante, celle-ci ne peut réagir puisque l'espèce titrée a été entièrement consommée. L'espèce titrante s'accumule donc sans pouvoir réagir, faute de partenaire de réaction : **l'espèce titrée constitue le réactif limitant**.

Ce changement de réactif limitant est à l'origine des brusques modifications dont le système chimique va être le siège au franchissement de l'équivalence (saut de pH, changement de couleur, rupture de pente sur les courbes de suivi conductimétrique).

On définit également la **demi-équivalence** d'un dosage comme le stade de ce dosage où la moitié de la quantité initiale d'espèce titrée a réagi avec l'espèce titrante.

La réaction de dosage étant supposée totale (cf. section suivante), cette demi-équivalence se produit pour une quantité d'espèce titrante versée, égale à la moitié de la quantité versée à l'équivalence. Ceci se retrouve dans le volume versé de solution titrante, c'est-à-dire que :

|                                  |                             |
|----------------------------------|-----------------------------|
| $V_{Y,1/2E} = \frac{V_{Y,E}}{2}$ | $V_{Y,1/2E}, V_{Y,E}$ en mL |
|----------------------------------|-----------------------------|

## — Quelles conditions une réaction chimique doit-elle remplir pour pouvoir servir de réaction de dosage ?

Les qualités de la réaction de dosage conditionnent naturellement celle du titrage réalisé : dans la mesure où sa cessation est l'unique critère dont nous disposons pour lier la quantité d'espèce titrante introduite à l'équivalence à celle d'espèce titrée initialement présente, il importe d'être sûr que cette fin est indiscutable.

Les écueils auxquels on risque le plus de se heurter sont :

1. **Le détournement d'espèce titrante** : cas où l'espèce titrante réagit non seulement avec l'espèce titrée, mais également avec une autre. La quantité d'espèce titrante versée pour faire réagir toute l'espèce titrée est alors surestimée, puisqu'il a fallu verser un supplément pour alimenter une réaction parasite.

2. **La réaction de dosage sans fin** : si la réaction de dosage aboutit à un équilibre, on peut toujours attendre l'équivalence : les réactifs étant constamment régénérés par la réaction inverse de la réaction de dosage, ils ne finiront jamais d'être consommés.
3. **La réaction de dosage qui prend son temps** : lorsque la réaction de dosage semble ne plus se faire, on doit pouvoir être sûr que c'est bien parce que l'espèce titrée n'est plus présente dans le milieu réactionnel. Et si l'on doit attendre une heure entre deux versements d'espèce titrante pour avoir la certitude que l'ajout précédent a bien été consommé en totalité, les TP risquent d'être longs.
4. **La réaction de dosage qui ne dit pas quand elle est finie** : tout le dosage repose sur la quantité de matière d'espèce titrante versée à l'équivalence. Encore faut-il pouvoir détecter cette équivalence, c'est-à-dire disposer d'un signe indiquant que, malgré l'ajout d'espèce titrante, la réaction de dosage ne se fait plus.

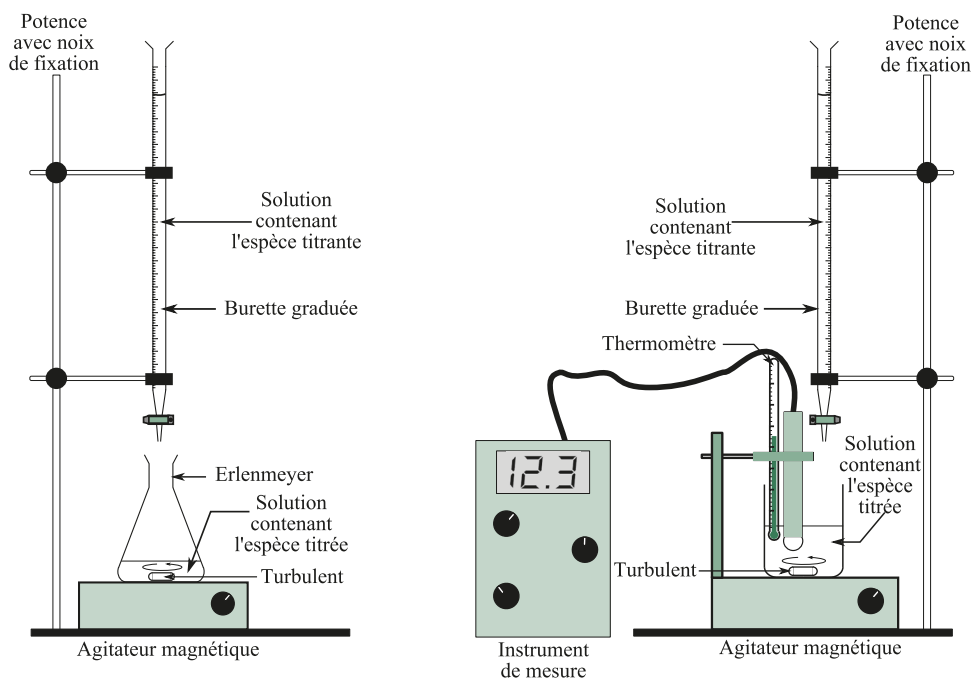
On peut ainsi déduire qu'une réaction ne peut être utilisée comme réaction de dosage que si elle est :

- |           |                      |
|-----------|----------------------|
| 1. Unique | 3. Rapide            |
| 2. Totale | 4. De fin détectable |

## — Comment procède-t-on dans le cas où la fin de la réaction de dosage est détectable ?

On utilise ici la méthode la plus simple : le **dosage direct**. C'est la méthode la plus simple : l'espèce titrante Y réagit avec l'espèce titrée X selon une réaction dont la fin est détectable d'une façon ou d'une autre.

Le dispositif est toujours le même, avec une variante selon que le dosage requiert ou non le suivi d'une grandeur physique :



- L'espèce titrée est en général en solution dans l'erlenmeyer. On prend cet élément de verrerie car son ouverture, moins large que celle d'un becher, minimise une éventuelle évaporation de la solution (le solvant et le soluté ne s'évaporent pas dans les mêmes proportions),

ainsi que le risque de perte par projection. Dans le cas où l'on effectue simultanément le suivi d'une grandeur physique (pH, conductivité), on peut remplacer l'ermeneyer par un becher, dont l'ouverture plus large permet l'introduction des sondes et d'un thermomètre en plus du bec de la burette destinée à verser la solution titrante.

- L'espèce titrante est en général en solution dans la burette graduée, qui permet une détermination précise (à 0,1 mL près le plus souvent) du volume versé, notamment à l'équivalence. Idéalement, on doit rincer la burette avec la solution titrante avant de s'en servir.
- Le milieu doit toujours être agité afin que le milieu reste homogène : la détection de fin de réaction doit concerner la totalité du milieu, non n'être qu'un événement local.

**Remarque :** il arrive, notamment lorsque l'on utilise des sondes de mesure, que le volume de solution dans lequel est réalisé le titrage ne permettent pas l'immersion complète de la sonde. On peut dans ce cas rajouter de l'eau sans états d'âme : celle-ci ne change pas la quantité de matière de soluté dans le milieu, or c'est celui-ci, non le solvant, qui est titré.

Pour être valable, un dosage doit être réalisé avec autant de précision que possible (à la goutte de solution titrante près). Ceci pose le problème des quantités de solution titrante versées :

- Si elles sont trop grandes, on risque de dépasser l'équivalence.
- Si elles sont trop faibles, le dosage risque de prendre des heures.

On résout parfois le problème en menant deux dosages consécutifs :

- Le premier grossièrement, qui vise à déterminer l'équivalence de manière approximative.
- Le second grossièrement jusqu'à ce que l'on approche de l'équivalence, puis de manière plus fine (on réduit les volumes de solution titrante versés) à mesure que les signes précurseurs de l'équivalence deviennent ostensibles (écarts croissants de pH dans un titrage acido-basique, délai croissant de disparition de la couleur de l'espèce titrante si elle est colorée, etc.).

## — Que faire si la fin de la réaction n'est pas détectable ?

Si la fin de la réaction de dosage n'est pas détectable à l'œil nu, on doit faire appel à des méthodes un peu plus sophistiquées. La première option est le **dosage indirect** :

1. On cherche, parmi les produits formés au cours de la réaction de dosage, s'il en est un (notons-le  $X'$ ) qui, lui, serait dosable par une espèce  $Y$ .
2. On introduit alors dans la première solution (celle contenant l'espèce titrée de départ  $X$ ) une quantité que l'on sait excessive d'un réactif transformant  $X$  en  $X'$ .
3. Sachant que toute l'espèce  $X$  a réagi, elle constitue le réactif limitant de cette première réaction. L'avancement final (inconnu à ce stade du raisonnement) de cette première réaction a donc pour valeur  $x_{l,f} = \frac{n_{X,i}}{\alpha_X}$ , en reprenant les notations utilisées plus haut.
4. Connaissant la stœchiométrie de cette première réaction de dosage, on peut alors exprimer la quantité de l'espèce  $X'$  qui a été formée comme :  $n_{X',f} = \alpha_{X'} \times x_{l,f} = \frac{\alpha_{X'}}{\alpha_X} n_{X,i}$ , en notant  $\alpha_{X'}$  le nombre stœchiométrique portant sur le produit  $X'$ .
5. On dose le produit  $X'$  par l'espèce  $Y$  (indifférente aux autres espèces en jeu, naturellement), et nous pouvons déterminer  $n_{X',f} = \frac{\alpha_{X'}}{\alpha_Y} n_{Y,E}$  ( $n_{Y,E}$  étant la quantité de  $Y$  versée à l'équivalence).
6. Il ne reste plus alors qu'à déduire  $n_{X,i} = \frac{\alpha_X}{\alpha_{X'}} n_{X',f}$ .

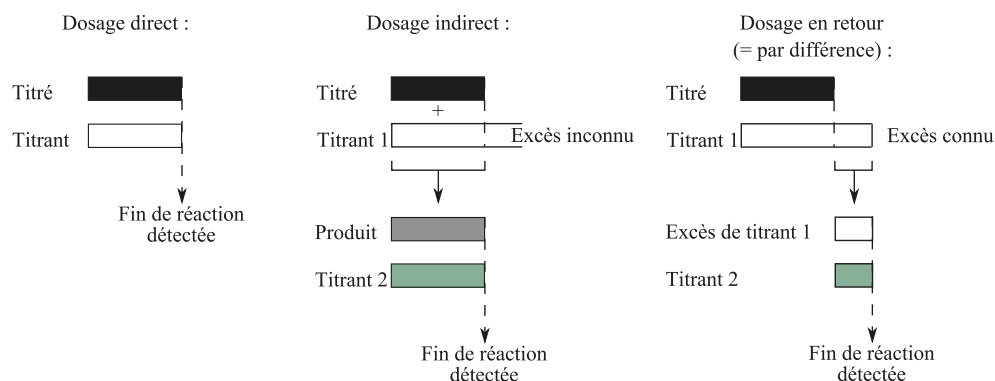
## — Que faire si la fin de la réaction n'est pas détectable et que seule l'espèce titrante est dosable ?

Nous avons vu dans le cas ci-dessus qu'il était parfois possible de doser, à défaut de l'espèce titrée proprement dite, une espèce dérivée de celle-ci. Toutefois cette option n'est pas toujours disponible et/ou il se trouve que l'espèce titrante peut elle-même être dosée par une autre espèce titrante  $Y'$ .

La deuxième option que l'on peut proposer est alors le **dosage en retour** ou **dosage par différence** :

1. On introduit l'espèce titrante en excès, connu cette fois avec précision (contrairement au dosage indirect, où le seul fait d'être en excès suffisait).
2. La réaction de dosage se fait donc jusqu'à consommation complète de l'espèce titrée, et il reste à la fin un excédent d'espèce titrante, que nous noterons  $n_{Y,f}$ .
3. Sachant que toute l'espèce titrée a réagi, elle constitue le réactif limitant de cette première réaction. L'avancement final (inconnu à ce stade du raisonnement) de cette première réaction a donc pour valeur  $x_{l,f} = \frac{n_{X,i}}{\alpha_X}$ .
4. Or, la quantité finale d'espèce Y peut s'exprimer :  $n_{Y,f} = n_{Y,i} - \alpha_Y x_{l,f} = n_{Y,i} - \frac{\alpha_Y}{\alpha_X} n_{X,i}$ .
5. Si l'on peut alors doser l'excédent d'espèce titrante Y par une autre espèce  $Y'$  (indifférente aux autres espèces en jeu, naturellement), nous pouvons déterminer  $n_{Y,f}$ .
6. Il ne reste plus alors qu'à déduire  $n_{X,i} = \frac{\alpha_X}{\alpha_Y} (n_{Y,i} - n_{Y,f})$ .

Le schéma suivant résume ces trois techniques :



## — 4.2 Méthodes de suivi d'un titrage

### — Quelles sont les grandeurs dont le suivi permet de détecter l'équivalence d'un titrage ?

Comme nous l'avons déjà dit, à l'équivalence d'un titrage se produit un changement de réactif limitant (avant l'équivalence, c'est le manque de titrant qui bloque la réaction de titrage,

après l'équivalence c'est le manque de titré). On observe donc également un changement de réactif en excès (avant l'équivalence : le titré, après l'équivalence : le titrant). Or le réactif en excès imposant ses propriétés au milieu réactionnel, on observera donc un changement dans les caractéristiques physiques (couleur, conductivité, pH, etc.) dudit milieu, lorsque l'équivalence sera franchie.

Le jeu des équilibres chimiques fait en outre que ce changement ne prendra en général effet que lorsque l'équivalence sera très proche et que le titré rendra son dernier soupir. Le changement en question sera donc souvent brusque, et c'est tant mieux : plus la marge sur laquelle il se produira sera étroite, plus l'incertitude sur la quantité de titrant nécessaire à l'atteinte de l'équivalence sera faible.

La/les grandeur(s) qu'il peut être utile de suivre dépend(ent) du contexte. Nous détaillons ci-dessous les différents types d'analyse que vous pourrez être amené(e) à rencontrer.

#### Analyses spectroscopiques :

- Si une ou plusieurs espèces intervenant dans l'équation de réaction est/sont colorée(s), on peut dans certains cas (apparition/disparition brusque d'une couleur) se contenter d'une appréciation purement visuelle (**titrage colorimétrique**).
- Si les circonstances nécessitent un suivi plus fin (apparition/disparition de la couleur trop progressive), on peut mener un **suivi spectrophotométrique** : moyennant la connaissance du coefficient d'étalonnage dans la loi de Beer-Lambert pour l'espèce colorée, une mesure d'absorbance nous donne accès à la concentration en solution de cette espèce et permet ainsi de détecter :
  - sa disparition si l'espèce est un réactif ;
  - son absence d'augmentation s'il s'agit d'un produit.
- Si l'espèce considérée possède une signature particulièrement marquée dans un **autre domaine spectroscopique** (IR que vous avez vu, RMN que vous verrez en CPGE), on peut également recourir à ces autres techniques, de manière analogue.
- Si les espèces intervenant dans l'équation de réaction ne permettent pas en elles-mêmes un suivi de type colorimétrique, spectroscopique ou spectrophotométrique, on peut avoir recours à des **indicateurs colorés**. Il s'agit en quelque sorte de mouchards colorimétriques : des couples d'espèces dont les membres présentent des couleurs différentes, sachant que le membre dominant dépend de la valeur d'une certaine grandeur dans le milieu réactionnel (typiquement : le pH, le potentiel électrique de la solution, etc.).

Si cette grandeur varie brutalement au franchissement de l'équivalence, l'espèce prédominante (et avec elle la couleur, le spectre en absorbance... qu'elle affecte au milieu réactionnel) changera avec la même brutalité, rendant ainsi l'équivalence détectable.

- On trouve des indicateurs colorés notamment dans les domaines acido-basique et d'oxydoréduction. Par ailleurs les complexes sont très souvent colorés et se prêtent donc en général bien au jeu.

**Analyse pH-métrique** : domaine très particulier (on se limite à apprécier l'abondance d'une seule espèce : les ions oxonium  $\text{H}_3\text{O}^+$ ), mais qui dans le même temps concerne une famille extrêmement vaste d'espèces (celle des acides et des bases). Ce type d'analyse sera étudié en détail dans les pages qui suivent.

**Analyse conductimétrique** : dès lors que la réaction de dosage engage des espèces électriquement chargées et libres de se déplacer dans le milieu réactionnel (des ions en solution, typiquement). Comme pour les analyses pH-métriques, ce cas se présente très souvent et fera l'objet d'un développement dans la suite de cette section.

**Analyse potentiométrique** : « pH » est littéralement l'abréviation de « potentiel hydrogène », le potentiel évoqué étant le potentiel électrique. Ainsi le pH-mètre mesure-t-il en réalité la différence de potentiel électrique entre la solution dans laquelle son électrode est

immergée, et une électrode de référence (généralement l'électrode au calomel saturée). Le concept a en fait une portée très générale, et il est possible de mesurer le potentiel électrique de n'importe quelle solution, potentiel qui dépend notamment des espèces s'y trouvant dissoutes, en particulier celles dotées de propriétés d'oxydoréduction.

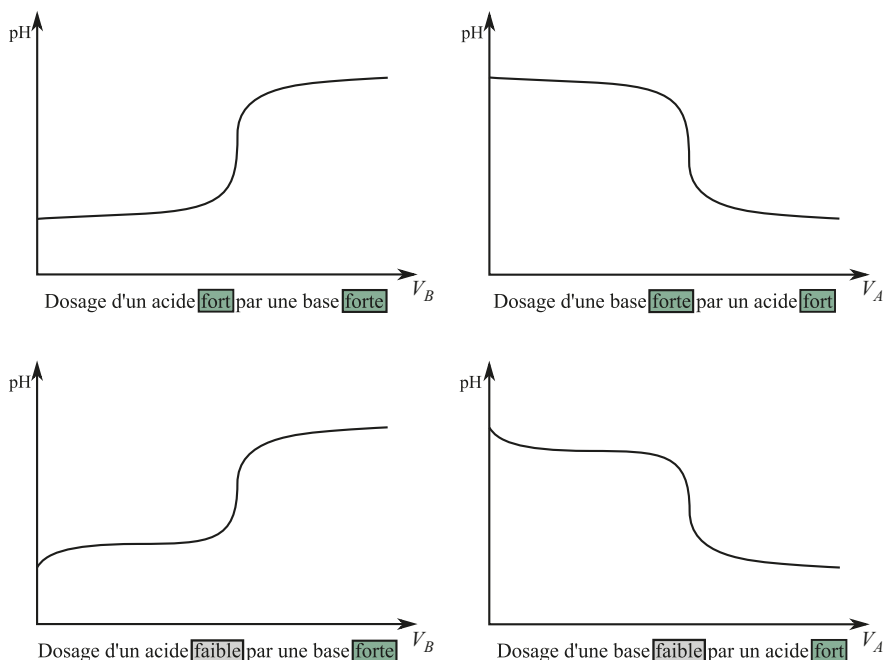
Ainsi si un titrage repose sur une réaction d'oxydoréduction, on retrouve la mécanique habituelle : le franchissement de l'équivalence signe le changement de réactif limitant (et donc de réactif en excès). Et si l'espèce en excès passe de l'oxydant d'un couple au réducteur d'un autre (ou l'inverse), on observera un changement brutal de potentiel, de la valeur imposée par le premier à celle imposée par le second. Une électrode mesurant ce potentiel révélera ce changement et permettra de détecter l'équivalence.

Notons par ailleurs qu'il est tout à fait possible de mener plusieurs de ces analyses de front, la cohérence des différents résultats entre eux renforçant mutuellement leur crédibilité (et dans le cas contraire, ouvrant la voie à plein de questions amusantes).

## — Comment détecte-t-on l'équivalence au cours d'un dosage pH-métrique ?

La plupart des acides sont incolores en solution aqueuse, ainsi que la plupart des bases. Les détectations de fin de réaction fondées sur la disparition ou la persistance de la couleur d'une espèce engagée dans la réaction support du titrage sont donc assez rares en ce domaine.

La base de travail sera ici le constat expérimental suivant : lorsque l'on réalise le dosage d'un acide par une base forte, ou d'une base par un acide fort, on obtient les courbes d'allure suivante :



Quelques remarques :

- L'espèce titrante est aussi souvent que possible un acide ou une base **fort(e)**, c'est-à-dire nivelé(e) respectivement aux ions  $\text{H}_3\text{O}^+$  ou  $\text{HO}^-$ . Les réactions ne sont donc pas, à stric-

tement parler, des réactions totales, mais aboutissent à des équilibres dont les constantes auront pour valeurs  $10^{pK_e - pK_A}$  (dosage d'un acide par les ions  $\text{HO}^-$ ) ou  $10^{pK_A}$  (dosage d'une base par  $\text{H}_3\text{O}^+$ ).

Ces valeurs sont donc en général suffisamment élevées (sauf si l'espèce titrée est vraiment très faible) pour que la réaction puisse être considérée comme quantitative, condition nécessaire pour être utilisée comme réaction de dosage.

Le dosage d'un acide faible par une base faible (ou l'inverse) est quant à lui plus risqué. Par exemple, si l'on dose l'acide éthanóique (acide relativement bon :  $pK_A = 4,8$ ) par l'ammoniac (base relativement bonne :  $pK_A = 9,2$ ), la réaction de dosage a pour constante  $10^{9,2-4,8} = 10^{4,4}$ . Les nombres stœchiométriques étant tous égaux à 1, la réaction peut alors être considérée comme quantitative.

En revanche, le dosage des ions éthanóate par l'ion ammonium offre une constante d'équilibre de  $10^{4,8-9,2} = 10^{-4,4}$  et est donc aussi mauvaise que la précédente était bonne.

- On peut noter que lorsque l'on dose une espèce forte par une autre, le début de la courbe part tout droit : on ne dose que des  $\text{H}_3\text{O}^+$  (ou des  $\text{HO}^-$ ) et rien d'autre.

Dans le cas en revanche où l'on dose une espèce faible, on observe un arrondi au départ de la courbe, qui correspond au dosage des quelques ions oxonium (ou hydroxyde) issus de la réaction de l'espèce dosée avec l'eau. Une fois ce début expédié, on attaque réellement le dosage de l'acide (ou de la base) introduit(e) au départ.

Notons cependant que ceci ne change rien à notre dosage puisque tout ion  $\text{H}_3\text{O}^+$  dosé au début correspond à une molécule d'acide dissociée, qui n'a donc plus à être dosée.

On constate, en travaillant avec une solution titrée de concentration connue à l'avance, que le saut ou la chute brutal(e) de pH observé(e) se produit précisément au niveau de l'équivalence. Une approche théorique (un peu trop conséquente pour figurer dans cet ouvrage) permet en outre d'interpréter ce saut à partir de considérations portant sur les équilibres dont le système est le siège tout au long du titrage.

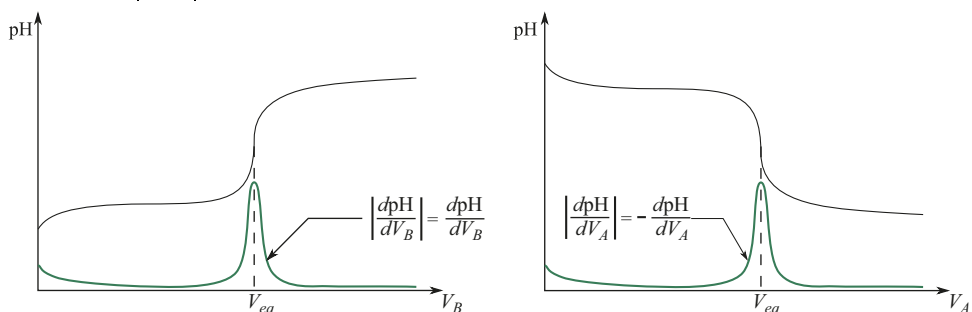
Cette brusque variation de pH est donc bien le témoignage du franchissement de l'équivalence. Sa détection peut se faire de trois façons :

### 1. Utilisation de la valeur absolue de la dérivée de la fonction $\text{pH} = f(V_Y)$ : le saut de pH

correspond à un pic de la valeur de  $\left| \frac{dpH}{dV_Y} \right|$ . On constate en effet que la tangente à la courbe représentative de cette fonction passe par une raideur maximale à l'équivalence, du fait de la brusque variation du pH.

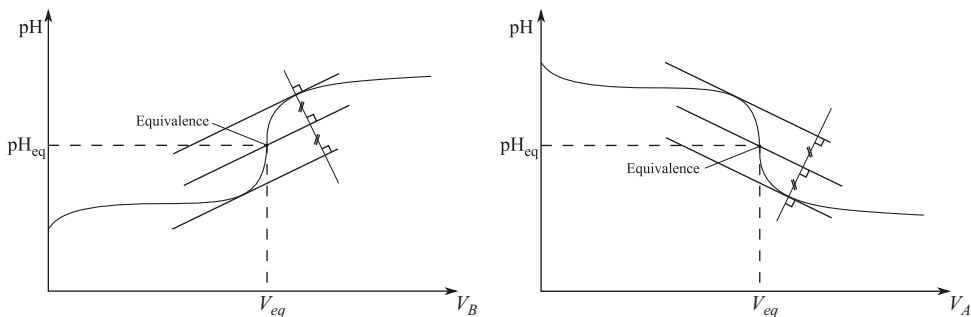
Si l'on dispose du logiciel adéquat, on peut alors saisir les coordonnées des points expérimentaux ( $V_Y$ , pH) et il ne reste plus qu'à approximer le nuage de points par un modèle. En faisant calculer et représenter les valeurs de la dérivée de ce modèle, on peut détecter

le point où  $\left| \frac{dpH}{dV_Y} \right|$  présente un maximum, et qui correspond à l'équivalence.



**2. Méthode des tangentes :** on procède de la manière suivante :

- (a) On trace deux tangentes à la courbe, l'une juste avant et l'autre juste après le saut ou la chute de pH, qui doivent être **parallèles entre elles**.
- (b) On trace ensuite une droite parallèle à ces deux droites, équidistante de celles-ci.
- (c) Le point d'intersection de cette droite avec la courbe constitue le point d'équivalence.

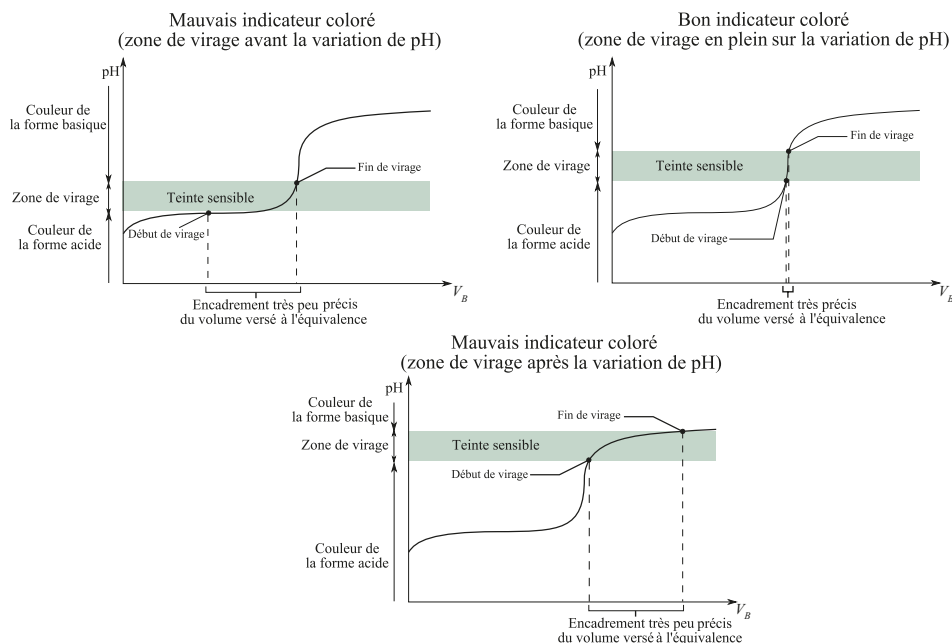


Cette méthode présente l'avantage, par rapport à la précédente, de permettre la détermination graphique du volume de solution titrante versé à l'équivalence, mais également de la valeur du pH à l'équivalence, encore qu'avec une précision discutable.

**3. Utilisation d'un indicateur coloré acido-basique :** les acides et bases les plus courants ne sont pas colorés mais certains, plus sophistiqués, le sont. Mieux, leurs formes acide et basique sont de couleurs différentes.

Or nous savons que, selon les valeurs relatives du pH et du  $pK_A$ , la forme acide du couple prédomine sur sa base conjuguée ( $pH < pK_A - 1$ ), ou bien le contraire ( $pH > pK_A + 1$ ). La forme prédominante de ce couple, et avec elle la couleur qu'elle confère à la solution, dépendent donc de la valeur du pH.

On appelle alors **zone de virage** l'intervalle des valeurs du pH dans lequel se fait le changement d'espèce prédominante. La zone de virage est en général haute de 2 à 3 unités de



pH. La couleur adoptée par la solution dans cette zone résulte de la synthèse soustractive des couleurs des formes acide et basique en proportions comparables et est appelée **teinte sensible**.

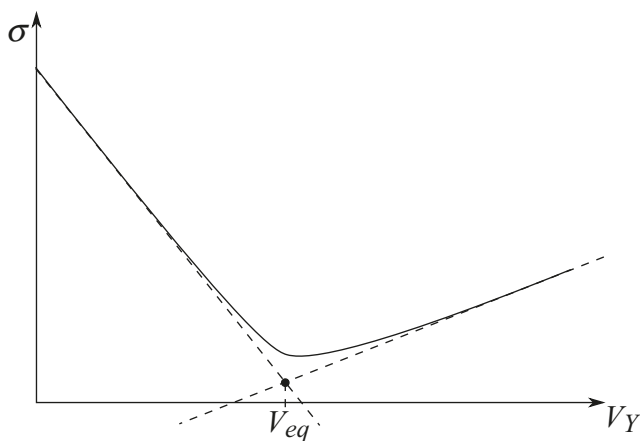
Donc si le pH du milieu dans lequel se trouve un indicateur coloré passe brusquement d'une valeur où prédominait sa forme acide à une autre où prédomine sa forme basique (ou l'inverse), la couleur de la solution va également brusquement passer de la couleur de la forme acide de cet indicateur à celle de sa forme basique (ou l'inverse).

Cette technique nous permet donc de greffer un critère colorimétrique sur une réaction de dosage n'engageant à l'origine aucune espèce colorée. Il convient cependant de noter deux conditions d'utilisation des indicateurs colorés acido-basiques :

- (a) La forme prédominante de l'indicateur coloré doit changer **brusquement**. L'indicateur coloré doit donc être choisi de sorte que sa **zone de virage soit incluse dans l'intervalle des valeurs couvertes par le saut de pH**.
- (b) L'indicateur coloré n'est présent que comme espèce témoin : il doit subir le pH imposé par la solution, surtout pas l'imposer lui-même, or il est par nature doté de propriétés acido-basiques puisqu'il EST un couple acide/base. Pour ce faire, on doit n'en introduire que quelques gouttes, qui suffisent d'ailleurs généralement à donner à la solution une couleur détectable à l'œil nu.

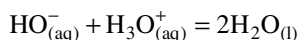
## — Comment détecte-t-on l'équivalence au cours d'un dosage conductimétrique ?

Au cours d'un dosage conductimétrique, on mesure la conductivité de la solution à mesure que l'espèce titrante y est introduite. Le critère de détection de l'équivalence est relativement simple : la courbe représentative de la fonction  $\sigma = f(V_Y)$  présente deux portions de droite de pentes différentes. Le point d'intersection de ces droites, où se produit la rupture de pente, marque l'équivalence.



L'interprétation de ce phénomène présente quelques subtilités, variables selon que les espèces titrée et/ou titrante sont/est électriquement chargée(s) ou non. Nous allons ici traiter l'exemple du dosage d'une solution aqueuse d'acide chlorhydrique par une solution aqueuse de soude (celui-ci pourrait d'ailleurs faire également l'objet d'un suivi pH-métrique).

La réaction de dosage est alors simplement :



1. Pour un stade quelconque du dosage précédent l'équivalence, où la soude apportée a été introduite en quantité  $n_{Y,v}$  les espèces présentes dans le milieu réactionnel sont :

- les ions  $\text{Cl}^-$ , présents tout du long en quantité  $n_{X,i}$  puisqu'ils sont spectateurs ;
- les ions  $\text{H}_3\text{O}^+$ , présents en quantité  $n_{X,i} - n_{Y,v}$  ;
- les ions  $\text{Na}^+$ , présents en quantité  $n_{Y,v}$ .

Les ions  $\text{HO}^-$ , limitants à ce stade du dosage, ont évidemment été entièrement consommés par la réaction de dosage.

On constate alors que, pour tout ion  $\text{H}_3\text{O}^+$  consommé, un ion  $\text{Na}^+$  est apporté. Du strict point de vue de la conductivité, tout se passe donc comme si l'on remplaçait progressivement les ions  $\text{H}_3\text{O}^+$  par des ions  $\text{Na}^+$ . Or les premiers sont parmi les meilleurs conducteurs du courant électrique que l'on puisse trouver en solution aqueuse ( $349,8 \cdot 10^{-4} \text{ S.m}^2.\text{mol}^{-1}$  à  $25^\circ\text{C}$ ), ce qui est loin d'être le cas des ions  $\text{Na}^+$  ( $50,1 \cdot 10^{-4} \text{ S.m}^2.\text{mol}^{-1}$  à  $25^\circ\text{C}$ ).

La conductivité du milieu décroît donc progressivement.

2. À l'équivalence, plus aucun ion  $\text{H}_3\text{O}^+$  ou  $\text{HO}^-$  n'est présent dans le milieu. Tout se passe comme si tous les ions  $\text{H}_3\text{O}^+$  avaient été remplacés par des ions  $\text{Na}^+$ , avec en plus une quantité égale d'ions  $\text{Cl}^-$  : la solution n'est plus que de l'eau salée, dont elle a la conductivité.

3. Pour un stade quelconque du dosage suivant l'équivalence, les espèces présentes dans le milieu réactionnel sont :

- les ions  $\text{Cl}^-$ , toujours présents en quantité  $n_{X,i}$  ;
- les ions  $\text{HO}^-$ , désormais présents en quantité  $n_{Y,v} - n_{X,i}$  ;
- les ions  $\text{Na}^+$ , présents en quantité  $n_{Y,v}$ .

On se met donc cette fois à introduire un surcroît d'ions ( $\text{Na}^+$  et surtout  $\text{HO}^-$ , très bon conducteur avec  $198,6 \cdot 10^{-4} \text{ S.m}^2.\text{mol}^{-1}$  à  $25^\circ\text{C}$ ), qui ne vont donner lieu à aucune réaction et contribueront donc d'autant plus à la conduction du courant électrique dans la solution.

La conductivité électrique du milieu augmente donc progressivement.

Il est à noter que, selon les cas, on pourra trouver des ruptures de pente dans l'autre sens (croissante puis décroissante), ou bien sans changement de signe de la pente, mais on observera toujours ce changement de direction.

**Remarque :** nous ne nous sommes intéressés qu'aux quantités de matière, négligeant le fait qu'à mesure que la solution titrante est ajoutée, le volume total du milieu réactionnel augmente, diminuant du même coup les concentrations et, avec elles la conductivité de la solution. Ce facteur est cependant de peu d'importance au regard des variations provoquées par les diminution/augmentation des quantités de matière d'espèce de conductivité ionique molaire élevée. Par ailleurs, on ajoute généralement un important volume d'eau dans la solution titrée avant de commencer, afin d'amoindrir les variations relatives de volume dues à l'ajout de solution titrante.

## — Comment détecte-t-on l'équivalence au cours d'un dosage d'oxydoréduction ?

Les espèces possédant des propriétés d'oxydoréduction sont parfois colorées, notamment lorsqu'il s'agit d'ions métalliques. Dans ce cas, un critère colorimétrique fait l'affaire.

Pour doser un réducteur, on utilisera souvent une solution aqueuse de permanganate de potassium. Les ions permanganate  $\text{MnO}_4^-$  constituent en effet un excellent oxydant, doté en

solution aqueuse d'une couleur violette très visible, même à des concentrations relativement faibles. On verse alors cette solution titrante et l'on observe sa décoloration sous l'effet de sa consommation, jusqu'à l'équivalence. Celle-ci est détectée à la première goutte pour laquelle la coloration violette persiste.

Pour doser un oxydant, on utilise souvent une solution aqueuse d'iodure de potassium, dans le cadre d'un dosage indirect. Les ions iodure  $I^-$  constituent un bon réducteur, et leur oxydant conjugué est le diiode  $I_2$ , qui confère à la solution, selon sa concentration, une couleur allant du jaune au brun (couleur de la fameuse teinture d'iode utilisée dans les milieux hospitaliers). La procédure détaillée est alors la suivante :

1. On introduit les ions iodure en excès : du diiode se forme et la solution brunit.
2. On dose ensuite le diiode formé par une solution aqueuse de thiosulfate de sodium.
3. Les ions thiosulfate  $S_2O_3^{2-}$  sont incolores en solution aqueuse, ainsi que leur oxydant conjugué, l'ion tétrathionate  $S_4O_6^{2-}$ . La teinte orangée que la solution doit au diiode s'affadit donc progressivement, tirant à la fin sur un jaune très léger.
4. Il est impossible de détecter précisément le stade où ce jaune disparaît complètement pour céder la place à une solution véritablement incolore. On utilise pour pallier ce problème un composé appelé empois d'amidon. Il s'agit d'une espèce présentant d'excellentes propriétés complexantes avec le diiode. Si faible soit la quantité de diiode présente dans un milieu réactionnel, il donne toujours un complexe d'un bleu-noir caractéristique très prononcé, même si les quantités de diiode et d'empois d'amidon introduites sont faibles.
5. Ainsi, lorsque la solution pâlit sérieusement, on introduit quelques gouttes d'empois d'amidon. La solution tourne au bleu-noir, et l'on continue à ajouter la solution titrante goutte à goutte, jusqu'à ce que la solution s'éclaircisse brusquement.

**Remarque :** le couple diiode/ion iodure est également utilisé pour des dosages en retour : on introduit du diiode en excès connu et l'on dose ensuite l'excédent par les ions thiosulfate, toujours en utilisant l'empois d'amidon pour la détection fine de l'équivalence.

Dans le cas où une réaction de dosage de type oxydoréduction ne présente pas une fin directement détectable sur le plan visuel, on peut utiliser d'autres méthodes :

- Un dosage conductimétrique ; mais pour être efficace, les ions engagés doivent avoir des conductivités ioniques molaires très contrastées, ce qui n'est pas toujours le cas.
- Un dosage potentiométrique : comme nous l'avons dit, le pH-mètre est en fait un volt-mètre, qui mesure la différence de potentiel entre la solution dont on souhaite déterminer le pH, et une électrode de référence. Le dosage potentiométrique généralise ce concept aux espèces autres que les ions oxonium. Ainsi, si l'on mesure le potentiel  $E$  d'une solution contenant un réducteur à mesure que l'on dose celui-ci par un oxydant, on obtient une courbe  $E = f(V_Y)$  en tout point similaire à la courbe obtenue lors du suivi pH-métrique de la réaction de dosage d'un acide par une base forte.

De même, si l'on dose un oxydant par une espèce réductrice, on obtient une courbe du même type que la courbe de dosage d'une base par un acide fort. Les méthodes de détermination de l'équivalence sont les mêmes :

- Tracé de l'évolution de  $\left| \frac{dE}{dV_Y} \right|$  et recherche de l'abscisse du maximum de la courbe.
- Méthode des tangentes.
- Utilisation d'un indicateur coloré oxydant/réducteur judicieusement choisi.

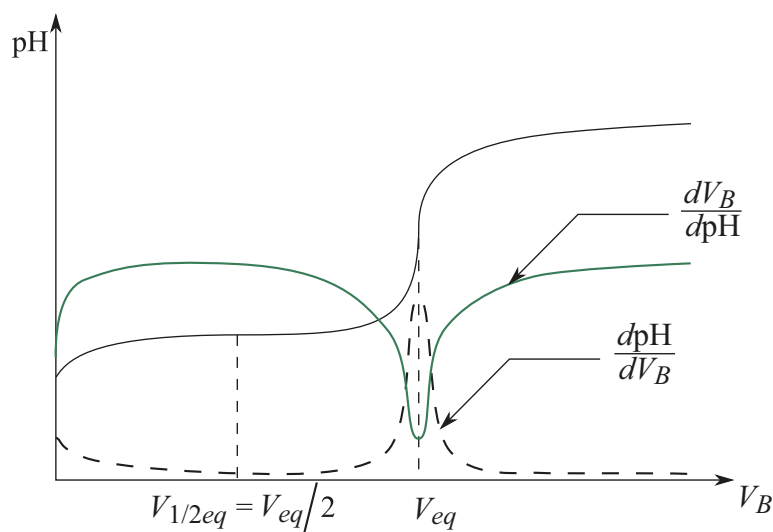
## — Qu'est-ce qu'une solution tampon et comment peut-on déduire ses propriétés d'une courbe de titrage acido-basique ?

Il est intéressant de noter qu'à la demi-équivalence d'un dosage pH-métrique la moitié de la quantité initiale d'acide s'est transformée en sa base conjuguée. Leurs concentrations sont donc égales et leur rapport est égal à 1, d'où l'on peut alors déduire :

|                                          |                                                    |
|------------------------------------------|----------------------------------------------------|
| $\text{pH}_{1/2\text{eq}} = \text{p}K_A$ | $\text{pH}_{1/2\text{eq}}, \text{p}K_A$ sans unité |
|------------------------------------------|----------------------------------------------------|

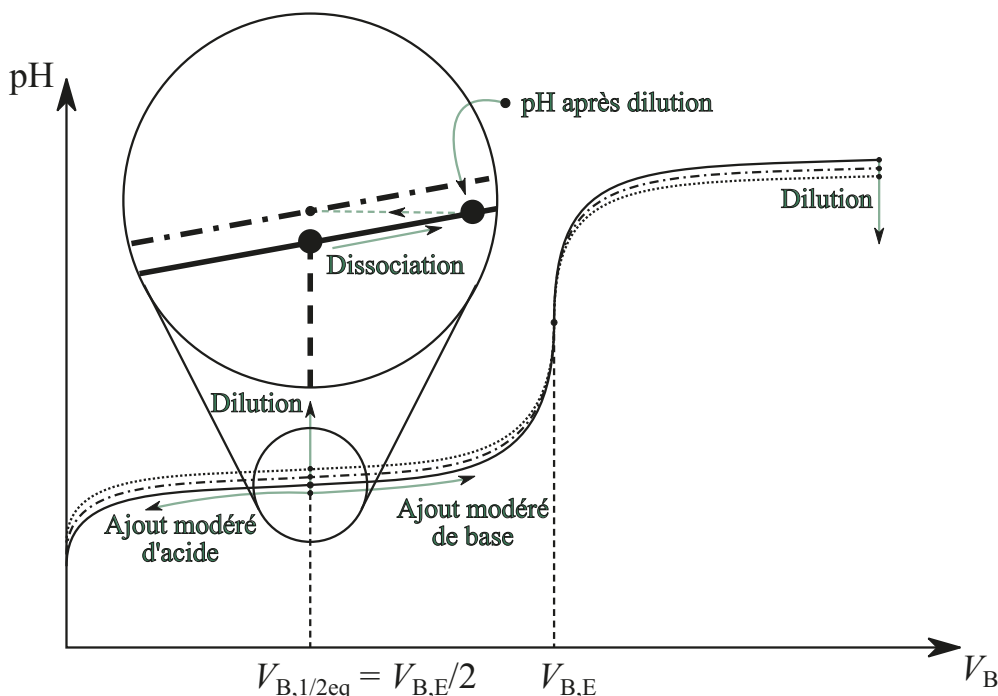
Une solution de ce type présente donc un pH particulièrement stable. On l'appelle **solution tampon** : il s'agit d'une solution dont le pH varie peu par un ajout modéré d'acide ou de base, ainsi que par dilution. On définit le pouvoir tampon  $\beta$  d'une solution comme l'inverse de la valeur absolue de la pente de la tangente à la courbe  $\text{pH} = f(V_{A/B})$ , où  $V_{A/B}$  représente un volume d'acide ou de base introduit.

|                                                      |                |
|------------------------------------------------------|----------------|
| $\beta = \left  \frac{dV_{A/B}}{d\text{pH}} \right $ | $\beta$ en L   |
|                                                      | $V_{A/B}$ en L |
|                                                      | pH sans unité  |



Le pouvoir tampon est donc d'autant plus élevé que la courbe est moins raide, et l'on peut montrer qu'il est maximal à la demi-équivalence, c'est-à-dire en fait pour une solution contenant **un acide et sa base conjuguée en quantités égales**.

Les solutions tampons présentent de nombreuses applications (solutions utilisées pour l'étalonnage des appareils, maintien du pH au sein d'un organe vivant à une valeur constante, et ce malgré l'introduction d'un médicament acide tel que l'aspirine...), et sont très présentes dans la nature, notamment dans la régulation du pH de milieux biologiques, comme le sang.



On remarque également qu'au niveau de la demi-équivalence d'un titrage acidobasique, la courbe représentative  $\text{pH} = f(V_Y)$  passe par un point d'inflexion (dérivée seconde nulle, dérivée première extrême). En d'autres termes, le pH de la solution ne varie jamais aussi peu que dans cette zone, en particulier si (illustration dans l'hypothèse du titrage d'un acide par une base) :

- on y ajoute un peu de base (cas où  $V_B$  augmente modérément) ;
- on y ajoute un peu d'acide (ce qui revient au pire à ajouter des ions oxonium (le plus fort des acides en solution aqueuse), lesquels vont consommer la base conjuguée de l'acide titré, autrement dit tout va se passer comme si nous revenions en arrière sur la courbe de titrage).

Il est intéressant de noter que cette stabilité du pH est également à l'épreuve de la dilution. En effet, que se passe-t-il si nous ajoutons du solvant dans une solution contenant des quantités égales entre elles d'un acide et de sa base conjuguée ? Le pH de la solution à l'équilibre (autrement dit  $[\text{H}_3\text{O}^+]_f$ ) est fixé par celles des réactions engageant les ions oxonium, dont la constante d'équilibre est la plus élevée (principe de la réaction prépondérante). Dans le cas présent, il s'agit de la réaction de l'acide titré avec l'eau ( $\text{H}_3\text{O}^+$  n'est engagé ni dans la réaction support du titrage, ni dans celle de l'acide titré avec sa base conjuguée) ; l'équation régissant le pH est donc, sur ce premier d'état d'équilibre :

$$\frac{[\text{A}^-]_{f(1)} \times [\text{H}_3\text{O}^+]_{f(1)}}{[\text{AH}]_{f(1)}} = K_A$$

Si nous diluons la solution, donc, le volume global va augmenter et ce faisant entraîner une diminution des concentrations de toutes les espèces présentes d'un même facteur (le facteur

de dilution  $D$ ). Le quotient de réaction, dont la valeur était fixée à  $Q_{r,f} = K_A$  à l'équilibre, va donc changer de valeur :

$$Q_r' = \frac{\frac{[A^-]_{f(1)}}{D} \times \frac{[H_3O^+]_{f(1)}}{D}}{\frac{[AH]_{f(1)}}{D}} = \frac{K_A}{D}$$

Le système se retrouve donc hors équilibre, avec un quotient de réaction inférieur à sa valeur d'équilibre, et la réaction va se faire dans le sens direct : l'acide va se dissocier davantage, entraînant donc une baisse de la quantité de matière de l'acide AH, corrélée à une augmentation des quantités de matière de base  $A^-$  et d'ions oxonium  $H_3O^+$ .

Concernant la concentration en acide  $[AH]$ , l'effet est évident : une baisse de la quantité de matière (réaction dans le sens direct) conjuguée à une augmentation du volume de la solution (dilution) font que  $[AH]$  ne peut que baisser.

Concernant les concentrations des produits, en revanche, et en particulier  $[H_3O^+]$ , la réponse est moins évidente : certes l'augmentation de volume pousse à dire que les concentrations vont baisser, mais l'augmentation de la quantité de matière pousse à dire le contraire.

Il suffit en fait de considérer le ratio des quantités de matière de la base conjuguée et de l'acide : le système ayant réagi en accentuant la réaction dans le sens direct, on a forcément, une fois atteint le nouvel équilibre :

$$\frac{n_{A^-,f(2)}}{n_{AH,f(2)}} \text{ qui a augmenté}$$

donc, en divisant en haut et en bas par le volume de la solution,

$$\frac{[A^-]_{f(2)}}{[AH]_{f(2)}} \text{ qui a augmenté également.}$$

Pour que le quotient de réaction retrouve sa valeur d'équilibre, le système va donc forcément réagir en faisant baisser  $[H_3O^+]_f$  par rapport à sa valeur avant dilution : certes la quantité d'ions oxonium aura augmenté, mais d'un facteur qui ne suffit pas à compenser l'augmentation de volume, et l'on a donc  $[H_3O^+]_{f(2)} < [H_3O^+]_{f(1)}$ .

Finalement, donc, la dilution ayant fait baisser  $[H_3O^+]_f$ , le pH aura augmenté. Nous aurions

pu donner cette conclusion dès constatation que  $\frac{[A^-]_f}{[AH]_f}$  avait augmenté, puisque ceci revient

à s'être déplacé vers la droite de la courbe de titrage.

Et comme dans le cas de l'ajout modéré d'acide ou de base, nous constatons que se trouver au voisinage de la demi-équivalence fait que la variation de pH résultant de cette perturbation de l'équilibre entre l'acide et sa base conjuguée ne va entraîner qu'une très faible variation du pH (très faible taux de croissance de la courbe à ce niveau de la courbe).

## Halte aux idées reçues

1. L'équivalence d'un titrage est le stade de ce titrage où la solution change de couleur.
2. Avec une solution de base de concentration connue et un suivi pH-métrique, il est toujours possible de titrer un acide par une base.
3. À l'équivalence d'un titrage acidobasique mené à 25 °C, le pH du milieu réactionnel vaut toujours 7.

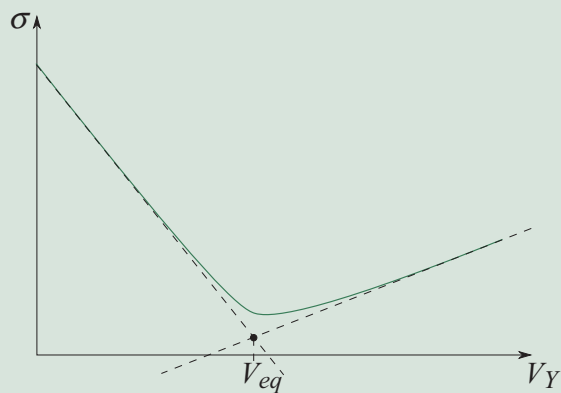
4. Dans un titrage conductimétrique, le thermomètre est là car la composition du milieu réactionnel à l'équilibre dépend de la température.

5. Dans un titrage, lorsque seule l'espèce titrante est colorée, la persistance de cette couleur dans le milieu réactionnel signe l'atteinte de l'équivalence.

## Du Tac au Tac

(Exercices sans calculatrice)

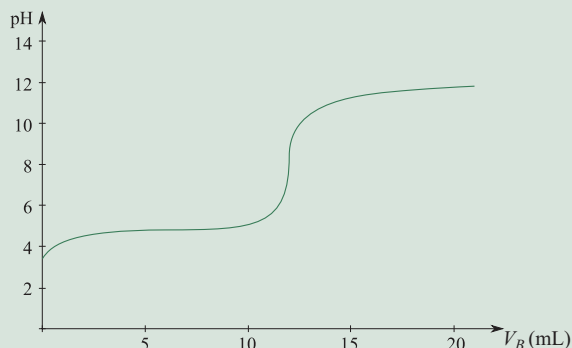
1. On réalise le suivi conductimétrique du titrage d'une solution aqueuse d'acide chlorhydrique ( $\text{H}_3\text{O}^+_{(\text{aq})}, \text{Cl}^-_{(\text{aq})}$ ) de concentration apportée en soluté  $c_1 = 1,0 \cdot 10^{-3} \text{ mol.L}^{-1}$ , par une solution aqueuse de soude ( $\text{Na}^+_{(\text{aq})}, \text{HO}^-_{(\text{aq})}$ ) de concentration apportée en soluté  $c_2 = 1,0 \cdot 10^{-1} \text{ mol.L}^{-1}$ . On obtient la courbe suivante :



On donne les conductivités ioniques molaires suivantes :  $\lambda_{\text{H}_3\text{O}^+} = 34,98 \text{ mS.m}^2.\text{mol}^{-1}$ ,  $\lambda_{\text{HO}^-} = 19,86 \text{ mS.m}^2.\text{mol}^{-1}$ ,  $\lambda_{\text{Na}^+} = 5,01 \text{ mS.m}^2.\text{mol}^{-1}$ ,  $\lambda_{\text{Cl}^-} = 7,63 \text{ mS.m}^2.\text{mol}^{-1}$ .

Interpréter l'évolution de la conductivité de la solution.

2. On considère le titrage d'un acide faible, par une solution aqueuse de soude. On obtient la courbe de titrage pH-métrique suivante :



où le pH indiqué est celui du mélange réalisé, et  $V_B$  désigne le volume de solution aqueuse de soude versé.

Sachant que l'on va devoir réitérer ce titrage à de nombreuses reprises, on souhaite utiliser un indicateur coloré acido-basique pour se dispenser du suivi pH-métrique les fois suivantes.

On dispose des indicateurs suivants :

|                         |          |     |                |     |              |
|-------------------------|----------|-----|----------------|-----|--------------|
| <b>Hélianthine</b>      | Rouge    | 3,1 | Zone de virage | 4,4 | Jaune orangé |
| <b>Rouge de méthyle</b> | Rouge    | 4,2 | Zone de virage | 6,2 | Jaune        |
| <b>BBT</b>              | Jaune    | 6,0 | Zone de virage | 7,6 | Bleu         |
| <b>Rouge de phénol</b>  | Jaune    | 6,8 | Zone de virage | 8,4 | Rouge        |
| <b>Phénolphthaléine</b> | Incolore | 8,0 | Zone de virage | 9,9 | Rose         |

Approximer les intervalles de valeurs de  $V_B$  sur lesquels vont virer ces indicateurs colorés, et en déduire lequel est le plus à même d'offrir une mesure précise du volume versé à l'équivalence.

3. On souhaite titrer les ions permanganate  $\text{MnO}_4^-$  (couple  $\text{MnO}_4^-/\text{Mn}^{2+}$ ) d'une solution aqueuse par les ions fer (II)  $\text{Fe}^{2+}$  (couple oxydant/réducteur  $\text{Fe}^{3+}/\text{Fe}^{2+}$ ) (réaction totale et quasi-instantanée).

En milieu acide, les ions manganèse (II) réagissent lentement avec les ions permanganate pour former un précipité marron de dioxyde de manganèse  $\text{MnO}_2$ .

Expliquer pourquoi il n'est pas possible de mener ce titrage selon le protocole habituel, et proposer une alternative en indiquant le critère de détection de l'équivalence.

Espèces colorées en solution aqueuse :

- $\text{MnO}_4^-$  : violet ;
- $\text{Fe}^{2+}$  : vert très pâle (quasi incolore) ;
- $\text{Fe}^{3+}$  : jaune léger.

4. La dureté d'une eau est définie comme la somme des concentrations de cette eau en ions calcium (II) et magnésium (II) (tous deux incolores en solution aqueuse). Pour la déterminer, on procède couramment par titrage complexométrique utilisant :

- une solution aqueuse d'un ligand appelé ion éthylène-diaminetétraacétate (EDTA, noté ci-après  $\text{Y}^{4-}$ ), présentant les propriétés suivantes :
  - il est incolore en solution aqueuse lorsqu'il n'est pas engagé dans un complexe,
  - il forme avec les ions  $\text{Ca}^{2+}$  ainsi qu'avec les ions  $\text{Mg}^{2+}$  les complexes  $[\text{CaY}]^{2-}$  et  $[\text{MgY}]^{2-}$ , tous deux incolores,

- en présence simultanée d'ions calcium et magnésium, l'EDTA forme le complexe  $[\text{CaY}]^{2-}$  prioritairement sur  $[\text{MgY}]^{2-}$  ;

- une solution de noir eriochrome T (NET), autre ligand, présentant les propriétés suivantes :

- le NET est bleu en solution aqueuse lorsqu'il n'est pas engagé dans un complexe,
- il forme, avec les ions  $\text{Mg}^{2+}$  seulement, un complexe  $[\text{MgNET}]^{2+}$  de couleur violette, moins stable que le complexe  $[\text{MgY}]^{2-}$ .

Proposer sur cette base un protocole de titrage colorimétrique de la dureté d'une eau.

5. Le titrage des ions chlorure contenus dans une solution par la méthode de Mohr s'appuie sur deux réactions distinctes :

- La précipitation des ions argent (I)  $\text{Ag}^+$  avec les ions chlorure  $\text{Cl}^-$  pour former un précipité de chlorure d'argent (solide blanc noirissant progressivement à la lumière).
- La précipitation des ions argent (I)  $\text{Ag}^+$  avec les ions chromate  $\text{CrO}_4^{2-}$  pour former un précipité de chromate d'argent (solide rouge).

La première est prioritaire sur la seconde, c'est-à-dire qu'elle ne peut se dérouler que si la première est terminée.

Établir les équations des deux réactions évoquées ci-dessus, puis proposer un protocole de titrage des ions chlorure présents dans une solution, s'appuyant sur ces réactions.

## Vers la prépa

On se propose de doser l'acide ascorbique contenu dans un cachet de vitamine C500. L'acide ascorbique s'oxyde très facilement au contact du dioxygène de l'air, aussi ne peut-on se permettre de prendre son temps pour réaliser l'expérience.

C'est pourquoi on réalise un dosage en retour : on broie un cachet dans un mortier, et on le dissout dans l'eau de sorte à obtenir un volume  $V_0 = 100,0$  mL de solution. On prélève ensuite un volume  $V_1 = 10,0$  mL de cette solution, que l'on introduit dans un becher contenant un volume  $V_2 = 25,0$  mL de solution aqueuse de diiode à la concentration  $c_2 = 2,00 \cdot 10^{-2}$  mol.L<sup>-1</sup>, afin d'oxyder immédiatement tout l'acide ascorbique.

On dose ensuite le diiode restant par une solution aqueuse de thiosulfate de sodium à la concentration  $c_3 = 2,50 \cdot 10^{-2}$  mol.L<sup>-1</sup>.

On donne les couples oxydant-réducteur suivants :  $\text{C}_6\text{H}_6\text{O}_6/\text{C}_6\text{H}_8\text{O}_6$ ,  $\text{I}_2/\text{I}^-$  et  $\text{S}_4\text{O}_6^{2-}/\text{S}_2\text{O}_3^{2-}$  (couple tétrathionate/thiosulfate).

1. Équilibrer l'équation de la réaction de l'acide ascorbique avec le diiode.

2. Équilibrer l'équation de la réaction des ions thiosulfate avec le diiode restant à l'issue de la première réaction.

3. Sachant que l'on trouve un volume de solution aqueuse de thiosulfate de sodium versé à l'équivalence  $V_{3,E} = 17,4$  mL, déterminer la masse d'acide ascorbique présente dans le cachet.

4. Que signifie selon vous la valeur « 500 » précisée dans le nom du cachet ?

## Halte aux idées reçues

**1.** Cette définition de l'équivalence possède deux défauts majeurs : premièrement elle propose une interprétation de l'équivalence qui ne met absolument pas en avant le changement de réactif limitant ou quelque considération sur les quantités de matière des espèces engagées qui la rendrait exploitable. Ensuite, pour peu qu'aucun des réactifs ni aucun des produits ne soit coloré en solution aqueuse (ce qui est le cas la plupart du temps), elle est hors de propos. Il est vrai que l'on peut souvent suppléer à cette absence de couleur par l'adjonction de quelques gouttes d'un indicateur coloré judicieusement choisi. Mais il est inenvisageable d'aborder une CPGE en se limitant à une approche aussi superficielle des choses.

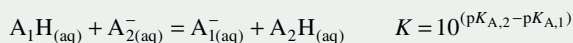
On retiendra donc que l'équivalence d'un titrage est le stade de ce titrage où le réactif titrant a été introduit en proportions stœchiométriques avec la quantité initiale de réactif titré.

À ce stade, le réactif limitant dans le milieu réactionnel change de nature : réactif titrant avant l'équivalence, réactif titré après. On peut également retenir cette propriété comme définition de l'équivalence, et présenter la version précédente comme une conséquence de cette autre définition.

Il importe surtout de retenir que ce changement de réactif limitant signifie également un changement des espèces présentes en solution, et donc changement des propriétés de cette solution. Ces modifications vont opérer un brusque changement de visage de la solution, changement qui servira de base à la détection de l'équivalence.

**2.** Nous savons que pour pouvoir servir de réaction de titrage, une réaction doit remplir 4 conditions, parmi lesquelles : être totale et être de fin détectable. Or ces deux points peuvent être mis en défaut si l'on choisit mal l'espèce titrante.

Concernant le caractère quantitatif de la réaction : on sait que la réaction de l'acide  $A_1H$  d'un couple de  $pK_A$  égal à  $pK_{A,1}$ , sur la base  $A_2^-$  d'un couple de  $pK_A$  égal à  $pK_{A,2}$ , a pour valeur :



On sait encore que pour une réaction engageant en tout et pour tout 2 réactifs et 2 produits, tous solvatés et intervenant avec le même nombre stœchiométrique (le cas présentement), la réaction peut être considérée comme totale si sa constante d'équilibre est supérieure à  $10^4$ . En d'autres termes, la base choisie doit appartenir à un couple dont le  $pK_A$  est supérieur d'au moins 4 unités de pH au  $pK_A$  du couple avec l'acide duquel elle réagit.

Par ailleurs, la détection de l'équivalence d'un titrage acido-basique repose sur le saut de pH provoqué par le changement de réactif limitant, saut dont l'amplitude en unité de pH est approximativement égal à  $pK_{A,2} - pK_{A,1}$ . On retrouve donc la nécessité, pour disposer d'une détection bien tranchée, d'un écart significatif entre les deux  $pK_A$ .

Cet écart fait que l'on opte en général pour une base forte : l'eau nivelant les bases fortes en ions hydroxyde  $HO^-$  (la plus forte des bases faibles dans l'eau  $pK_A(H_2O/HO^-) = 14$  à  $25^\circ C$ ), on met toutes les chances de son côté.

**3.** Lorsque l'on atteint l'équivalence d'un titrage acidobasique (titrage d'un monoacide par une base forte, par exemple), ont été introduits dans le milieu réactionnel :

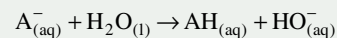
- une certaine quantité d'acide ;
- une quantité égale de base forte.

La réaction prépondérante, entre ces deux espèces, donne donc :

- une quantité égale de la base conjuguée de l'acide de départ ;
- une quantité égale d'eau (acide conjugué de l'ion hydroxyde, auquel se trouvent nivelées les bases fortes dans l'eau), négligeable devant la quantité de solvant.

**Remarque :** pour rappel, l'eau représente  $56 \text{ mol.L}^{-1}$ , donc dès lors qu'une solution est le siège d'une réaction s'appuyant sur des espèces dont les concentrations sont significativement inférieures à  $1 \text{ mol.L}^{-1}$ , toute consommation ou formation d'eau par cette réaction sera négligeable devant la quantité initiale de solvant.

Tout se passe donc comme si l'on avait introduit la base conjuguée seule, dans l'eau. La deuxième réaction qui va avoir lieu sera donc l'équilibre de cette base avec l'eau, aboutissant à un équilibre régi par le  $pK_A$  de son couple :



$$K = K_B = \frac{K_e}{K_A} = \frac{[AH]_f \times [HO^-]_f}{[A^-]_f}$$

Si nous supposons la quantité d'ions hydroxydes formée par l'autoprotolyse de l'eau, négligeable devant celle issue de l'équilibre ci-dessus, nous pouvons écrire :

$$[AH]_f = [HO^-]_f \left( = \frac{x_f}{V_{\text{tot}}} \right)$$

où  $V_{\text{tot}}$  désigne le volume total du milieu réactionnel, soit donc le volume de solution titrée, plus le volume de solution titrante versé à l'équivalence.

Par ailleurs, nous avons également :

$$[A^-]_f = \frac{C_A \times V_A - x_f}{V_{\text{tot}}}$$

En ajoutant les hypothèses que :

- la base réagit peu avec l'eau ( $x_f \ll C_A \times V_A$ ) ;
- la solution titrante est beaucoup plus concentrée que la solution titrée ( $V_{B,E} \ll V_A$ , donc  $V_{\text{tot}} = V_A + V_{B,E} \approx V_A$ ) cette expression devient :  $[A^-]_f \approx C_A$  d'où finalement :

$$\begin{aligned} \frac{K_e}{K_A} &= \frac{1}{C_A} \times [HO^-]_f^2 \Leftrightarrow \frac{K_e}{K_A} \\ &= \frac{1}{C_A} \times \left( \frac{K_e}{[H_3O^+]_f} \right)^2 \Leftrightarrow [H_3O^+]_f \\ &= \sqrt{\frac{K_A \times K_e}{C_A}} \end{aligned}$$

En prenant  $-\log$  des membres de gauche et de droite, on trouve alors :

$$pH_f = \frac{1}{2}(pK_e + pK_A - \log C_A)$$

Avec  $pK_e = 14$ , on obtient finalement :

$$pH_f = 7 + \frac{1}{2}(pK_A - \log C_A)$$

Le pH à l'équivalence ne prend en réalité la valeur  $pH_E = 7$ , que si l'on titre un acide fort par une base forte : dès lors qu'il s'agit d'un acide faible par une base forte, la valeur du pH à l'équivalence dépend non seulement de la nature de l'acide titré, mais également de sa concentration initiale dans le milieu réactionnel.

**4.** Il est exact de dire que la composition du milieu réactionnel à l'équilibre dépend de la température. Rappelons cependant que la réaction support d'un titrage se doit d'être **totale** : le concept d'équilibre chimique est donc en soi hors de propos dans ce contexte.

Il est cependant vrai que les réactions véritablement totales sont en réalité rares, et que l'on s'appuie souvent sur des réactions simplement quantitatives, c'est-à-dire aboutissant à des équilibres dont les constantes sont très élevées, entraînant *de facto* des concentrations résiduelles négligeables en l'un au moins des réactifs qu'elles engagent. La composition du système aura beau varier si la température change, cette variation ne mènera en général pas à devoir prendre en compte une concentration qui se trouvait négligeable quelques kelvins plus haut ou plus bas.

Si d'ailleurs la chose se produisait, la réaction deviendrait aussitôt inéligible en tant que réaction de titrage, le fait

d'aboutir à un équilibre médian empêchant d'atteindre un semblant d'équivalence.

Reste la question de savoir à quoi sert ce thermomètre : il sert à vérifier non la valeur, mais simplement la **constance** de la température. Les conductivités ioniques molaires sont en effets très sensibles à celle-ci, et la courbe de suivi conductimétrique est exploitable uniquement si elle est constituée de deux portions de droite.

Le caractère affine de l'évolution de la conductivité en fonction du volume de solution titrante versé suppose que les coefficients de la loi de Kohlrausch soient constants, or ceux-ci reposent justement sur les conductivités ioniques molaires. D'où la nécessité de s'assurer que la température ne varie pas, afin d'obtenir des données cohérentes entre elles et donc à même d'être exploitées.

**5.** Cette affirmation semble en soi plutôt sensée : si seule l'espèce titrante est colorée, cette couleur ne pourra persister dans le milieu réactionnel avant l'équivalence (lorsque l'espèce titrante constitue le réactif limitant), et en revanche s'imposera après l'équivalence (lorsque l'espèce titrante sera en excès).

Cependant tout est dans cet « après » : si la couleur est visible, c'est que l'espèce titrante est présente. Et si l'espèce titrante est présente c'est que l'on a **dépassé** l'équivalence. Donc en toute rigueur, lorsque l'équivalence approche (donc lorsque l'ajout de solution titrante commence à mettre beaucoup de temps à se décolorer, témoignant de la raréfaction de l'espèce titrée), il faudrait :

- noter le volume avant chaque nouvelle goutte versée ;
- ajouter la goutte ;
- en cas de persistance de la coloration, retenir non le volume lisible sur la burette à ce stade, mais le dernier à avoir précédé ladite persistance.

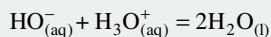
Ceci souligne une nouvelle fois ce point trop souvent oublié : à l'équivalence, le milieu réactionnel ne doit contenir ni espèce titrante, ni espèce titrée.

## Du Tac au Tac

**1.** Avant toute chose, une remarque s'impose concernant les valeurs des concentrations fournies : la solution titrante est 100 fois plus concentrée, que la solution titrée. Sachant en outre que les réactifs réagissent mole à mole, on peut aisément montrer que le volume de solution titrante sera égal à un centième du volume de solution titrée traité. Par conséquent l'effet de dilution consécutif à l'ajout de solution titrante, qui tendrait en toute rigueur à faire baisser la conductivité globale de la solution, aura finalement très peu d'impact puisqu'engageant une variation seulement de l'ordre du pourcent, des concentrations des espèces considérées. Cette façon de faire est très fréquente dans les titrages conductimétriques, la prise en compte des variations de volume de la solution de travail compliquant significativement l'exploitation des résultats.

Concernant la courbe proprement dite, nous constatons que dans la première partie la conductivité diminue assez

rapidement. Analysons ce qu'il se passe concrètement dans cette première phase du titrage : on ajoute une solution aqueuse de soude ( $\text{Na}^+_{(\text{aq})}, \text{HO}^-_{(\text{aq})}$ ) à la solution d'acide chlorhydrique ( $\text{H}_3\text{O}^+_{(\text{aq})}, \text{Cl}^-_{(\text{aq})}$ ). Ceci entraîne une réaction acido-basique entre les ions hydroxyde et les ions oxonium, selon l'équation :



Une première idée pour expliquer la chute de la conductivité de la solution, une fois exclu l'effet de dilution mentionné au début, serait la consommation des ions oxonium pour ne générer que de l'eau, espèce non chargée et donc non conductrice du courant électrique. Ce serait cependant négliger le fait que pour tout ion hydroxyde introduit, un contre-ion sodium est également introduit. Ainsi la consommation des ions oxonium ne se fait-elle pas sans contre-partie, puisqu'ils sont remplacés par une quantité égale d'ions sodium. Tout se passe donc au final comme si l'on remplaçait progressivement les ions oxonium par des ions sodium. Cependant, si l'on compare les conductivités ioniques molaires des ions oxonium et des ions sodium, on constate que les premiers conduisent le courant électrique presque sept fois mieux que les seconds (il peut être utile de garder à l'esprit que les ions oxonium sont les meilleurs conducteurs que l'on puisse trouver en solution aqueuse, d'ailleurs immédiatement suivis par les ions hydroxyde). Ce remplacement par des ions moins bons conducteurs explique ainsi cette chute rapide de la conductivité de la solution.

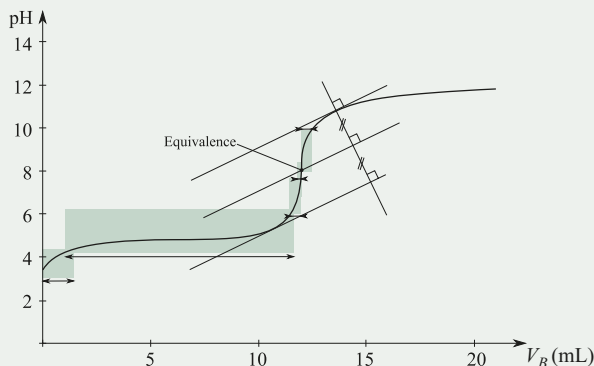
Dans un deuxième temps, on constate que la conductivité augmente. Le phénomène de dilution a, ici, encore moins droit de cité que dans la première partie, puisque s'il a un effet, ce ne peut être que de faire baisser la conductivité. Il suffit de voir que dans cette deuxième partie, l'équivalence a déjà été atteinte. En l'absence d'ions oxonium, les ions hydroxyde introduits (et avec eux les ions sodium) s'accumulent, sans plus pouvoir réagir avec quoi que ce soit, provoquant donc l'augmentation de la conductivité de la solution.

**2.** Le problème est ici de choisir un indicateur coloré dont la zone de virage soit incluse dans le saut de pH qui accompagne l'équivalence : s'il se met à virer pour un pH se trouvant encore sur la partie plate précédant l'équivalence, le changement de couleur va s'étirer sur un intervalle très large de valeurs de  $V_B$ , et la connaissance du volume versé à l'équivalence sera très approximative. À l'autre extrémité, si la fin de la zone de virage se situe à un pH au-dessus du saut de pH, la fin du changement de couleur va s'étirer de la même façon, mais après l'équivalence.

Nous trouvons par la méthode des tangentes que l'équivalence a lieu pour un volume de base versé, de  $V_{B,E} = 12,0 \text{ mL}$ , et que le pH prend alors pour valeur  $\text{pH}_E = 8,0$ .

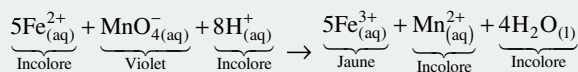
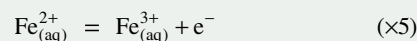
La zone de virage de l'indicateur coloré choisi devra donc inclure cette valeur de pH de 8,0, et un intervalle de valeurs correspondantes de  $V_B$  aussi serré que possible, donc une zone de virage aussi restreinte que possible en valeur de pH.

Avec les indicateurs colorés proposés, nous constatons que les intervalles de valeurs de  $V_B$  sur lesquels se fait le changement de couleur sont respectivement :



- Pour l'hélianthine : entre 0 et 1,4 mL, soit un intervalle de 1,4 mL. Outre le fait que cette incertitude est supérieure aux valeurs communément exigées pour un titrage (de l'ordre du dixième de mL, typiquement), on constate que le changement de couleur est terminé alors que le saut de pH n'est même pas commencé. Cet indicateur est donc à exclure.
- Pour le rouge de méthyle : entre 1,0 et 11,6 mL, soit un intervalle absurde de 10,6 mL. On constate en outre une nouvelle fois que le changement de couleur cesse avant même que ne soit atteinte l'équivalence. Cet indicateur est donc à exclure également.
- Pour le bleu de bromothymol (BBT) : entre 11,4 mL et 11,9 mL, soit un intervalle de 0,5 mL. L'intervalle, sans être idéal, commence à devenir raisonnable. Mais les intervalles de pH autant que de  $V_B$  excluent encore le point d'équivalence, et le BBT ne peut pas non plus être retenu.
- Pour le rouge de phénol : entre 11,8 mL et 12,0 mL, soit un intervalle de 0,2 mL. Cet intervalle n'est pas encore le meilleur que l'on puisse espérer, puisque les burettes graduées sont précises à 0,1 mL près. Il est cependant acceptable, et a le mérite d'inclure le point d'équivalence. Le rouge de phénol peut donc être retenu pour ce titrage.
- Pour la phénolphthaléine : entre 12,0 mL et 12,5 mL. L'intervalle recommence à croître, mais comprend malgré tout le point d'équivalence. Il est en soi également éligible, mais l'encadrement de  $V_B$  étant moins bon que celui proposé par le rouge de phénol, on lui préférera ce dernier.

**3.** D'après l'énoncé, on souhaite titrer les ions  $\text{Fe}^{2+}_{(\text{aq})}$  par les ions  $\text{MnO}_4^-_{(\text{aq})}$ . Ce faisant, nous allons former les espèces conjuguées de ces deux réactifs, soit  $\text{Fe}^{3+}_{(\text{aq})}$  et  $\text{Mn}^{2+}_{(\text{aq})}$ , selon la réaction d'équation :



En suivant la méthode classique de titrage, nous placerions la solution contenant les ions  $\text{MnO}_4^-$  dans un erlenmeyer, acidifierions le milieu (besoin d'ions  $\text{H}^+$ ) et, sous agitation, verserions progressivement une solution aqueuse d'ions  $\text{Fe}^{2+}$  de concentration connue, au moyen d'une burette graduée.

La solution titrée, violette, se décolorerait progressivement à mesure que les ions  $\text{MnO}_4^-$  se feraient consommer, la disparition finale de cette teinte au profit de la seule couleur jaune des ions  $\text{Fe}^{3+}$  signant la fin de la réaction support et donc l'équivalence.

Cependant cette façon de faire pose effectivement problème : le réactif initialement présent dans l'erlenmeyer se trouvant en excès dans le milieu réactionnel durant toute la première phase du titrage (entre le début et l'atteinte de l'équivalence), les ions  $\text{Mn}^{2+}$  formés par la réaction support devraient dans ce cas coexister avec les ions  $\text{MnO}_4^-$  qui n'auront pas encore été consommés, le tout dans un milieu nécessairement acide. Or nous savons qu'une telle coexistence s'accompagne d'une réaction entre ces deux espèces, pour former un précipité de dioxyde de manganèse.

Cette réaction parasite va avoir une double conséquence :

- la formation du précipité marron de dioxyde de manganèse  $\text{MnO}_2(\text{s})$  évoqué par l'énoncé, qui risque de compromettre la détection du changement de couleur ;
- mais surtout, la consommation par l'un des produits de la réaction de titrage (les ions  $\text{Mn}^{2+}$ ), d'une partie du réactif titré, violant ainsi le critère d'unicité que doit vérifier la réaction support pour offrir un titrage valide.

Les ions permanganate seraient ainsi consommés par deux réactions ; la quantité d'ions fer (II) utilisée ne couvrirait donc pas la totalité des ions permanganate initialement présents, menant à une sous-estimation de leur quantité.

On pourrait imaginer de prendre en compte cette surconsommation, cependant il faudrait être sûr que cette seconde réaction est également totale et rapide (or nous savons qu'elle est lente), et resterait le problème de détection de l'équivalence liée à la présence du précipité.

Il est possible de procéder plus simplement : ainsi que nous l'avons dit, le problème vient du fait que les ions permanganate se trouvent en excès, et sont donc à même de réagir avec les produits issus de la réaction de titrage (en l'occurrence avec  $\text{Mn}^{2+}$ , donc). Il nous suffit alors d'intervertir les contenants des deux solutions : si nous plaçons la solution aqueuse d'ions  $\text{Fe}^{2+}$  dans l'erlenmeyer et la solution acidifiée d'ions permanganate dans la burette, les ions  $\text{Mn}^{2+}$  formés par la réaction coexisteront uniquement avec les ions  $\text{Fe}^{2+} + \text{Fe}^{3+}$ , sans réaction parasite possible.

À chaque ajout de solution titrante, les ions permanganate réagiront instantanément (critère de rapidité de la réaction support) avec les ions  $\text{Fe}^{2+}$ , devant le pion aux ions manganèse (II) (réaction lente).

Dans ces conditions, on observera cette fois la solution violette se décolorant au contact de la solution contenue dans l'erlenmeyer (consommation des ions permanganate),

laquelle se colorera progressivement en jaune (formation d'ions  $\text{Fe}^{3+}$ ). L'équivalence sera alors signée par la persistance de la coloration violette dans l'erlenmeyer, témoignant du changement de réactif limitant.

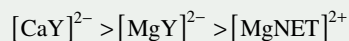
4. La première chose à bien noter dans cet exercice est le fait que nous cherchons à déterminer une concentration globale en ions calcium ET magnésium : peu nous importe donc de faire la différence entre les concentrations individuelles de l'un et de l'autre.

L'énoncé nous informe que ces deux ions sont incolores en solution aqueuse, et que l'EDTA forme avec chacun d'eux un complexe également incolore (en premier lieu avec les ions calcium, puis s'il ne reste aucun de ceux-ci à complexer, avec les ions magnésium).

L'EDTA seule ne peut guère nous aider, puisque les réactions auxquelles elle donne lieu avec les ions qui nous intéressent ne produisent aucun changement colorimétrique visible dans le milieu réactionnel.

Examinons alors ce que nous pouvons tirer du NET : celui-ci est bleu tant qu'il n'est pas engagé dans un complexe, et forme avec les seuls ions magnésium (II) un complexe violet, qui cependant est moins stable que celui formé par l'EDTA avec ces mêmes ions.

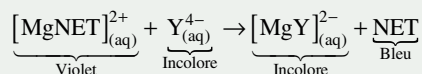
À ce stade, nous avons donc identifié trois complexes dont nous connaissons les priorités de formation :



Nous en déduisons que nous ne tirerons aucun bénéfice du NET s'il est introduit après l'EDTA : il ne donnera rien avec les ions calcium, et ne pourra pas réagir avec ceux des ions magnésium qui auront déjà été complexés par l'EDTA.

Si en revanche nous introduisons du NET en premier, celui-ci va complexer une partie des ions magnésium, formant le complexe  $[\text{MgNET}]^{2+}$  qui confèrera à la solution sa couleur violette. En introduisant ensuite une solution d'EDTA, celle-ci va réagir avec les différentes espèces présentes en suivant scrupuleusement l'ordre de priorité des complexes, autrement dit du plus stable au moins stable. Seraient ainsi complexés, dans l'ordre :

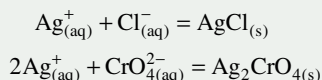
- les ions calcium  $\text{Ca}^{2+}$  ;
- les ions magnésium libres  $\text{Mg}^{2+}$  (plus disponibles que ceux complexés par le NET, puisqu'aucune dissociation n'est requise) ;
- et finalement les ions magnésium engagés dans le complexe  $[\text{MgNET}]^{2+}$ , moins stable que  $[\text{MgY}]^{2-}$ . Ce faisant,  $[\text{MgNET}]^{2+}$  va poliment se dissocier au profit de  $[\text{MgY}]^{2-}$  selon une réaction que l'on pourrait résumer par l'équation :



En conclusion, nous allons pouvoir titrer l'ensemble des ions calcium et magnésium par l'EDTA, en usant du NET comme indicateur coloré, selon le protocole suivant :

1. On introduit quelques gouttes de NET dans la solution titrée, qui complexe des ions magnésium et confère à la solution une couleur violette.
2. On introduit une solution d'EDTA de concentration connue ; celle-ci réagit successivement avec les ions calcium (aucun changement apparent), les ions magnésium libres (aucun changement apparent), et finalement dissocie le complexe  $[\text{MgNET}]_{(\text{aq})}^{2+}$ , dont la disparition finale sera marquée par le changement de couleur de la solution, qui va virer du violet au bleu, marquant ainsi l'équivalence du titrage.
3. On détermine la quantité de matière d'EDTA nécessaire à l'atteinte de l'équivalence ; l'EDTA réagissant mole à mole avec les ions calcium et magnésium, cette quantité de matière sera égale la somme des quantités de matière de ces deux ions, qui rapportées au volume d'eau titré donneront la somme des concentrations correspondantes.

5. Les précipités étant des solides, ils doivent être électriquement neutres. Nous en déduisons, compte tenu des charges des différents ions considérés, les stœchiométries suivantes :



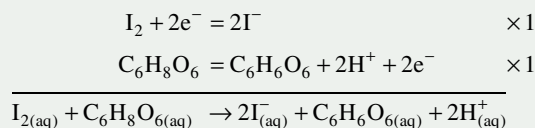
La seule espèce à réagir avec les ions chlorure est l'ion argent. On peut donc imaginer que ce dernier constituera l'espèce titrante. Reste que, le précipité se formant en volume dans la solution, il est impossible de dire avec précision quand sa formation prend fin. On constate par ailleurs que les ions argent peuvent également former avec les ions chromate un précipité rouge, mais seulement si la réaction précédente ne peut plus se faire, c'est-à-dire dans notre cas s'il n'y a plus en solution d'ions chlorure pour réagir avec les ions argent. On peut dès lors envisager d'utiliser les ions chromate comme indicateurs de fin de réaction, et proposer le protocole suivant :

- (a) Introduire dans un erlenmeyer (titrage colorimétrique) un volume  $V_0$  de solution aqueuse contenant les ions chlorure à titrer.  $V_0$  sera mesuré à la pipette jaugée si possible, graduée sinon, afin de pouvoir calculer avec une précision optimale la concentration des ions chlorure lorsque nous rapporterons leur quantité de matière à  $V_0$ .
- (b) Ajouter quelques mL de solution aqueuse de chromate de potassium. À ce stade, aucun ion argent ne se trouve dans le milieu réactionnel et rien ne se produit.
- (c) Rincer puis remplir une burette graduée (besoin de précision sur le volume de solution titrante versé, pour la répercuter sur la quantité de matière d'ions argent introduite) avec une solution aqueuse de nitrate d'argent de concentration connue.
- (d) Verser cette solution dans l'erlenmeyer par ajouts successifs tout en agitant, jusqu'à ce que le contenu de l'erlenmeyer prenne la teinte rouge du chromate d'argent : on sait alors qu'il n'y a plus d'ions chlorure à titrer.
- (e) Connaissant le volume  $V_{1,\text{E}}$  de solution titrante versé à l'équivalence, nous pouvons déterminer la quantité de

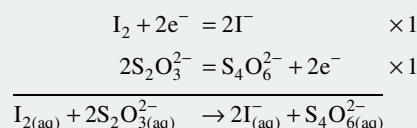
matière d'ions argent introduite, dont nous déduirons la quantité d'ions chlorure initialement présente dans l'erlenmeyer. Il ne restera plus qu'à rapporter cette dernière au volume  $V_0$  pour connaître la concentration en ions chlorure de la solution de départ.

## Vers la prépa

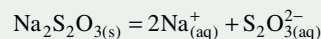
1. Bien que la formule brute de l'acide ascorbique ne soit pas supposée être connue, il était facile de l'identifier. En effet, l'énoncé précisait qu'au cours de la réaction envisagée l'acide ascorbique subissait une oxydation. Il s'agit donc d'un réducteur, donc le second membre du couple  $\text{C}_6\text{H}_8\text{O}_6/\text{C}_6\text{H}_6\text{O}_6$ . Commençons par équilibrer chacune des demi-équations électroniques associées aux couples auxquels appartiennent les réactifs engagés :



2. Écrivons de même les demi-équations des couples engagés dans la deuxième partie du dosage :



3. Le dosage du diiode restant aboutit à l'équivalence pour un volume  $V_{3,\text{E}} = 17,4$  mL. Comme la dissociation du thiosulfate de sodium donne un ion thiosulfate à chaque fois :



nous pouvons écrire que, dans la solution titrante,  $[\text{S}_2\text{O}_3^{2-}]_3 = c_3 = 2,50 \cdot 10^{-2} \text{ mol} \cdot \text{L}^{-1}$ , et en déduire que la quantité d'ions thiosulfate nécessaire pour faire réagir tout le diiode restant à l'issue de la première réaction était donc :

$$n_{\text{S}_2\text{O}_3^{2-},\text{E}} = [\text{S}_2\text{O}_3^{2-}]_3 V_{3,\text{E}} = 4,35 \cdot 10^{-4} \text{ mol}$$

Or nous savons qu'à l'équivalence l'espèce titrante a été introduite en proportions stœchiométriques avec l'espèce titrée, ce que nous pouvons traduire ici par l'égalité (issue des annulations des quantités de matière de diiode et d'ions thiosulfate) :

$$x_{2,\text{E}} = \frac{n_{\text{I}_2,\text{rest}}}{1} = \frac{n_{\text{S}_2\text{O}_3^{2-},\text{E}}}{2}.$$

Nous en déduisons qu'après avoir réagi avec la vitamine C, il restait dans le milieu une quantité de matière de diiode égale à :  $n_{\text{I}_2,\text{rest}} = 2,18 \cdot 10^{-4} \text{ mol}$ .

4. Nous pouvons affirmer que le diiode restant est le diiode qui n'a pas été consommé. C'est-à-dire que la somme des

quantités de matière de diiode consommé par la vitamine C et de diiode restant à l'issue de cette consommation est égale à la quantité de matière de diiode initialement introduite :

$$n_{I_2,i} = n_{I_2,rest} + n_{I_2,CPC} \Leftrightarrow n_{I_2,CPC} = n_{I_2,i} - n_{I_2,rest}$$

où  $n_{I_2,CPC}$  représente la quantité de matière de diiode consommé par la vitamine C.

Il nous suffit donc de déterminer la quantité de matière de diiode initialement introduite avec la solution fabriquée à partir du cachet, soit :  $n_{I_2,i} = c_2 V_2 = 5,00 \cdot 10^{-4}$  mol, et nous trouvons alors que  $n_{I_2,CPC} = 2,82 \cdot 10^{-4}$  mol.

Nous pouvons alors écrire l'avancement final de la première réaction par annulation des quantités de matière de vitamine C

et de diiode consommé, comme étant :  $x_{1,f} = \frac{n_{I_2,CPC}}{1} = \frac{n_{VC,i}}{1}$ ,

quantité de matière de vitamine C initialement présente dans le volume de jus prélevé.

Nous en déduisons :  $n_{VC,i} = 2,82 \cdot 10^{-4}$  mol.

Comme par ailleurs cette quantité de matière était contenue dans un volume  $V_1 = 10,0$  mL, nous en déduisons la concentration en acide ascorbique de la solution réalisée :

$$c_1 = \frac{n_{VC,i}}{V_1} = 2,82 \cdot 10^{-2} \text{ mol.L}^{-1}.$$

De là nous tirons enfin la quantité de matière totale  $n_0$  d'acide ascorbique contenue dans la fiole jaugée, c'est-à-dire dans le cachet de départ :  $n_0 = c_1 V_0 = 2,82 \cdot 10^{-3}$  mol.

Comme la masse molaire de l'acide ascorbique a pour valeur  $M_{VC} = 6M_C + 8M_H + 6M_O = 176,0$  g.mol<sup>-1</sup>, nous trouvons la masse d'acide ascorbique correspondante :  $m_{VC} = n_0 M_{VC} = 496$  mg.

On peut donc supposer que le « 500 » du nom du médicament représente la masse d'acide ascorbique que contient un cachet, exprimée en mg.

### 5.1 Facteurs cinétiques

#### — Qu'est-ce qu'une réaction lente ?

Les notions de lenteur ou de rapidité d'une réaction chimique sont difficiles à définir, car subjectives. En pratique, pour ce qui est de la chimie, on considérera qu'une réaction est lente si les états intermédiaires par lesquels passe le système réactionnel entre son état initial et son état final sont appréciables par l'expérimentateur (autrement dit au moins de l'ordre de la seconde, et plus confortablement supérieur à la minute).

#### — Quelles sont les hypothèses de la cinétique chimique ?

Un problème pour lequel beaucoup de paramètres varient simultanément peut rapidement devenir inextricable. Ainsi, il est nécessaire de bien cadrer notre étude avant d'aller plus loin. On se limite donc, dans le cadre de la cinétique chimique, à l'étude d'un système réactionnel contenu dans un **réacteur fermé et de composition homogène**. Ces deux hypothèses nous assurent :

- que le système n'échange pas de matière avec l'extérieur, ce qui permettra un suivi simple des quantités de matière ;
- que la vitesse de la réaction est bien la même en tout point du système réactionnel (expérimentalement, on réalise cette condition à l'aide d'un agitateur...)

#### — Quelles techniques permettent de suivre la cinétique d'une réaction chimique ?

Pour suivre la cinétique d'une réaction en solution, il suffit de pouvoir suivre l'évolution de la concentration de l'une des espèces qu'elle engage. À partir de cette concentration et du volume du système, on pourra remonter à la quantité de matière de l'espèce considérée et de là à l'avancement. En répercutant celui-ci sur les autres quantités de matière, nous pourrions ainsi suivre le bilan de matière du système tout au long du déroulé temporel de la réaction. Le choix de l'espèce sur laquelle on va fonder ce suivi sera essentiellement conditionné par son accessibilité à la mesure dans les conditions de l'expérience.

En pratique, le **suivi cinétique** d'une réaction lente peut se faire par l'un des biais suivants :

- colorimétrie ;
- spectrophotométrie ;
- conductimétrie ;
- dosages successifs d'une espèce donnée.

Un suivi colorimétrique ou spectrophotométrique ne sera bien sûr envisageable que si la transformation étudiée engage une espèce colorée ; le suivi conductimétrique pourra se faire si des espèces chargées sont mises en jeu dans la réaction. Enfin, pour envisager un suivi par dosages successifs, on devra réaliser avant chaque dosage une **tremp**e d'un échantillon du système réactionnel afin de stopper son évolution le temps du dosage.

## — Quels facteurs influencent la vitesse d'une réaction ?

Les facteurs influençant la vitesse d'une réaction sont appelés **facteurs cinétiques**. On en dénombre essentiellement trois :

- la **température** du système siège de la réaction ;
- l'éventuelle utilisation d'un **catalyseur** ;
- la **concentration** des réactifs ou du catalyseur dans le milieu réactionnel.

Plus la **température** du système chimique est **élevée**, plus **rapide** est l'évolution de celui-ci vers son état final (indépendamment du fait que, s'il s'agit d'un équilibre, celui-ci puisse également être déplacé si la température change). À l'inverse, on va parfois volontairement refroidir tout ou partie du système pour ralentir fortement et pratiquement stopper les transformations chimiques dont il est le siège : c'est l'idée de base qui régit l'utilisation d'un réfrigérateur ou d'un congélateur afin de conserver en état de consommation les aliments (en ralentissant au maximum les transformations chimiques responsables de leur dégradation). Au laboratoire, on utilise également une méthode appelée **trempe** du système, qui consiste à le refroidir (et/ou le diluer) brutalement afin de le figer dans son état à un instant donné et ainsi pouvoir analyser son contenu (par un dosage par exemple).

L'utilisation d'un **catalyseur** permet d'accélérer notablement l'évolution du système vers son état final. Rappelons simplement qu'un catalyseur est un corps dont la présence dans le milieu réactionnel augmente la vitesse d'évolution de la réaction dont celui-ci est le siège, **sans en modifier le bilan**. Il n'apparaît donc pas dans l'équation bilan et est totalement régénéré en fin de réaction. En pratique, on distingue deux types de catalyse :

- **catalyse homogène** pour laquelle les réactifs et le catalyseur forment une seule phase ;
- **catalyse hétérogène** lorsque réactifs et catalyseur forment deux phases distinctes.

Enfin, une **augmentation de la concentration** des réactifs **accélère** la réaction. On peut en déduire l'évolution générale de la vitesse  $v(t)$  de la réaction. Au fur et à mesure que l'avancement de la réaction augmente, les réactifs sont consommés et leurs concentrations diminuent ; la vitesse volumique de la réaction  $v(t)$  diminue donc avec le temps.

**Remarque :** cette tendance générale admet toutefois une exception : celle des réactions autocatalysées. Dans ce type de réaction, l'un des produits formés est aussi catalyseur de la réaction, d'où une accélération de celle-ci à son début...

## — La température peut-elle modifier l'état final du système réactionnel ?

La réponse n'est pas simple, en ce sens qu'elle dépend de **grandeurs thermodynamiques** associées à la réaction. En pratique, on distingue thermodynamiquement parlant trois types de réactions :

- les réactions **exothermiques**, qui dégagent de l'énergie sous forme thermique ;
- les réactions **endothermiques**, qui absorbent de l'énergie thermique ;
- les réactions **athermiques**, qui ne s'accompagnent d'aucun transfert thermique.

Augmenter la température permet d'atteindre plus vite l'état d'équilibre tout en :

- le déplaçant dans le sens direct de la réaction si la réaction étudiée est endothermique ;
- le déplaçant dans le sens de la réaction inverse si la réaction étudiée est exothermique ;
- ne le modifiant pas si la réaction étudiée est athermique.

## — Qu'est-ce que la durée de demi-réaction ?

On appelle **durée de demi-réaction** la durée nécessaire pour que l'avancement de la réaction atteigne la moitié de sa valeur finale. Dans le cas d'une **réaction réversible**, la durée de demi-réaction se réfère bien sûr à l'avancement final, lui-même inférieur à l'avancement maximal...

# 5.2 Cinétique chimique : loi de vitesse d'ordre 1

## — Comment définir la vitesse d'une réaction chimique ?

De manière générale, une vitesse est toujours définie comme le taux d'accroissement d'une grandeur physique par rapport au temps (vitesse moyenne), ou comme la limite de ce taux d'accroissement lorsque l'intervalle de temps sur lequel il est calculé devient très petit à l'échelle de temps caractéristique du phénomène considéré (vitesse instantanée). On modélise alors cette limite par la dérivée temporelle de la grandeur considérée.

Dans le cas d'une réaction chimique, nous savons que la grandeur quantifiant le point auquel elle se fait est son avancement  $x$ . On définit ainsi la vitesse d'une réaction chimique comme :

$$v = \frac{dx}{dt}$$

$x$  en mol,  
 $t$  en secondes,  
 $v$  en  $\text{mol.s}^{-1}$ .

La vitesse ainsi définie donne donc littéralement le nombre de fois (le nombre de moles de fois pour être exact) que se fait la réaction par unité de temps, et donc la mesure du rythme auquel vont être consommés les réactifs et formés les produits qu'elle engage.

## — Comment relier cette vitesse aux quantités de matière de réactifs et de produits ?

Dans les faits on travaille rarement avec l'avancement proprement dit : la grandeur est moins concrète qu'une quantité de matière de réactif ou de produit, qui se déduisent plus facilement de grandeurs mesurables. En remplaçant  $x$  par son expression en fonction d'une quantité de

matière de réactif ou de produit, on trouve que : 
$$v = -\frac{1}{r} \frac{dn_R(x)}{dt} = +\frac{1}{p} \frac{dn_P(x)}{dt}$$

En réalité, tant que nous disposons d'une expression analytique liant l'avancement à une grandeur mesurable (en particulier dans le cas des modèles affines, voire linéaire, qui constituent une grande partie des relations que nous utilisons), nous pouvons exprimer la vitesse de réaction à partir de la dérivée temporelle de cette grandeur, multipliée par le facteur multiplicatif liant cette grandeur à l'avancement.

Si par exemple nous mesurons la masse d'une espèce  $X$ , nous pouvons écrire que :

$$n_X(x) = \frac{m_X(x)}{M_X}$$

et il vient alors, selon que l'espèce considérée est un réactif ( $X = R$ ) ou un produit ( $X = P$ ) :

$$v = -\frac{1}{r \times M_R} \times \frac{dm_R}{dt} = +\frac{1}{p \times M_P} \times \frac{dm_P}{dt}$$

Si la grandeur suivie est la **concentration effective** d'une espèce dans un volume  $V$  de solution, on a alors :

$$n_X(x) = [X](x) \times V$$

et il vient alors, pour peu que la réaction se déroule à volume constant :

$$v = -\frac{V}{r} \times \frac{d[R]}{dt} = +\frac{V}{p} \times \frac{d[P]}{dt}$$

Ce cas se rencontre très fréquemment, à tel point que l'on définit souvent la vitesse volumique de réaction comme :

$$v_V = \frac{1}{V} \times v = -\frac{1}{r} \times \frac{d[R](x)}{dt} = +\frac{1}{p} \times \frac{d[P](x)}{dt}$$

Unités usuelles :  $[R], [P]$  en  $\text{mol.L}^{-1}$ ,  $t$  en secondes,  $v_V$  en  $\text{mol.L}^{-1}.\text{s}^{-1}$ .

**Remarque :** le concept est extensible à l'absorbance (il suffit de connaître le coefficient d'étalement dans la loi de Beer-Lambert), ainsi qu'à la conductivité (*via* les coefficients de la loi de Kohl-Rausch), etc.

## — Comment mesure-t-on concrètement une vitesse de réaction ?

Comme nous venons de le voir, dès lors que l'on suit le déroulement d'une réaction *via* l'abondance d'un réactif ou d'un produit dans le milieu réactionnel, la mesure de la vitesse de réaction peut toujours se ramener à celle de la dérivée par rapport au temps d'une grandeur mesurée, moyennant la connaissance de la relation entre la grandeur en question et l'avancement de la réaction.

Reste la question de savoir comment mesurer la dérivée de cette grandeur par rapport au temps. La solution est désarmante de simplicité : il suffit d'effectuer le suivi temporel de cette grandeur et de le modéliser par une courbe (recours à un logiciel dédié ou lissage à la main à partir d'un nuage de points expérimentaux). On sait alors que la dérivée de la fonction à une date donnée s'identifie simplement au **coefficient directeur de la tangente à la courbe représentative de cette fonction, à la date considérée**.

En résumé, pour connaître la vitesse de réaction à une date  $t_0$  :

- 1) On trace la courbe représentative des variations de la grandeur  $G(t)$  suivie au cours du temps.
- 2) On trace la tangente à cette courbe à la date  $t_0$ .
- 3) On détermine la pente de cette tangente : elle s'identifie directement à  $\left(\frac{dG}{dt}\right)_{t=t_0}$ .
- 4) Connaissant la relation liant  $G(t)$  à  $x(t)$ , on peut calculer  $\left(\frac{dx}{dt}\right)_{t=t_0}$  à partir de  $\left(\frac{dG}{dt}\right)_{t=t_0}$ .

**Remarque :** les grandeurs que l'on peut associer à un réactif (concentration, masse, volume, absorbance, conductivité, etc.) vont diminuer à mesure que le réactif est consommé. Les courbes seront

donc descendantes et les tangentes à ces courbes présenteront des pentes négatives. Nous avons cependant vu que la vitesse de réaction engageait toujours un signe « - », qui garantira donc une vitesse positive.

## — Qu'appelle-t-on « ordre » d'une réaction ?

On constate expérimentalement que dans certains cas, la vitesse volumique d'une réaction lente engageant deux réactifs  $R_1$  et  $R_2$  en solution aqueuse suit une loi (dite loi de vitesse) de la forme :

$$v_V(t) = k \times [R_1]^{\alpha_1} \times [R_2]^{\alpha_2}$$

$v_V$  en  $\text{mol.L}^{-1}.\text{s}^{-1}$ ,  $[R_1]$ ,  $[R_2]$  en  $\text{mol.L}^{-1}$ .

Seules les concentrations des réactifs interviennent (ainsi éventuellement que celle du catalyseur le cas échéant), à l'exclusion notamment des concentrations en produits.

Dans cette expression :

- le paramètre  $k$  est appelé **constante de vitesse** (attention à ne pas confondre avec la constante d'équilibre, en majuscule) ; pour une réaction donnée, sa valeur dépend uniquement de la température à laquelle sont portés les réactifs, et croît évidemment avec celle-ci ;
- son unité dépend des valeurs des coefficients  $\alpha_1$  et  $\alpha_2$  ;
- ces coefficients (toujours positifs) sont appelés **ordres partiels** de la réaction vis-à-vis des réactifs ;
- on appelle **ordre total** de la réaction la valeur de la somme  $\alpha_1 + \alpha_2$ .

Il est alors intéressant de constater que si nous mettons en vis-à-vis, dans une égalité :

- la définition (formelle) de la vitesse volumique de réaction en s'appuyant sur la concentration en l'un ou l'autre des réactifs ;
- une loi de vitesse du type vu ci-dessus ;

on obtient alors une **équation différentielle** par rapport au temps, portant sur la concentration du réactif en question. Les solutions de cette équation différentielle vont alors nous permettre d'identifier une expression analytique de l'évolution temporelle de la concentration du réactif. Vous verrez en CPGE que l'étude de ce type de loi constitue ensuite une base précieuse pour interpréter le(s) mécanisme(s) réactionnel(s) à l'œuvre au cours d'une transformation.

## — Quelles sont les solutions de l'équation différentielle obtenue ?

Vous apprendrez en CPGE à résoudre cette équation différentielle dans le cas de lois de vitesse d'ordre 0, 1 et 2. À l'issue de la terminale, vous êtes supposé(e) avoir déjà traité le cas d'ordre 1, c'est-à-dire la situation où la loi de vitesse répond à une expression du type :

$$v_V(t) = k \times [R]$$

En plaçant cette expression en vis-à-vis avec l'expression générale :

$$v_V = -\frac{1}{r} \times \frac{d[R]}{dt}$$

on obtient :

$$-\frac{1}{r} \times \frac{d[R]}{dt} = k \times [R] \Leftrightarrow \frac{d[R]}{dt} + kr \times [R] = 0$$

Nous reconnaissons une équation différentielle linéaire d'ordre 1 à coefficients constants, sans second membre. Nous savons que les solutions d'une telle équation sont de la forme :

$$[R](t) = A \times e^{-\frac{t}{\tau}}$$

avec  $\tau = \frac{1}{k \times r}$ , constante de temps caractéristique de la loi et  $A$  constante déterminée par les conditions aux limites ; typiquement dans ce cas de figure, on utilisera les conditions initiales :

$$[R]_{(t=0)} = [R]_i$$

On obtient ainsi l'égalité :

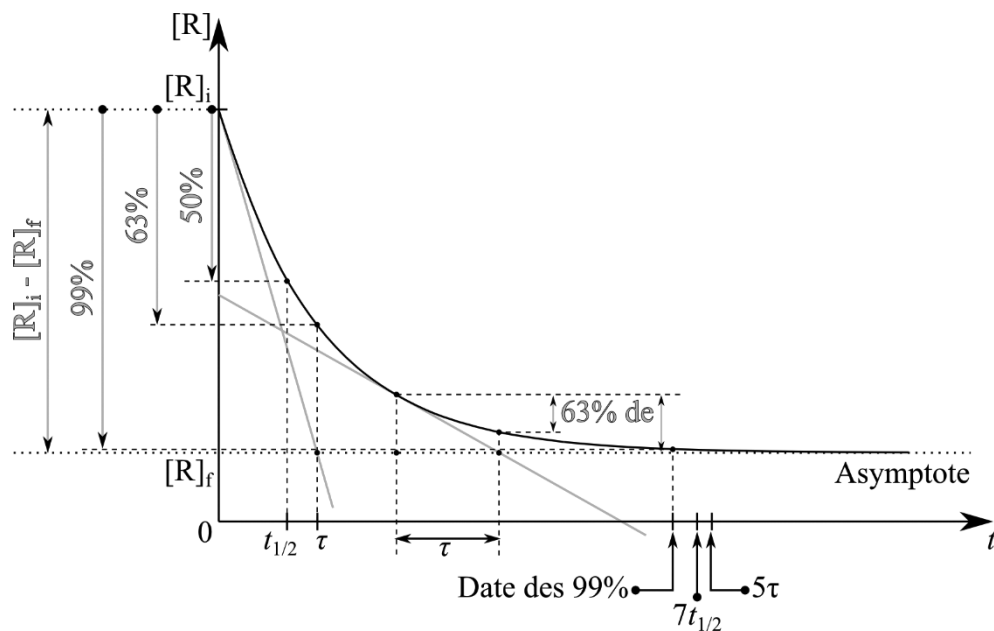
$$A \times e^{-\frac{t=0}{\tau}} = [R]_i \Leftrightarrow A = [R]_i$$

On peut alors affirmer que la concentration du réactif  $R$  suit la loi suivante :

$$[R](t) = [R]_i \times e^{-\frac{t}{\tau}}$$

Le suivi au cours du temps des valeurs de  $[R](t)$  nous permet ainsi, sachant que  $\tau = \frac{1}{k \times r}$ , de remonter à la valeur de la constante de vitesse de la réaction, valeur encore une fois riche d'information et permettant d'élaborer des modèles sur le mécanisme régissant la transformation étudiée.

Par ailleurs, cette loi s'accompagne naturellement des diverses propriétés inhérentes aux fonctions exponentielles :



- dérivée de la fonction directement proportionnelle à la fonction elle-même : si la concentration est divisée par 2, par 3, par 4, etc. la vitesse de réaction subit le même sort ;
- asymptote horizontale pour  $t \rightarrow \infty$  ;
- durée de demi-réaction indépendante de la concentration initiale, ayant pour effet que toute attente d'une même durée entraînera la division de la concentration par une même valeur, indépendamment de la concentration de départ ;
- lecture de  $\tau$  comme écart temporel entre une date quelconque, et l'intersection de la tangente à cette courbe à la date en question, avec l'asymptote ;
- durée de demi-réaction liée au paramètre  $\tau$  par la relation :

$$t_{1/2} = \tau \times \ln(2) ;$$

- pourcentage de chute selon la durée :
  - indépendant de la valeur initiale,
  - 50% au bout d'une durée  $t_{1/2}$  (par définition),
  - 63% au bout d'une durée  $\tau$ ,
  - >99% au bout d'une durée  $5\tau$  (reste 0,7%) ou  $7t_{1/2}$  (reste 0,8%).

## Halte aux idées reçues

Expliquer, pour chacune des affirmations suivantes, pourquoi elle est fausse :

1. Lorsque la vitesse d'une réaction chimique diminue, la quantité de matière de produit dans le milieu réactionnel diminue également.
2. Lorsque la loi de vitesse associée à une réaction dépend de la concentration d'un unique réactif, la durée de demi-réaction est indépendante de la concentration initiale en ce réactif.

3. Au bout d'une durée de demi-réaction, il reste dans le milieu réactionnel la moitié de la quantité initiale de chaque réactif.

4. La vitesse d'une réaction suivant une loi de vitesse d'ordre global égal à 1 dépend nécessairement de la concentration d'un seul et unique réactif.

5. Pour obtenir un maximum de produit à partir d'un mélange réactionnel, il est toujours souhaitable de traiter un mélange réactionnel à la plus haute température possible.

## Du Tac au Tac

(Exercices sans calculatrice)

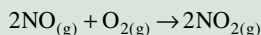
1. À température ambiante, le peroxyde d'hydrogène (couple oxydant/réducteur  $\text{H}_2\text{O}_{2(\text{aq})}/\text{H}_2\text{O}_{(\text{l})}$ ) réagit lentement avec les ions iodure (couple  $\text{I}_{2(\text{aq})}/\text{I}^-_{(\text{aq})}$ ).

Par ailleurs, en ambiance fraîche (quelques °C au-dessus de 0° C), le diiode (orange en solution aqueuse) :

- réagit selon une réaction totale et rapide avec les ions thiosulfate (couple  $\text{S}_4\text{O}_{6(\text{aq})}^{2-}/\text{S}_2\text{O}_{3(\text{aq})}^{2-}$ ) ;
- est indifférent au peroxyde d'hydrogène.

Proposer une méthode de suivi cinétique de la première réaction en vous fondant sur la deuxième, et exprimer la vitesse à partir d'une grandeur clé à mesurer.

2. Le monoxyde d'azote peut s'oxyder en dioxyde d'azote au contact du dioxygène, selon la réaction d'équation :



Montrer qu'une mesure en continu de la pression dans une enceinte rigide de volume  $V$  thermostatée à la température

$T$  suffit au suivi cinétique de la transformation, et détailler la méthode à suivre pour déterminer la vitesse de réaction à un instant donné, sur la base de la courbe représentative de l'évolution de la pression en fonction du temps.

3. Lorsque l'on souhaite savoir si une réaction présente un ordre par rapport à un unique réactif, on procède souvent en traçant le logarithme népérien de la vitesse de réaction en fonction de celui de la concentration dudit réactif.

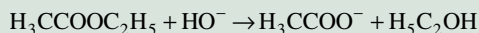
Expliquer en quoi cette façon de faire permet de répondre à la question et, dans l'affirmative, quelles informations relatives à la loi de vitesse il est possible d'en déduire.

4. La réaction d'oxydation de l'acide oxalique (couple oxydant/réducteur  $\text{CO}_{2(\text{g})}/\text{H}_2\text{C}_2\text{O}_{4(\text{aq})}$ ) par les ions permanganate (couple  $\text{MnO}_{4(\text{aq})}^-/\text{Mn}_{(\text{aq})}^{2+}$ ) est catalysée par les ions  $\text{Mn}_{(\text{aq})}^{2+}$ . Cette réaction est dite autocatalysée.

Justifier ce qualificatif, décrire l'évolution de sa vitesse au cours du temps et en déduire l'allure de la courbe représentative de l'évolution de l'avancement au cours du temps.

## Vers la prépa

L'hydrolyse basique de l'acétate d'éthyle peut être modélisée par la réaction d'équation :



On travaille avec une concentration initiale en acétate d'éthyle  $[\text{AdE}]_i = 1,0 \cdot 10^{-2} \text{ mol.L}^{-1}$  et l'on étudie la cinétique de la réaction avec deux solutions aqueuses de soude distinctes. On constate alors qu'au démarrage :

- pour  $[\text{HO}^-]_i = 1,0 \text{ mol.L}^{-1}$ , la réaction semble suivre une cinétique d'ordre 1 vis-à-vis du seul acétate d'éthyle, avec alors une constante de vitesse  $k_1 = 1,0 \cdot 10^{-3} \text{ s}^{-1}$  ;
- pour  $[\text{HO}^-]_i = 1,0 \cdot 10^{-4} \text{ mol.L}^{-1}$ , la réaction semble suivre une cinétique d'ordre 1 vis-à-vis des seuls ions hydroxyde, avec cette fois une constante de vitesse  $k_2 = 6,0 \cdot 10^{-4} \text{ min}^{-1}$ .

**1.** Montrer que ces résultats sont en accord avec l'hypothèse d'une loi de vitesse d'ordre partiel 1 vis-à-vis de chacun des deux réactifs.

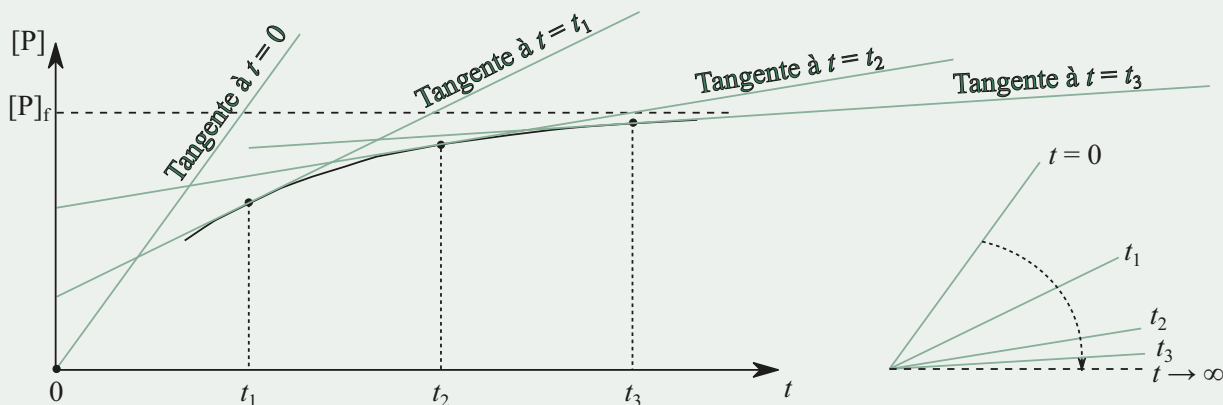
**2.** En déduire la loi de vitesse générale, en particulier la valeur de sa constante de vitesse.

**3.** Proposer sur cette base une méthode expérimentale permettant d'identifier de manière générale, s'ils existent, les ordres partiels d'une réaction vis-à-vis des différents réactifs qu'elle engage.

**4.** Déterminer l'expression en fonction du temps, du rapport des concentrations en ions hydroxyde et en acétate d'éthyle pour une concentration initiale  $[\text{HO}^-]_i$  quelconque.

## Halte aux idées reçues

1. Il est ici question d'une confusion qui, à ce stade de vos études, ne doit plus se produire, à savoir celle entre la valeur d'une fonction et celle de sa dérivée. La vitesse pouvant être définie à une constante multiplicative près comme la dérivée de la concentration en produit dans le milieu réactionnel, une perte de vitesse témoigne simplement d'une baisse du rythme de formation du produit au cours du temps, non de la quantité formée. En effet le rythme a beau baisser, il n'en demeure pas moins qu'il atteste d'une formation : si la quantité de produit augmente moins vite que précédemment, elle continue néanmoins d'augmenter.



2. Ayant *a priori* étudié uniquement des cinétiques d'ordre 1 à ce stade de vos études, l'amalgame entre unicité du réactif et loi de vitesse exponentielle (cas dans lequel l'affirmation de l'énoncé est exacte) serait sinon juste, du moins compréhensible.

Vous verrez cependant qu'il existe toutes sortes de lois de vitesse, qui peuvent atteindre d'importants niveaux de complexité pour peu que la réaction considérée repose en réalité sur plusieurs mécanismes réactionnels consécutifs.

Sans aller chercher ces situations (dont l'étude constituera un appétissant chapitre de votre programme en CPGE), nous pouvons tout de même traiter l'exemple d'une réaction obéissant à une cinétique d'ordre 2 vis-à-vis d'un unique réactif. On a dans ce cas :

$$v_V(t) = k \times [R]^2$$

Comme par ailleurs on a par définition :

$$v_V(t) = -\frac{1}{r} \times \frac{d[R]}{dt}$$

On peut interpréter ce point graphiquement : l'étude de la quantité de matière de produit formé au cours d'une réaction montre en général une montée rapide au début (les réactifs abondent, les collisions sont fréquentes, formant beaucoup de produit), puis l'augmentation se poursuit mais avec de moins en moins de vigueur (il y a de moins en moins de réactif).

Si l'on observe l'évolution au cours du temps de la tangente à la courbe représentative, on constate qu'elle est d'abord très raide (montée rapide de la courbe), autrement dit la vitesse est élevée. À mesure que le temps passe, la courbe se rapproche de l'horizontale, d'où une tangente également horizontale et donc une pente (et par association, une vitesse) qui tend vers 0.

on trouve, en combinant les deux expressions et en effectuant une séparation de variables :

$$-\frac{1}{r} \times \frac{d[R]}{dt} = k \times [R]^2 \Leftrightarrow -\frac{1}{[R]^2} d[R] = kr dt$$

L'intégration de cette égalité entre l'instant initial et un instant de date quelconque donne alors :

$$\frac{1}{[R]_{(t)}} - \frac{1}{[R]_{(t=0)}} = kr \times t \Leftrightarrow [R]_{(t)} = \frac{[R]_{(t=0)}}{1 + kr \times [R]_{(t=0)} \times t}$$

La durée de demi-réaction est donc cette fois, en supposant une réaction totale, la durée  $t_{1/2}$  pour laquelle :

$$[R]_{(t=t_{1/2})} = \frac{[R]_{(t=0)}}{2} \Leftrightarrow 1 + kr \times [R]_{(t=0)} \times t_{1/2} = 2$$

$$\Leftrightarrow t_{1/2} = \frac{1}{kr \times [R]_{(t=0)}}$$

dont nous constatons bien la dépendance explicite vis-à-vis de la concentration initiale. On retiendra donc que l'indépen-

dance de la durée de demi-réaction, vis-à-vis de la concentration initiale en réactif, constitue l'exception beaucoup plus que la règle, et qu'elle caractérise une cinétique d'ordre 1.

3. La durée de demi-réaction est la durée au bout de laquelle l'avancement a atteint la moitié de sa valeur finale. Or pour que reste seulement la moitié d'un réactif, il est nécessaire que :

$$n_R(t_{1/2}) = \frac{n_{R,i}}{2} = n_{R,i} - r \times \frac{x_f}{2} \leftrightarrow x_f = \frac{n_{R,i}}{r}$$

Nous reconnaissons dans l'expression  $\frac{n_{R,i}}{r}$  la valeur maximale autorisée à l'avancement par le réactif R. L'affirmation que propose l'énoncé revient donc à assimiler la valeur finale de l'avancement à la valeur maximale autorisée à l'avancement par ce réactif.

Or nous savons que la valeur finale de l'avancement correspond à la valeur maximale autorisée par un réactif à deux conditions :

- la réaction doit être totale (pour que la valeur finale dépende uniquement des quantités de matière, sans modulation par un éventuel équilibre chimique) ;
- le réactif considéré doit être un réactif limitant (pour que ce réactif impose bien la valeur maximale).

Or aucune de ces deux hypothèses n'était posée dans l'énoncé, qui donc est au mieux incomplet et au pire, faux. Il aurait fallu l'exprimer sous la forme :

« Au bout d'une durée de demi-réaction, il reste dans le milieu réactionnel la moitié de la quantité initiale de chaque réactif **limitant**, à supposer que cette réaction soit totale. »

4. Nous savons que l'ordre global d'une cinétique est défini comme la somme de ses ordres partiels. Si ces ordres sont entiers, alors effectivement un ordre 1 ne peut faire intervenir qu'un seul réactif. Cependant vous verrez que les lois de vitesse peuvent être très variées, certaines adoptant même des formes ne répondant pas à la définition d'ordre. On peut par exemple tout-à-fait envisager une cinétique de la forme :

$$v = k \sqrt{[R_1]} \times [R_2]$$

Dans ce cas la loi de vitesse admet un ordre partiel égal à  $\frac{1}{2}$  vis-à-vis de chacun des réactifs, soit un ordre total  $\frac{1}{2} + \frac{1}{2} = 1$ , tout en dépendant des concentrations de deux réactifs distincts.

Un point intéressant à ce stade, et que vous détaillerez en CPGE, est le fait que si une réaction est constituée d'un unique acte élémentaire (c'est-à-dire résulte d'une seule collision), alors sa loi de vitesse est simplement proportionnelle aux concentrations des réactifs concernés, chacune élevée à une puissance égale à son nombre stœchiométrique.

On peut ainsi :

- connaissant les actes élémentaires dont l'enchaînement donne une certaine réaction, déterminer sa loi de vitesse globale en établissant les lois de vitesse associées respectivement à chacun de ces actes (on obtient un système

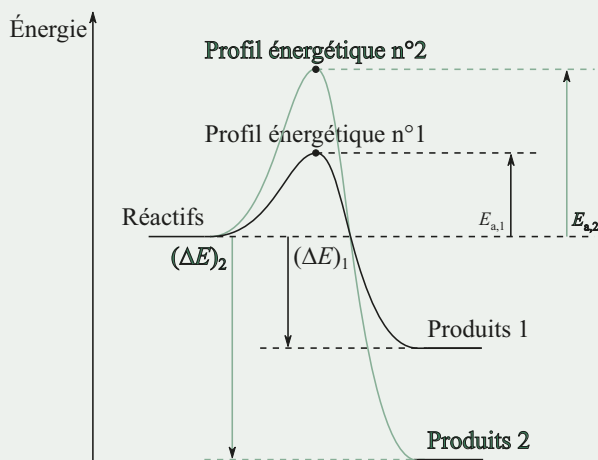
d'équations différentielles que l'on simplifie moyennant certaines hypothèses, et dont la résolution donne la loi globale) ;

- plus intéressant : en analysant la cinétique d'une réaction, chercher une forme mathématique modélisant convenablement la loi de vitesse, et en déduire quel enchaînement d'actes élémentaires pourrait donner un système d'équations dont cette loi de vitesse serait une solution.

5. Cette affirmation est fausse à deux niveaux :

- en premier lieu, nous savons que la composition finale d'un système siège d'un équilibre chimique dépend de la température (via sa constante d'équilibre), et qu'une élévation de température provoque un déplacement d'équilibre dans le sens endothermique. Donc si la réaction sur laquelle repose la formation d'un produit est exothermique, une élévation de température va au contraire réduire la quantité finale de produit obtenu ;
- il est cependant vrai que ce problème peut être contourné par diverses techniques (notamment extraction *in situ* du produit formé, de façon à maintenir le quotient de réaction en-dessous de sa valeur d'équilibre, et ce faisant forcer la réaction à se poursuivre indéfiniment) ;
- mais surtout, à un mélange réactionnel donné peuvent être associées plusieurs réactions, qui n'aboutiront pas au même produit, l'une étant favorisée par rapport à l'autre selon les conditions dans lesquelles est menée la réaction.

On peut notamment trouver des situations telles que celle-ci :



Nous constatons que :

- les produits les plus stables (favorisés thermodynamiquement) sont ceux issus de la réaction n° 2, énergétiquement plus bas que ceux issus de la réaction n° 1 ( $|(\Delta E)_2| > |(\Delta E)_1|$ ) ;
- les produits formés le plus vite (favorisés cinétiquement) sont ceux issus de la réaction n° 1, réclamant une énergie d'activation plus élevée ( $E_{a,2} > E_{a,1}$ ).

Si nous chauffons autant qu'il est possible sans altération des espèces (autre que la réaction, s'entend), alors l'énergie coulera à flots et nous formerons essentiellement les produits

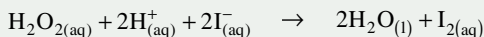
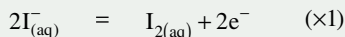
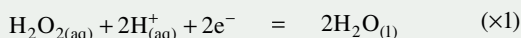
issus de la réaction n° 2 ; on dit dans ce cas que la transformation est menée sous **contrôle thermodynamique** : on obtient les **espèces les plus stables**.

Si en revanche nous limitons les apports énergétiques au strict minimum, les espèce n° 2 seront peu formées, les molécules dotées d'une énergie d'activation suffisante pour les donner étant statistiquement peu présentes dans le milieu réactionnel. Ce sont donc cette fois les produits issus de la réaction n° 1 qui seront favorisés ; on dit alors que la transformation est menée sous **contrôle cinétique** : on obtient les **espèces qui se forment le plus rapidement**.

Vous verrez, notamment en chimie organique, comment il est possible d'orienter le résultat d'une synthèse de manière très significative en jouant sur le type de contrôle adopté.

## Du Tac au Tac

1. Commençons par établir l'équation de la première réaction :



Pour en étudier la cinétique, nous avons besoin d'un indicateur de son état d'avancement, par exemple la concentration dans le milieu de l'un des réactifs ou de l'un des produits qu'elle engage. Dans le cas présent, le diiode semble un bon candidat puisque sa quantité de matière formée s'identifie directement à l'avancement de la réaction, qu'il est coloré et surtout qu'il réagit avec les ions thiosulfate selon une réaction dont l'énoncé précisait qu'elle est :

- totale ;
- rapide ;
- sans interférence (du moins à froid) avec le peroxyde d'hydrogène.

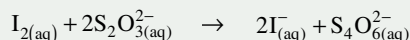
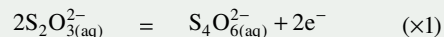
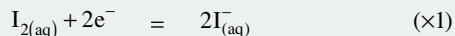
Tout nous invite à considérer l'idée d'une série de titrages, menés au fur et à mesure du déroulement de la première réaction.

Titre une solution siège d'une réaction chimique (autre que la réaction de titrage) pose en soi problème, mais l'énoncé précise que la première réaction est lente à température ambiante. On peut donc imaginer sans peine qu'elle soit très lente à froid.

Il suffit donc d'amorcer la première réaction à température ambiante en mélangeant les deux réactifs (solution aqueuse de peroxyde d'hydrogène et solution aqueuse d'iodure de potassium) puis à intervalle de temps régulier, de prélever un échantillon de volume fixé dans le milieu réactionnel et de l'inonder d'eau froide (on dit que l'on effectue une *trempe*).

Ce faisant, la réaction dont nous effectuons le suivi cinétique sera bloquée et la quantité de matière de diiode formée restera telle qu'elle était au moment du refroidissement.

Il ne nous restera plus qu'à titrer le diiode formé au moyen d'une solution aqueuse de thiosulfate de sodium (équivalence détectée à la décoloration de la solution, mise en exergue au moyen d'empois d'amidon), selon la réaction d'équation :



La quantité de matière d'ions thiosulfate versée l'équivalence sera donc le double de la quantité de matière de diiode effectivement présente, et nous aurons, sur l'échantillon titré dont la composition à la date  $t$  est restée figée par la trempe :

$$x(t) = n_{\text{I}_2}(t) = \frac{n_{\text{S}_2\text{O}_3^{2-}, \text{E}}(t)}{2} = \frac{[\text{S}_2\text{O}_3^{2-}]_{\text{titrante}}}{2} \times V_{\text{E}}(t) \Leftrightarrow$$

$$v = \frac{dx}{dt} = \frac{[\text{S}_2\text{O}_3^{2-}]_{\text{titrante}}}{2} \times \frac{dV_{\text{E}}}{dt}$$

Une fois obtenue la courbe représentative des variations du volume versé à l'équivalence en fonction de la date de prélèvement de l'échantillon titré, nous pourrions déterminer la vitesse de réaction à une date  $t_0$  donnée. Il nous suffira de tracer la tangente à ladite courbe à la date en question, et à en déterminer la pente, dont la valeur correspondra donc à  $\left(\frac{dV_{\text{E}}}{dt}\right)_{t=t_0}$ . Il ne restera plus enfin qu'à multiplier cette valeur par  $\frac{[\text{S}_2\text{O}_3^{2-}]_{\text{titrante}}}{2}$ .

2. Nous savons que la vitesse de réaction s'exprime :

$$v = \frac{dx}{dt}$$

Il nous suffit donc d'exprimer cet avancement en fonction de la pression globale pour répondre à la consigne. Commençons donc par exprimer la pression dans l'enceinte d'après la loi du gaz parfait, nous avons :

$$P \times V = n_{\text{gaz}} \times R \times T \quad \text{avec} \quad n_{\text{gaz}} = n_{\text{NO}} + n_{\text{O}_2} + n_{\text{NO}_2}$$

Détaillons alors les quantités de matière des différents gaz en fonction de l'avancement :

$$n_{\text{NO}} = n_{\text{NO},i} - 2x \quad n_{\text{O}_2} = n_{\text{O}_2,i} - x \quad n_{\text{NO}_2} = n_{\text{NO}_2,i} + 2x$$

Nous obtenons alors :

$$n_{\text{gaz}} = n_{\text{NO},i} + n_{\text{O}_2,i} + n_{\text{NO}_2,i} - x \quad \Leftrightarrow \quad x$$

$$= n_{\text{NO},i} + n_{\text{O}_2,i} + n_{\text{NO}_2,i} - \frac{V}{R \times T} \times P$$

En dérivant par rapport au temps, nous trouvons alors que la vitesse à une date  $t_0$  s'exprime :

$$v_{(t=t_0)} = \left(\frac{dx}{dt}\right)_{t=t_0} = -\frac{V}{R \times T} \times \left(\frac{dP}{dt}\right)_{t=t_0}$$

Nous savons que  $\left(\frac{dP}{dt}\right)_{t=t_0}$  n'est autre que la pente de la tangente à la courbe représentative de  $P(t)$  à la date  $t = t_0$  ; il suffit donc de tracer la tangente en question, d'en déterminer la pente (négative, puisque la réaction produit moins de gaz qu'elle n'en consomme, d'où une pression décroissante) et finalement de multiplier la valeur ainsi déterminée par le terme  $\left(-\frac{V}{R \times T}\right)$ .

3. Nous savons que si une réaction présente un ordre  $\alpha$  par rapport à un unique réactif X, alors sa vitesse suit une loi de la forme :

$$v = k \times [X]^\alpha$$

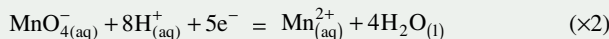
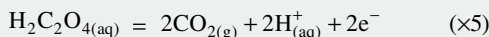
Le logarithme d'une telle expression donne :

$$\ln v = \alpha \times \ln [X] + \ln k$$

$\ln v$  est donc une fonction affine de  $\ln [X]$  et nous en déduisons que si les points du nuage représentatif des variations de  $\ln v$  en fonction de  $\ln [X]$  sont alignés entre eux, alors la vitesse de réaction peut effectivement être modélisée par une loi de la forme vue plus haut, dont :

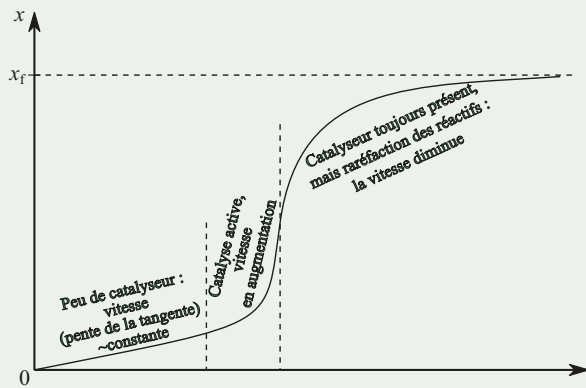
- l'ordre s'identifie directement à la pente de la droite modélisant le nuage de points ;
- la constante de vitesse est égale à l'exponentielle de l'ordonnée à l'origine.

4. La réaction envisagée est catalysée par les ions manganèse (II), or ces derniers constituent le réducteur conjugué des ions permanganate utilisés comme réactifs, et figurent donc parmi les produits qu'elle forme :



Nous comprenons ainsi que la réaction forme son propre catalyseur, d'où le qualificatif « autocatalysée ». Ce faisant, nous pouvons anticiper les grandes lignes de son déroulement :

- dans un premier temps, elle va démarrer à un certain rythme, sans influence des ions manganèse (II) (initialement absents du milieu réactionnel) ;
- à mesure qu'elle se déroule, elle va former des ions manganèse (II) qui vont commencer à la catalyser. Leur quantité croissant à mesure que la réaction avance, sa vitesse (donc la raideur de la courbe  $x(t)$ ) va s'accroître de plus en plus ;
- nous savons par ailleurs que la concentration des réactifs est également un facteur cinétique, or à mesure que la réaction se déroule, les ions manganèse (II) vont être de moins en moins nombreux. La vitesse de la réaction va donc subir deux influences contradictoires : un catalyseur dont la concentration augmente, et des réactifs dont les concentrations diminuent ;
- cependant si les réactifs peuvent réagir sans catalyseur, le catalyseur ne peut rien accélérer sans réactif à catalyser, et au bout d'un moment la réaction va décélérer.



## Vers la prépa

1. Commençons par expliciter les lois de vitesse décrites dans l'énoncé.

Exprimons donc les lois de vitesse décrites par l'énoncé :

- dans le premier cas :  $v_V = k_1 \times [\text{AdE}]$  avec  $k_1 = 1,0 \cdot 10^{-3} \text{ s}^{-1}$

$$\text{lorsque } [\text{HO}^-]_{i,1} = 1,0 \text{ mol.L}^{-1};$$

- dans le second cas :  $v_V = k_2 \times [\text{HO}^-]$  avec  $k_2 = 6,0 \cdot 10^{-6} \text{ min}^{-1}$

$$\text{lorsque } [\text{HO}^-]_{i,2} = 1,0 \cdot 10^{-4} \text{ mol.L}^{-1}.$$

Imaginons donc comme nous le suggère l'énoncé, que la loi de vitesse soit en réalité d'ordre 1 par rapport à chacun des deux réactifs :

$$v_V = k \times [\text{AdE}] \times [\text{HO}^-]$$

et examinons en quoi l'influence de la concentration en ions hydroxyde serait susceptible de changer l'ordre de cette cinétique :

- Dans le premier cas, on note que  $[\text{HO}^-]_i \gg [\text{AdE}]_i$  : les deux espèces réagissant mole à mole, l'avancement de la réaction va, en termes relatifs, beaucoup moins affecter la concentration en acétate d'éthyle, que celle en ions hydroxyde. Si par exemple l'acétate d'éthyle était consommé en totalité, la concentration des ions hydroxyde n'aurait baissé que de 1%. On aurait ainsi :  $v_V = k \times \underbrace{[\text{HO}^-]}_{\approx \text{cte}} \times [\text{AdE}]$  ;
- en notant alors  $k_1 = k \times [\text{HO}^-]_{i,1}$ , la loi de vitesse donne effectivement l'illusion d'une cinétique d'ordre 1 vis-à-vis de l'acétate d'éthyle ;
- dans le second cas, c'est l'inverse : on a  $[\text{HO}^-]_i \ll [\text{AdE}]_i$ . C'est donc cette fois l'acétate d'éthyle dont la concentration varie très peu en termes relatifs. On aura donc cette fois :  $v_V = k \times \underbrace{[\text{AdE}]}_{\approx \text{cte}} \times [\text{HO}^-]$  ;
- en notant alors  $k_2 = k \times [\text{AdE}]_i$ , la loi de vitesse donne cette fois l'illusion d'une cinétique d'ordre 1 vis-à-vis des ions hydroxyde.

Cette situation est classique en cinétique chimique : on dit que l'ordre est **dégénéré**, c'est-à-dire que la dépendance de la loi de vitesse vis-à-vis d'un réactif (et donc l'ordre global de la cinétique si elle en possède un) n'apparaît pas, du fait que ce réactif est trop présent pour que ses variations soient sensibles. On ne perçoit alors que des ordres apparents, auxquels sont associées des constantes de vitesse également apparentes.

2. D'après la question précédente, la constante  $k$  peut être calculée comme :

$$k = \frac{k_1}{[\text{HO}^-]_{i,1}} = 1,0 \cdot 10^{-3} \text{ s}^{-1} \text{ et } k = \frac{k_2}{[\text{AdE}]_i} = 6,0 \cdot 10^{-2} \text{ min}^{-1} = 1,0 \cdot 10^{-3} \text{ s}^{-1}$$

3. Comme bien souvent dans le domaine scientifique, nous pouvons utiliser à notre profit ce qui de prime abord apparaît comme un défaut : certes la possibilité pour un ordre d'être dégénéré complique notre tâche, cependant le fait de ne pas voir l'existence de certains ordres nous permet de bloquer (tout du moins en apparence) l'influence d'un paramètre, et ce faisant de mieux examiner l'influence des autres.

Dans une étude cinétique, il sera donc intéressant de mener diverses expériences, en fixant chaque fois la concentration d'un réactif beaucoup plus bas que celles des autres. Ainsi une seule concentration variera significativement tandis que les autres demeureront sensiblement constantes, nous permettant d'identifier un à un les ordres partiels de la réaction vis-à-vis de chacun d'eux.

4. Pour expliciter la vitesse de réaction, il nous suffit alors d'intégrer l'une ou l'autre des équations différentielles suivantes :

$$v_v = \frac{d\left(\frac{x}{V}\right)}{dt} = k \times [\text{AdE}] \times [\text{HO}^-] \Leftrightarrow \frac{d\left(\frac{x}{V}\right)}{dt} = k \times \left([\text{AdE}]_i - \frac{x}{V}\right) \times \left([\text{HO}^-]_i - \frac{x}{V}\right)$$

En notant  $x_v = \frac{x}{V}$  l'avancement volumique de la réaction, cette équation différentielle devient :

$$\frac{dx_v}{([\text{AdE}]_i - x_v) \times ([\text{HO}^-]_i - x_v)} = k \times dt$$

Pour l'intégrer, on commence par décomposer le quotient du premier membre en facteurs simples :

$$\frac{1}{([\text{AdE}]_i - x_v) \times ([\text{HO}^-]_i - x_v)} = \frac{A}{[\text{AdE}]_i - x_v} + \frac{B}{[\text{HO}^-]_i - x_v}$$

On trouve alors par identification :

$$A = -B = \frac{1}{[\text{HO}^-]_i - [\text{AdE}]_i}$$

L'équation différentielle devient ainsi :

$$\frac{1}{[\text{HO}^-]_i - [\text{AdE}]_i} \times \left[ \frac{1}{[\text{AdE}]_i - x_v} - \frac{1}{[\text{HO}^-]_i - x_v} \right] dx_v = k \times dt$$

L'intégration entre le début et un instant de date  $t$  quelconque donne :

$$\frac{1}{[\text{HO}^-]_i - [\text{AdE}]_i} \times \left[ \ln \left( \frac{[\text{HO}^-]_i - x_v}{[\text{AdE}]_i - x_v} \right) - \ln \left( \frac{[\text{HO}^-]_i}{[\text{AdE}]_i} \right) \right] = kt$$

Ou encore :

$$\frac{1}{[\text{HO}^-]_i - [\text{AdE}]_i} \times \ln \left( \frac{[\text{HO}^-]}{[\text{HO}^-]_i} \times \frac{[\text{AdE}]_i}{[\text{AdE}]} \right) = kt$$

Et finalement :

$$\frac{[\text{HO}^-]}{[\text{AdE}]} = \frac{[\text{HO}^-]_i}{[\text{AdE}]_i} \times e^{([\text{HO}^-]_i - [\text{AdE}]_i) \times kt}$$

Notons que si l'une ou l'autre des concentrations initiales est très inférieure à l'autre, la solution correspond bien à une cinétique d'ordre 1 (décroissance exponentielle) pour l'espèce concernée.

| $[\text{HO}^-]_i \gg [\text{AdE}]_i$                                                                                                          | $[\text{HO}^-]_i \ll [\text{AdE}]_i$                                                                                                          |
|-----------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------|
| $[\text{HO}^-] \approx [\text{HO}^-]_i$ , d'où :<br>$[\text{AdE}] = [\text{AdE}]_i \times e^{([\text{AdE}]_i - [\text{HO}^-]_i) \times kt}$   | $[\text{AdE}] \approx [\text{AdE}]_i$ , d'où :<br>$[\text{HO}^-] = [\text{HO}^-]_i \times e^{([\text{HO}^-]_i - [\text{AdE}]_i) \times kt}$   |
| $[\text{AdE}]_i - [\text{HO}^-]_i \approx -[\text{HO}^-]_i$ , d'où :<br>$[\text{AdE}] = [\text{AdE}]_i \times e^{-[\text{HO}^-]_i \times kt}$ | $[\text{HO}^-]_i - [\text{AdE}]_i \approx -[\text{AdE}]_i$ , d'où :<br>$[\text{HO}^-] = [\text{HO}^-]_i \times e^{-[\text{AdE}]_i \times kt}$ |
| On pose $k_1 = k \times [\text{HO}^-]_i$ :<br>$[\text{AdE}] = [\text{AdE}]_i \times e^{-k_1 t}$                                               | On pose $k_2 = k \times [\text{AdE}]_i$ :<br>$[\text{HO}^-] = [\text{HO}^-]_i \times e^{-k_2 t}$                                              |

### 6.1 La cinématique

#### — Quel est l'objet de la cinématique ?

La cinématique consiste en une étude purement descriptive du mouvement, c'est-à-dire en particulier sans aucune considération pour les éventuelles causes dudit mouvement. En effet, le concept même de mouvement ne va pas sans certaines subtilités (en témoigne la prodigieuse quantité de contre-vérités et de non-sens accumulée en son temps à ce sujet par un certain Aristote) ; sa description réclame donc, pour être rigoureuse, une structure, un vocabulaire, des conventions... dont nous allons poser ici les bases.

#### — Qu'est-ce qu'un système ponctuel ?

Un premier niveau de complexité dans la description du mouvement d'un système, provient du fait qu'un système peut se déplacer en restant toujours parallèle à lui-même (on parle dans ce cas d'un mouvement de **translation**, au cours duquel tous les points du système décrivent la même trajectoire), mais il peut également tourner sur lui-même (au mouvement de translation précédemment décrit se superpose alors un mouvement de **rotation**, dont la description peut réclamer jusqu'à 3 angles indépendants).

Toutefois, si le système est suffisamment petit devant l'étendue de son mouvement, on peut le **modéliser par un unique point matériel** et dans ce cas résumer son mouvement à celui du point en question. On peut montrer de manière générale qu'en modélisant un système par un point matériel, le mouvement de celui-ci se résume à celui du **centre de masse** du système en question.

Comme souvent en Physique, il ne s'agit pas d'un critère strict, mais d'un modèle (celui du système résumé à un unique point matériel) assujéti à une approximation (système très petit par rapport à l'étendue du mouvement), et qui décrira d'autant mieux la réalité, que cette approximation sera mieux satisfaite.

#### — Par rapport à quoi définit-on le mouvement d'un point matériel ?

En pratique, le repérage du mouvement d'un point matériel se fait par rapport à un objet solide servant de référence. Ce solide est appelé **référentiel d'étude** du mouvement.

Le mouvement est une notion relative. Un même point matériel observé par rapport à deux référentiels distincts aura *a priori* des mouvements différents par rapport à chacun d'eux. Un objet fixé sur le siège d'une voiture en marche est immobile par rapport au référentiel de la voiture, mais en mouvement par rapport à une borne kilométrique.

Pour procéder à une analyse quantitative du mouvement, on dote le référentiel d'étude d'un repère, constitué d'un point solidaire du référentiel tenant lieu d'**origine**, et de 3 axes définissant **3 directions fixes** par rapport à ce référentiel.

**Remarque :** on trouve très souvent l'expression d'une étude menée **dans** tel ou tel référentiel. Cette expression, sans conséquence pour qui comprend bien son sujet, peut faire des ravages si l'on n'y prend garde. En effet, le fait de se trouver « dans » référentiel n'a pas plus de sens ni d'importance, que de se trouver « sur », « sous » ou « derrière » le référentiel d'étude. La seule question qui se pose est celle de savoir si le repérage des positions occupées par le système au cours du temps a été effectué depuis un point qui était **solidaire** d'un certain référentiel. Aussi dans le cadre de cet ouvrage éviterons-nous systématiquement l'expression « **dans** le référentiel... » pour lui préférer « **par rapport** au référentiel », et vous invitons-nous à faire de même.

## — Que signifie concrètement le fait de travailler par rapport à un certain référentiel ?

On peut définir arbitrairement un objet comme étant immobile, et considérer que c'est le reste de l'univers qui est en mouvement par rapport à lui. Mathématiquement, le référentiel est caractérisé par le repère qui lui est associé. L'immobilité du référentiel est donc définie à travers l'invariance des caractéristiques du repère en question, soit :

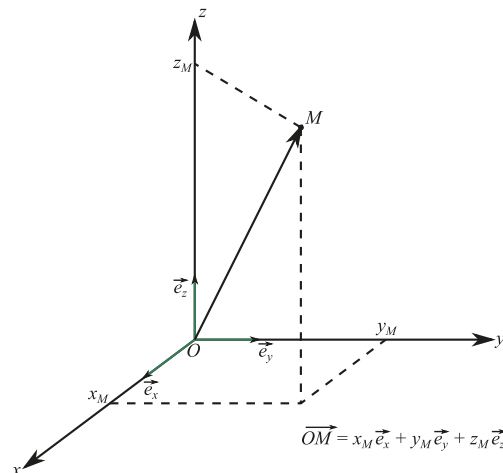
- l'origine de ce repère ;
- ses 3 axes orientés, définis par les 3 vecteurs unitaires constituant sa base orthonormée.

Choisir de travailler par rapport à un certain référentiel, c'est donc considérer que tout repère solidaire de ce référentiel ne varie pas au cours du temps. En particulier, cela signifie que **les dérivées par rapport au temps des vecteurs unitaires liés au référentiel d'étude sont égales au vecteur nul**. Nous verrons les implications concrètes de cette propriété un peu plus loin, lorsque nous définirons entre autres la vitesse et l'accélération.

## — Quel(s) système(s) de coordonnées peut-on utiliser pour repérer la position d'un point matériel ?

Le repérage d'un point dans un espace nécessite autant de coordonnées que cet espace possède de dimensions. Le long d'une courbe par exemple, une seule coordonnée suffit. Sur une surface, il en faut deux ; dans un espace à trois dimensions, il en faut trois.

Vous avez vu jusqu'ici le système de coordonnées cartésiennes. Dans ce système, les trois coordonnées d'un point matériel  $M$  sont les projections du vecteur  $\vec{OM}$  sur trois axes  $x$ ,  $y$  et  $z$  orientés respectivement par des vecteurs unitaires  $\vec{e}_x$ ,  $\vec{e}_y$  et  $\vec{e}_z$ ,  $O$  étant l'origine du système de coordonnées. Les coordonnées cartésiennes  $(x_M, y_M, z_M)$  du point  $M$  sont alors définies comme les coordonnées du vecteur  $\vec{OM}$ , c'est-à-dire les longueurs algébriques des projections orthogonales de ce vecteur, respectivement sur chacun des trois axes :



Notons au passage que dans un espace à 3 dimensions, les 3 vecteurs d'une base orthonormée peuvent être orientés de 2 façons différentes. Ainsi si l'on prend une base et que l'on intervertit deux de ses vecteurs, on constate que même en la faisant tourner de toutes les manières possibles, la nouvelle base obtenue n'est pas superposable à celle de départ. Dans les faits, on choisit toujours par convention une orientation telle que les vecteurs unitaires de la base forment un trièdre direct, c'est-à-dire qu'ils vérifient l'une ou l'autre des règles suivantes :

- **Règle de la main droite** : la main droite est orientée suivant le premier vecteur, le deuxième lui sortant de la paume ; le sens du troisième est donné par le pouce.
- **Règle du tire-bouchon (pour droitiers)** : le tire-bouchon tourne du premier vecteur vers le deuxième ; le sens dans lequel il se déplace est celui du troisième.

**Remarque** : ces deux règles sont parfaitement équivalentes pour vérifier qu'un trièdre est bien direct. Elles interviennent à de nombreux niveaux en Physique ; il est indispensable d'en connaître au moins une, et de la maîtriser parfaitement.

En CPGE, vous serez amenés à faire connaissance avec deux autres systèmes de coordonnées (les coordonnées cylindro-polaires et les coordonnées sphériques), qui engagent notamment des variables radiales et angulaires.

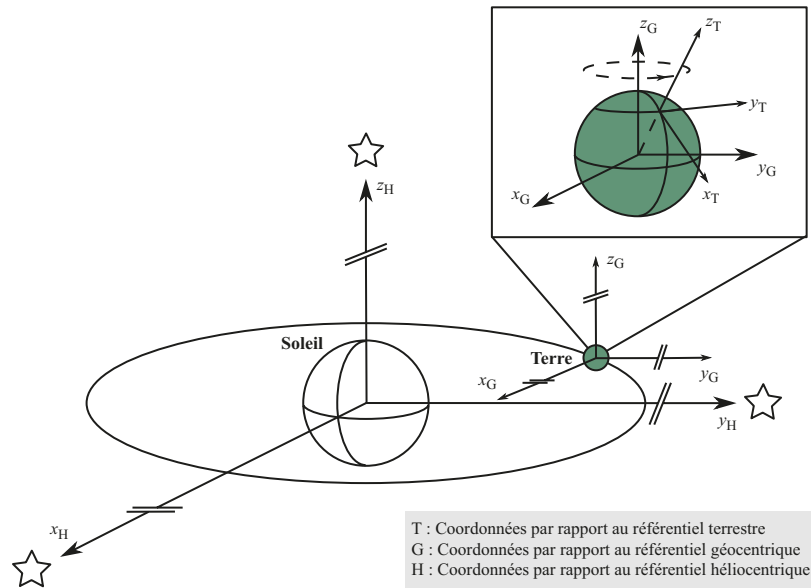
## — Quels sont les référentiels d'usage courant ?

Même s'il existe en soi une infinité de référentiels, nous verrons qu'en réalité ce sont un peu toujours les mêmes qui seront mis à l'honneur (et pour cause : nos observations et expériences sont la plupart du temps menées depuis les mêmes référentiels, et la plupart des lois qui permettent de les interpréter sont valables uniquement par rapport à ces derniers.

On retiendra donc essentiellement :

| Référentiel           | Solide de référence                                                                                                                                          | Origine du repère               | Axes du repères                                                                                                                             |
|-----------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Terrestre</b>      | La Terre.                                                                                                                                                    | Un point solidaire de la Terre. | <ul style="list-style-type: none"> <li>• Un axe vertical</li> <li>• 2 axes horizontaux (un nord-sud et un est-ouest par exemple)</li> </ul> |
| <b>Géocentrique</b>   | Terre virtuelle accompagnant la Terre dans son mouvement de translation circulaire au tour du soleil, mais pas dans son mouvement de rotation sur elle-même. | Centre de la Terre.             | 3 axes pointant vers des étoiles suffisamment lointaines pour pouvoir être considérées comme fixes au cours du temps.                       |
| <b>Héliocentrique</b> | Soleil virtuel coïncidant à chaque instant avec le soleil mais ne l'accompagnant pas dans son mouvement de rotation sur lui-même.                            | Centre du Soleil.               | 3 axes pointant vers des étoiles suffisamment lointaines pour pouvoir être considérées comme fixes au cours du temps.                       |

Les axes du référentiel géocentrique ne tournent pas par rapport au référentiel héliocentrique, aussi le mouvement du premier par rapport au second est-il un mouvement de **translation circulaire**, et non un mouvement de rotation comme on pourrait le croire. Le référentiel terrestre, lui, est en rotation par rapport au référentiel géocentrique, et à la fois en rotation et en translation par rapport au référentiel héliocentrique.



## — Quelles sont les grandeurs cinématiques usitées pour caractériser le mouvement d'un point matériel ?

Pour décrire le mouvement d'un point matériel  $M$  dans un espace à 3 dimensions, nous avons donc besoin des évolutions temporelles des 3 coordonnées décrivant la position de ce point au cours du temps :  $x_M(t)$ ,  $y_M(t)$  et  $z_M(t)$ . Or ces 3 coordonnées indépendantes se trouvent être les composantes du vecteur  $\overrightarrow{OM}(t)$  selon les vecteurs de la base orthonormée  $(\overrightarrow{e}_x, \overrightarrow{e}_y, \overrightarrow{e}_z)$ .

On appelle alors :

- **Vecteur position**  $\overrightarrow{OM}(t)$  d'un point matériel  $M$  par rapport à un référentiel donné, le vecteur liant l'origine  $O$  du repère associé à ce référentiel, au point  $M$  en question.
- **Vecteur vitesse**  $\vec{v}_M(t)$  de ce point matériel par rapport à ce référentiel, la dérivée du vecteur position  $\overrightarrow{OM}(t)$  de ce point par rapport au temps :

$$\vec{v}_M(t) = \frac{d\overrightarrow{OM}}{dt}$$

$$v_M \quad \text{en m.s}^{-1}$$

- **Vecteur quantité de mouvement**  $\vec{p}_M(t)$  de ce point matériel par rapport à ce référentiel, le produit de la masse de ce point, par son vecteur vitesse :

$$\vec{p}_M(t) = m \times v_M(t) = m \frac{d\overrightarrow{OM}}{dt}$$

$$p_M \quad \text{en kg.m.s}^{-1}$$

- **Vecteur accélération**  $\vec{a}_M(t)$  de ce point matériel par rapport à ce référentiel, la dérivée du vecteur vitesse  $\vec{v}_M(t)$  de ce point par rapport au temps :

$$\vec{a}_M(t) = \frac{d\vec{v}_M}{dt} = \frac{d^2\overline{OM}}{dt^2}$$

$$a_M \quad \text{en m.s}^{-2}$$

Considérons l'espace, muni d'un repère  $(O, \vec{e}_x, \vec{e}_y, \vec{e}_z)$ . Le repérage et l'évolution au cours du temps de la position de  $M$  de coordonnées  $x$ ,  $y$  et  $z$  (les notations étant à présent clairement définies, nous prenons le parti de les alléger en ne précisant plus qu'elles se rapportent au point  $M$ , et en sous-entendant leur dépendance vis-à-vis du temps) se fait à l'aide des grandeurs suivantes :

**Sa position**, est repérée par les composantes du vecteur  $\overline{OM}$  :

$$\overline{OM} = x\vec{e}_x + y\vec{e}_y + z\vec{e}_z = \begin{pmatrix} x \\ y \\ z \end{pmatrix}$$

$$x, y, z \quad \text{en m}$$

**Sa vitesse**, dérivée de  $\overline{OM}$  par rapport au temps :

$$\vec{v} = \frac{d\overline{OM}}{dt} = \begin{pmatrix} \dot{x} \\ \dot{y} \\ \dot{z} \end{pmatrix}$$

$$\dot{x}, \dot{y}, \dot{z} \quad \text{en m.s}^{-1}$$

**Remarque :** en Physique, on note  $\dot{f}$  la fonction dérivée par rapport au temps  $t$  de la fonction  $f(t)$

$$\text{Soit } \dot{f} = \frac{df}{dt}.$$

**Son accélération**, dérivée de  $\vec{v}$  par rapport au temps :

$$\vec{a} = \frac{d\vec{v}}{dt} = \frac{d^2\overline{OM}}{dt^2} = \begin{pmatrix} \ddot{x} \\ \ddot{y} \\ \ddot{z} \end{pmatrix}$$

$$\ddot{x}, \ddot{y}, \ddot{z} \quad \text{en m.s}^{-2}$$

**Remarque :** en Physique, on note  $\ddot{f}$  la fonction dérivée seconde par rapport au temps  $t$  de la fonction  $f(t)$ . Soit  $\ddot{f} = \frac{d^2f}{dt^2}$ .

**Remarque :** le vecteur position étant défini comme  $\overline{OM}(t) = x\vec{e}_x + y\vec{e}_y + z\vec{e}_z$ , sa dérivée temporelle devrait donc en toute rigueur donner  $\vec{v}_M(t) = \dot{x}\vec{e}_x + \dot{y}\vec{e}_y + \dot{z}\vec{e}_z + x\dot{\vec{e}}_x + y\dot{\vec{e}}_y + z\dot{\vec{e}}_z$ . Le fait que les dérivées des vecteurs unitaires s'annulent est autorisé uniquement du fait que ces derniers sont supposés être constants au cours du temps. Nous trouvons ici la conséquence concrète d'avoir clairement défini notre référentiel. En CPGE, vous serez amené(e)s à travailler avec des systèmes de coordonnées et/ou par rapport à des référentiels plus complexes. Les termes supplémentaires apparaissant au cours de ces opérations de dérivation ne seront alors pas toujours nuls, et il en résultera des composantes de vitesse ou d'accélération qui permettent notamment d'interpréter certains caractères non-galiléens des référentiels d'usage (forces de Coriolis et forces de marées, entre autres).

## 6.2 La dynamique et les lois de Newton

### — Quel est le propos de la dynamique ?

Avec la cinématique, nous avons appris à décrire le mouvement d'un point matériel à l'aide des grandeurs cinématiques que sont la position, la vitesse et l'accélération de ce point par rapport à un référentiel donné. La dynamique y ajoute la prise en compte des **causes** du mouvement de ce point : les forces auxquelles il est soumis.

L'objet des grands théorèmes de la mécanique (lois de Newton, théorème de l'énergie cinétique) consiste donc à lier quantitativement les causes du mouvement, aux grandeurs qui permettent de décrire ce mouvement en détail : prévoir le mouvement d'un système connaissant les forces auxquelles il est soumis, ou déterminer les caractéristiques d'une force à partir du mouvement d'un système qui en subit l'exercice.

**Remarque :** vous apprendrez en CPGE, notamment lorsque vous aborderez le cas des solides non ponctuels, d'autres causes du mouvement que sont les moments de force, ainsi que d'autres théorèmes permettant de les traiter, en tête desquels le théorème du moment cinétique.

### — Les lois de la dynamique sont-elles valables par rapport à n'importe quel référentiel ?

La plupart des lois de la dynamique engagent des grandeurs cinématiques, dont les caractéristiques dépendent du référentiel par rapport auquel elles sont calculées. Il est donc nécessaire de définir le cadre d'application de ces lois, c'est-à-dire entre autres les référentiels par rapport auxquels doivent être calculées les grandeurs cinématiques qu'elles engagent, pour que ces lois soient valides.

L'expérience prouve en effet qu'il existe une certaine catégorie de référentiels, appelés **référentiels d'inertie** ou **référentiels galiléens**, par rapport auxquels ces lois sont exactes.

### — Quels référentiels peuvent être considérés comme galiléens ?

En l'état actuel de nos connaissances, il n'est pas possible de mettre en évidence un référentiel parfaitement galiléen (ceci reviendrait à pouvoir définir un point de l'univers comme étant immobile dans l'absolu, autrement dit fixer un centre à l'univers, tâche dont on a commencé à mesurer la difficulté lorsque certaines croyances on laissé place à un peu de raisonnement...). On constate cependant que de nombreux référentiels constituent de bonnes approximations de référentiel galiléen (c'est-à-dire que les lois de la Dynamique sont bien vérifiées par rapport à ceux-ci). Dans les faits, nous considérerons que, de la moins bonne à la meilleure approximation : le référentiel terrestre ; le référentiel géocentrique ; le référentiel héliocentrique peuvent en général être considérés comme galiléens.

On retiendra en outre que **tout référentiel en mouvement de translation rectiligne uniforme (MTRU) par rapport à un référentiel galiléen est lui-même galiléen.**

Ceci signifie que si ces lois sont exactes par exemple par rapport au référentiel terrestre alors elles le sont également par rapport à tout référentiel en mouvement de translation rectiligne uniforme par rapport au référentiel terrestre, notamment tout véhicule se déplaçant tout droit et à vitesse constante par rapport à celui-ci.

**Remarque :** la proposition énoncée ci-dessus est une condition nécessaire et suffisante : si un référentiel n'est pas en MTRU par rapport à un référentiel galiléen, alors il n'est pas galiléen.

## — Que représente la quantité de mouvement d'un point matériel ?

La quantité de mouvement est le produit des deux grandeurs conférant à un système sa capacité à résister à l'influence d'une force extérieure (soit la masse et la vitesse dont il est doté). On voit donc que par sa définition, cette grandeur quantifie l'**inertie** de ce système.

Elle est par ailleurs consacrée par les lois de Newton qui, au-delà de cette approche intuitive, lui donnent une légitimité expérimentale.

## — Qu'est-ce que la 1<sup>re</sup> loi de Newton, ou principe d'inertie ?

Elle peut être énoncée de la façon suivante : il existe une catégorie particulière de référentiels, dits référentiels **galiléens** ou d'inertie, par rapport auxquels le centre d'inertie d'un système persévère dans un état de repos ou de mouvement rectiligne uniforme, si et seulement si la résultante vectorielle des forces s'exerçant sur lui est égale au vecteur nul.

**Remarque :** le principe d'inertie énonce simplement que si, parmi les trois conditions « le référentiel d'étude est galiléen », « les forces s'exerçant sur le système se compensent entre elles » et « le centre d'inertie du système persévère dans un état de repos ou de mouvement rectiligne uniforme », deux sont vérifiées alors la troisième l'est automatiquement.

## — Qu'est-ce que la 2<sup>e</sup> loi de Newton, ou principe fondamental de la dynamique (PFD) ?

Le **principe fondamental de la dynamique** s'énonce de la manière suivante : l'accélération du centre d'inertie d'un système de masse constante par rapport à un référentiel galiléen est égale au rapport de la résultante vectorielle des forces extérieures s'exerçant sur ce système, à la masse de celui-ci.

$$\vec{a} = \frac{\sum \vec{F}_{\text{ext}}}{m} \Leftrightarrow m\vec{a} = \sum \vec{F}_{\text{ext}}$$

|                  |                      |
|------------------|----------------------|
| $m$              | en kg                |
| $a$              | en $\text{m.s}^{-2}$ |
| $F_{\text{ext}}$ | en N                 |

**Remarque :** en réalité, la grandeur clé sur les 2 premières lois de Newton est la quantité de mouvement : c'est elle qui est conservée dans la première, et elle encore dont le taux de variation instantané par rapport au temps (autrement dit la dérivée) s'identifie à la résultante des forces subie par le système auquel elle se rapporte. Pour un système de masse constante (le cas que vous rencontrerez le plus souvent), la masse passe à travers l'opérateur de dérivation et l'on retrouve les énoncés ci-dessus. En revanche pour un système ouvert (typiquement : une fusée, dont la propulsion repose précisément sur l'éjection à grand débit des produits de combustion de ses ergols), il est nécessaire

de faire appel à cette version plus complète. Elle s'écrit alors :  $\frac{d\vec{p}}{dt} = \sum \vec{F}$ .

**Remarque :** il importe de bien noter que les seules forces à prendre en compte sont celles qui sont extérieures au système, c'est-à-dire qui sont exercées sur tout ou partie du système, par un objet **ne**

**faisant pas partie** du système. Pourtant les différents éléments constitutifs d'un système agissent les uns sur les autres (ne serait-ce que via les liaisons nucléaires, de covalence... assurant la cohésion dudit système). Nous verrons cependant avec la 3<sup>ème</sup> loi de Newton que les forces internes se compensent toutes deux à deux.

## — Qu'est-ce que la troisième loi de Newton, ou principe des actions réciproques ?

Le **principe des actions réciproques** s'énonce de la manière suivante : lorsqu'un système  $A$  exerce sur un système  $B$  une force  $\vec{F}_{A/B}$ , alors  $B$  exerce en retour sur  $A$  une force  $\vec{F}_{B/A}$ , telle que :

$$\vec{F}_{B/A} = -\vec{F}_{A/B}$$

$$F_{A/B}, F_{B/A} \text{ en N}$$

**Remarque :** cette troisième loi ne faisant intervenir aucune grandeur cinématique, elle est vraie quel que soit le référentiel d'étude, même si celui-ci n'est pas galiléen.

**Remarque :** on voit donc ici qu'à partir du moment où un élément du système exerce une force sur un autre (force interne au système, donc), alors cet autre élément exerce sur le premier une force exactement opposée. Bilan : ces deux forces antagonistes s'annulent et n'ont donc aucune influence sur le comportement dynamique du système. Ceci permet de souligner l'importance du choix du système pour le traitement d'un problème de mécanique : selon que tel ou tel élément matériel s'y trouve intégré ou non, on devra omettre ou au contraire prendre en compte les forces qu'exerce cet élément sur le reste du système.

## — Comment traite-t-on un problème de mécanique ?

L'étude dynamique d'un point matériel réclame toujours, pour être menée proprement, de s'interroger avant toute autre chose sur le « SRBdFext » :

- Quel est le **Système** étudié ?
- Par rapport à quel **Référentiel** l'étude va-t-elle être menée ?
- Quel est le **Bilan des Forces Extérieures** à ce système et s'exerçant sur celui-ci ?

Le système sera choisi de manière à permettre une étude aussi simple que possible. En particulier, on fera en sorte qu'il ne soit pas soumis à des forces peu ou pas connues, qui compliqueraient le problème en y introduisant un surcroît d'inconnues.

**Remarque :** vous aurez naturellement soin de choisir un référentiel qui soit galiléen, puisque dans le cas contraire la plupart des théorèmes de la mécanique ne seraient plus applicables.

# 6.3 Les interactions

## — À quelles actions un système mécanique peut-il être soumis ?

À l'heure actuelle, on a pu mettre en évidence **4 interactions fondamentales** : la gravitationnelle, l'électromagnétique, l'interaction nucléaire forte et l'interaction nucléaire faible. Il s'agit cependant d'interactions fondamentales, dont l'action est définie à l'échelle particulaire. Les forces avec lesquelles nous devons traiter la plupart du temps ne sont que des résultantes macroscopiques intégrant une ou plusieurs de ces forces, appliquées aux milliards de milliards de particules dont sont constitués les systèmes macroscopiques que nous étudions le plus souvent.

| Nom                                              | Subie par                                                       | Exercée par                                                                                                                     | Direction                        | Sens                                                                                                                                                                                                    | Valeur                               | Expression synthétique                                              | Paramètres                                                                                                                                                                                                                                                            |
|--------------------------------------------------|-----------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------|----------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------|---------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Force gravitationnelle</b><br>$\vec{F}_g$     | Tout objet de centre de masse $B$ et de masse $m_B$             | Tout objet ponctuel ou de centre de masse $A$ , et dont la masse $m_A$ est répartie selon une distribution à symétrie sphérique | Droite ( $AB$ )                  | De $B$ vers $A$                                                                                                                                                                                         | $\mathcal{G} \frac{m_A m_B}{(AB)^2}$ | $\vec{F}_{gA/B} = -\mathcal{G} \frac{m_A m_B}{(AB)^2} \vec{e}_{AB}$ | <ul style="list-style-type: none"> <li><math>\mathcal{G} = 6,67 \cdot 10^{-11}</math> USI</li> <li>Masses en kg</li> <li><math>AB</math> : distance (m)</li> </ul>                                                                                                    |
| <b>Poids</b><br>$\vec{P}$                        |                                                                 | Tout astre générant un champ de pesanteur $\vec{g}_A$                                                                           | Celle de $\vec{g}_A$ (verticale) | Celui de $\vec{g}_A$ (haut $\rightarrow$ bas)                                                                                                                                                           | $m_B \times \ \vec{g}_A\ $           | $\vec{P}_{A/B} = m_B \vec{g}_A$                                     | <ul style="list-style-type: none"> <li><math>m_B</math> : masse de <math>B</math> (kg)</li> <li><math>\vec{g}_A</math> : accélération de pesanteur (USI : <math>N \cdot kg^{-1}</math> ou <math>m \cdot s^{-2}</math>)</li> </ul>                                     |
| <b>Force coulombienne</b><br>$\vec{F}_c$         | Tout objet ponctuel $B$ porteur d'une charge électrique $q_B$   | Tout objet ponctuel $A$ porteur d'une charge électrique $q_A$                                                                   | Droite ( $AB$ )                  | <ul style="list-style-type: none"> <li><math>q_A q_B &gt; 0</math> : <math>A \rightarrow B</math> (répulsif)</li> <li><math>q_A q_B &lt; 0</math> : <math>B \rightarrow A</math> (attractif)</li> </ul> | $k \frac{ q_A q_B }{(AB)^2}$         | $\vec{F}_{eA/B} = k \frac{q_A q_B}{(AB)^2} \vec{e}_{AB}$            | <ul style="list-style-type: none"> <li><math>k = \frac{1}{4\pi\epsilon_0} = 8,99 \cdot 10^9</math> USI dans le vide</li> <li>Charges électriques en C</li> <li><math>AB</math> : distance (m)</li> </ul>                                                              |
| <b>Force électrique</b><br>$\vec{F}_e$           |                                                                 | Tout corps $A$ générant un champ électrique $\vec{E}_A$                                                                         | Celle de $\vec{E}_A$             | <ul style="list-style-type: none"> <li><math>q_B &gt; 0</math> : sens de <math>\vec{E}_A</math></li> <li><math>q_B &lt; 0</math> : opposé à <math>\vec{E}_A</math></li> </ul>                           | $ q_B  \times \ \vec{E}_A\ $         | $\vec{F}_{eA/B} = q_B \vec{E}_A$                                    | <ul style="list-style-type: none"> <li><math>q_B</math> : charge électrique de <math>B</math> (C)</li> <li><math>\vec{E}_A</math> : champ électrique généré par le corps <math>A</math> (USI : <math>N \cdot C^{-1}</math> ou <math>V \cdot m^{-1}</math>)</li> </ul> |
| <b>Poussée d'Archimède</b><br>$\vec{\Pi}_A$      | Tout objet immergé dans un fluide                               | Tout fluide subissant un champ de pesanteur $\vec{g}_A$                                                                         | Celle de $\vec{g}_A$ (verticale) | Opposé à $\vec{g}_A$ (bas $\rightarrow$ haut)                                                                                                                                                           | $\rho_f V_{\text{im}} g_A$           | $\vec{\Pi}_A = -\rho_f V_{\text{im}} \vec{g}_A$                     | <ul style="list-style-type: none"> <li><math>\rho_f</math> : masse volumique du fluide</li> <li><math>V_{\text{im}}</math> : volume immergé du système</li> <li><math>\vec{g}_A</math> : accélération de pesanteur</li> </ul>                                         |
| <b>Force de rappel harmonique</b><br>$\vec{F}_r$ | Tout système accroché à l'extrémité $B$ d'un ressort            | Tout ressort d'extrémités $A$ et $B$ et de longueur $l$                                                                         | Droite ( $AB$ )                  | <ul style="list-style-type: none"> <li>Traction : <math>B \rightarrow A</math></li> <li>Pression : <math>A \rightarrow B</math></li> </ul>                                                              | $k \times  l - l_0 $                 | $\vec{F}_r = -k(l - l_0) \vec{e}_{AB}$                              | <ul style="list-style-type: none"> <li><math>k</math> : constante de raideur (USI : <math>N \cdot m^{-1}</math>)</li> <li><math>l_0</math> : longueur à vide</li> </ul>                                                                                               |
| <b>Réaction normale</b><br>$\vec{R}_n$           | Tout objet au contact d'un support solide                       | Le support en question                                                                                                          | $\perp$ à la tangente au support | Du support vers le système                                                                                                                                                                              |                                      |                                                                     | Variable selon les cas                                                                                                                                                                                                                                                |
| <b>Frottements statiques</b><br>$\vec{f}_s$      |                                                                 |                                                                                                                                 | Tangente au support              | Empêche le mouvement                                                                                                                                                                                    |                                      |                                                                     |                                                                                                                                                                                                                                                                       |
| <b>Frottements dynamiques</b><br>$\vec{f}_d$     | Tout objet en mouvement par rapport à une matière à son contact | La matière en question                                                                                                          | Parallèle au vecteur vitesse     | Opposé au vecteur vitesse                                                                                                                                                                               |                                      |                                                                     |                                                                                                                                                                                                                                                                       |
| <b>Tension</b><br>$\vec{T}$                      | Tout objet lié à un(e) fil/tige                                 | Le fil ou la tige en question                                                                                                   | Le long du fil/tige              | Système vers fil/tige                                                                                                                                                                                   |                                      |                                                                     |                                                                                                                                                                                                                                                                       |
| Actions à distance                               |                                                                 |                                                                                                                                 |                                  |                                                                                                                                                                                                         |                                      |                                                                     |                                                                                                                                                                                                                                                                       |
| Actions de contact                               |                                                                 |                                                                                                                                 |                                  |                                                                                                                                                                                                         |                                      |                                                                     |                                                                                                                                                                                                                                                                       |

On retiendra donc que, même si la cuisine universelle se résume *in fine* à ces 4 ingrédients fondamentaux, ceux-ci se déclinent en une multitude de recettes, et que vous devez connaître les caractéristiques de chacune.

## 6.4 Mouvement d'un point matériel dans un champ uniforme

### Dans quel(s) cas peut-on être amené à traiter un système dans un champ de force uniforme ?

Parmi les forces intervenant dans de très nombreux problèmes, certaines ont le bon goût de conserver leurs caractéristiques sur toute l'étendue spatiale et/ou pendant toute la durée du mouvement. On trouve notamment dans ce domaine les mouvements dans un champ de pesanteur uniforme (situation vérifiée en bonne approximation tant que l'étendue du mouvement n'excède pas quelques dizaines de kilomètres), et les mouvements de particules électriquement chargées dans des zones de champ électrique uniforme (typiquement entre deux plaques parallèles séparées d'une distance faible devant leurs dimensions).

### Comment procède-t-on et à quelles lois aboutit-on ?

Après avoir clairement défini le système et le référentiel d'étude (évidemment un référentiel galiléen), on identifie le champ de force uniforme auquel est soumis le système. L'application de la 2<sup>e</sup> loi de Newton nous donne alors les composantes d'accélération (constante sur l'axe du champ du champ de force, nulle sur l'autre). Une première intégration donne les lois horaires de vitesse (fonction affine sur l'axe du champ de force, constante sur l'autre). Et une seconde intégration nous donne les lois horaires de position (fonction quadratique sur l'axe du champ de force, fonction affine sur l'autre).

|                                                                                      | Accélération                     | Vitesse                                                   | Position                                                                                    |
|--------------------------------------------------------------------------------------|----------------------------------|-----------------------------------------------------------|---------------------------------------------------------------------------------------------|
| Direction du champ de force (axe $Oz$ )                                              | $a_z = \frac{F}{m} = \text{cte}$ | $v_{z(t)} = \int a_z dt$<br>$= a_z \times t + v_{z(t=0)}$ | $z_{(t)} = \int v_z dt$<br>$= \frac{1}{2} a_z \times t^2 + v_{z(t=0)} \times t + z_{(t=0)}$ |
| Direction de la composante initiale de vitesse normale au champ de force (axe $Ox$ ) | $a_x = 0$                        | $v_{x(t)} = \int a_x dt$<br>$= \text{cte} = v_{x(t=0)}$   | $x_{(t)} = \int v_x dt$<br>$= v_{x(t=0)} \times t + x_{(t=0)}$                              |

### Que se passe-t-il dans le cas d'une force dépendant de la position et/ou de la vitesse du système ?

Si certaines forces sont constantes (ou peuvent en bonne approximation être considérées comme telles), d'autres dépendent significativement de la position du système et/ou de sa

vitesse. Elles vont donc évoluer à mesure que celui-ci se déplace ; leur évolution va affecter les équations du mouvement dans lesquelles elles figurent, entraînant du même coup une évolution de la vitesse et/ou de la position du système, etc. Ce serpent mathématique en train de se mordre la queue aboutit en général à une (ou plusieurs) équation différentielle. Dans ce cas les équations du mouvement seront les solutions de cette équation, dont les constantes seront déterminées, comme d'habitude, par des conditions particulières (souvent les conditions initiales) par lesquelles est passé le système.

## 6.5 Mécanique céleste

### — Que devient l'expression de la force gravitationnelle lorsqu'elle est exercée par un corps non ponctuel ?

On peut montrer que le champ gravitationnel généré par un corps dont la masse est répartie selon une distribution à symétrie sphérique, est identique à celui que génèrerait ce corps si toute sa matière se trouvait concentrée en son centre.

Or pour bon nombre de corps (planètes, étoiles...) la sphère pleine constitue souvent une très bonne approximation de la forme réelle, et l'expression vue dans le cas de deux corps ponctuels constitue un modèle d'interaction tout-à-fait satisfaisant.

En conclusion et jusqu'à nouvel ordre (que vous pourrez rencontrer à la faveur de certains exercices en CPGE), vous pourrez toujours considérer que la force exercée par un corps céleste sur un autre est semblable à l'expression que vous connaissez :

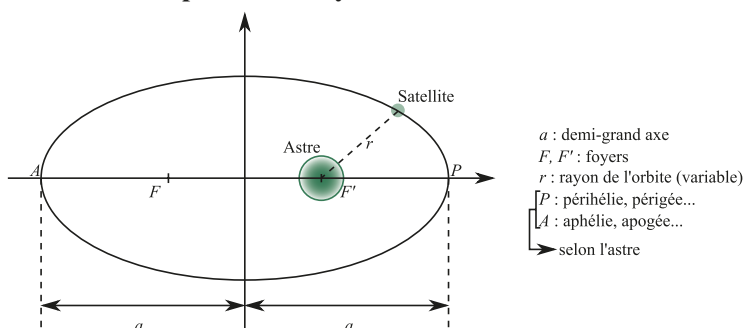
$$\vec{F}_{g,A/B} = -G \frac{m_A \times m_B}{r^2} \vec{e}_{AB}$$

|            |        |
|------------|--------|
| $G$        | en USI |
| $F_{A/B}$  | en N   |
| $m_A, m_B$ | en kg  |
| $r$        | en m   |

en redéfinissant les points  $A$  et  $B$  comme étant non plus les positions de deux corps ponctuels, mais celles des centres de deux corps dont les masses sont réparties selon des distributions à symétrie sphérique.

### — Quel est le mouvement d'un satellite autour d'un astre ?

Kepler énonce dans sa première loi que **les orbites des planètes autour du Soleil sont des ellipses, dont le Soleil occupe l'un des foyers** :

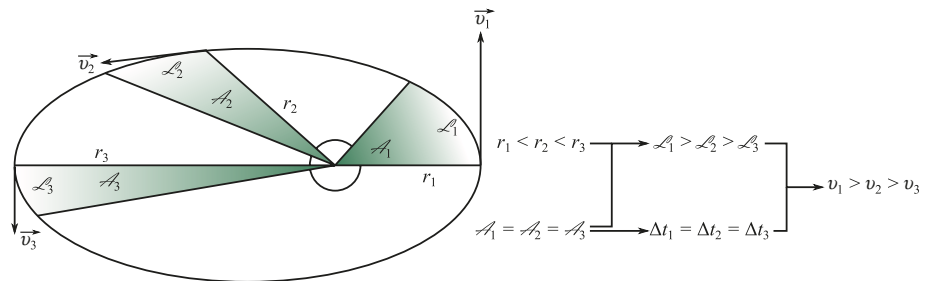


Vous avez vu en classe de Terminale que la loi de la gravitation universelle permettait d'expliquer la possibilité pour une planète d'avoir une orbite circulaire autour du Soleil. On retiendra que, de manière générale, un satellite gravitant autour d'un astre décrit une orbite elliptique dont cet astre occupe l'un des foyers (vous verrez la démonstration en CPGE).

**Remarque :** vous verrez également que dans le cas d'un corps ne bouclant pas une orbite autour de l'astre mais se contentant de passer à proximité, la trajectoire peut également être une parabole voire une hyperbole, toutes ces trajectoires (cercle, ellipse, parabole, hyperbole) étant les différentes déclinaisons d'une même famille de courbes appelées *coniques*.

## À quel rythme une planète évolue-t-elle sur son orbite ?

Kepler énonce dans sa deuxième loi que **le rayon liant une planète au Soleil balaye des aires égales en des durées égales** :



On appelle **vitesse aréolaire** la dérivée par rapport au temps de l'aire  $\mathcal{A}$  balayée par le rayon liant le Soleil à la planète considérée. On peut montrer qu'elle s'exprime :  $\frac{d\mathcal{A}}{dt} = \omega \frac{r^2}{2}$  où  $\omega$  représente la vitesse angulaire du satellite.

Nous constatons que, pour un mouvement circulaire ( $R = \text{cte}$ ), la vitesse angulaire  $\omega$  est constante : l'angle couvert par le rayon est donc dans ce cas proportionnel à la durée de balayage, autrement dit le mouvement est uniforme.

On retiendra que de manière générale, pour un satellite en orbite elliptique autour d'un astre, à la distance  $r(t)$  de celui-ci et animé d'une vitesse angulaire  $\omega(t)$  on peut écrire :

$$\frac{d\mathcal{A}}{dt} = \omega(t) \frac{r^2(t)}{2} = \text{cte}$$

|               |                        |
|---------------|------------------------|
| $\mathcal{A}$ | en $\text{m}^2$        |
| $t$           | en s                   |
| $\omega$      | en $\text{rad.s}^{-1}$ |
| $r$           | en m                   |

## Existe-t-il une constante du mouvement, commune à toutes les planètes du système solaire ?

Kepler énonce dans sa troisième loi que **le rapport du carré de la période de révolution  $T$  d'une planète autour du Soleil, au cube du demi-grand axe  $a$  de l'ellipse de son orbite, possède la même valeur quelle que soit la planète considérée** :

$$\frac{T^2}{a^3} = \text{cte}$$

|     |      |
|-----|------|
| $T$ | en s |
| $a$ | en m |

On retiendra que de manière générale, pour un satellite en orbite sur une ellipse de demi-grand axe  $a$  et accomplissant une révolution en une durée  $T$  autour d'un astre de masse  $M_A$ , on a :

$$\frac{T^2}{a^3} = \frac{4\pi^2}{M_A \mathcal{G}}$$

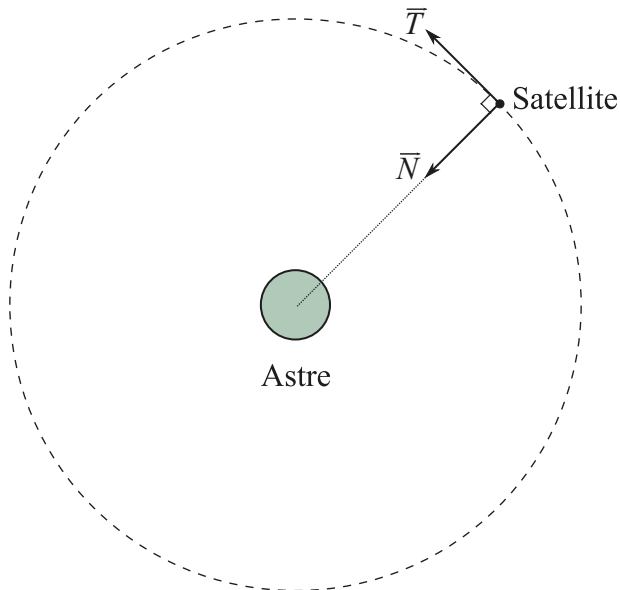
|               |        |
|---------------|--------|
| $T$           | en s   |
| $a$           | en m   |
| $M_A$         | en kg  |
| $\mathcal{G}$ | en USI |

**Remarque :** Kepler avait seulement découvert la constance de ce rapport, sans toutefois pouvoir en expliquer la valeur. La forme explicite ci-dessus, trouvée sur la base des travaux de Newton, fut l'une des grandes avancées dues à ce dernier. Elle permet notamment de déterminer la masse d'un astre par simple analyse du mouvement de ses satellites.

## — Quelles sont les caractéristiques du mouvement circulaire d'un corps autour de son astre attracteur ?

Bien que la première loi de Kepler stipule un mouvement elliptique des planètes autour du Soleil, l'approximation par un cercle est très souvent licite. C'est le cas notamment pour le mouvement de la Terre autour du Soleil, mais aussi pour celui de nombreux satellites autour de l'astre autour duquel ils orbitent (en fait les orbites très excentrées, telles que celles des comètes, constituent beaucoup plus l'exception que la règle). Dans ce cas, l'application de la deuxième loi de Newton nous permet d'établir un certain nombre de relations (cf. « Du Tac au Tac »).

On utilisera à cet effet une base de projection mobile, c'est-à-dire qui suit le système étudié dans son mouvement : la **base de Frénet** ( $\vec{T}, \vec{N}$ ).



**Remarque :** dans le cas d'un mouvement circulaire, la base de Frénet coïncide avec la base polaire.

On montre alors que **le mouvement circulaire de rayon  $R$  d'un corps autour de son astre attracteur de masse  $M$  est uniforme** à la vitesse constante  $v_{\text{sat}}$  :

$$v_{\text{sat}} = \sqrt{\frac{GM}{R}}$$

$M$  en kg  
 $R$  en m  
 $\mathcal{G}$  en USI

Cette vitesse, appelée **vitesse de satellisation**, est la vitesse minimale dont doit être doté un objet situé à une distance  $R$  du centre de l’astre, pour pouvoir boucler une orbite et devenir un satellite de cet astre plutôt que de s’écraser dessus.

Une autre vitesse intéressante est la **vitesse de libération** de l’astre, qui est la valeur minimale de la vitesse dont doit être doté un objet pour échapper à l’attraction gravitationnelle (c’est-à-dire partir à l’infini) de l’astre en question. Elle a pour expression :

$$v_{\text{lib}} = \sqrt{\frac{2GM}{R}} = \sqrt{2} \times v_{\text{sat}}$$

L’orbite du satellite a alors pour longueur  $2\pi R$  et pour période  $T$  :

$$T = \frac{2\pi R}{v_{\text{sat}}} = 2\pi \sqrt{\frac{R^3}{GM}}$$

$M$  en kg  
 $R$  en m  
 $\mathcal{G}$  en USI

**Remarque :** les propriétés énoncées ici sont très réduites du fait que nous nous limitons au cas des mouvements circulaires, ainsi qu’à une approche vectorielle des mouvements dans les champs newtoniens. L’intégration des divers types de coniques et des aspects énergétiques, ainsi que l’extension à d’autres interactions (n’oublions pas que l’interaction coulombienne varie également en  $\frac{1}{r^2} \vec{e}_r$ ) constitue une mine de résultats et d’applications extrêmement riches, que vous développerez abondamment en CPGE.

## Halte aux idées reçues

Expliquer, pour chacune des affirmations suivantes, pourquoi elle est fausse :

1. Le principe d'inertie n'est vrai, pour un système ponctuel, que si la quantité de mouvement de ce dernier est conservée par rapport à un référentiel galiléen.
2. Si le vecteur vitesse d'un système ponctuel est nul par rapport à un référentiel galiléen, alors la résultante des forces extérieures qui s'exercent sur lui est nulle.
3. Un pèse-personne subit le poids de ce que l'on pose sur lui.

4. Un système soumis uniquement à l'attraction exercée par un corps se déplace toujours vers celui-ci.

5. Le fait que la durée de révolution d'une planète autour du Soleil soit d'autant plus longue que le rayon de son orbite est plus grand provient simplement de ce que cette orbite est plus grande et offre donc un chemin plus grand à parcourir (on considérera que les orbites sont circulaires).

## Du Tac au Tac

Exercices sans calculatrice. On approximera l'accélération de pesanteur terrestre par  $g = 10 \text{ m.s}^{-2}$ .

1. Un fusil est rigidement fixé à un bloc solide, l'ensemble possédant une masse  $m_1 = 200 \text{ kg}$ . Le bloc est posé sur un support horizontal n'exerçant sur lui aucun frottement (patinoire idéale), et se trouve initialement au repos par rapport au référentiel terrestre, supposé galiléen. Le fusil tire une balle de masse  $m_2 = 20 \text{ g}$ , expulsée du canon avec une vitesse  $v_2 = 300 \text{ m.s}^{-1}$ .

Caractériser qualitativement et quantitativement le mouvement du bloc suite à ce tir.

2. Un parachutiste équipé, de masse totale  $m = 100 \text{ kg}$ , tombe verticalement à une vitesse de valeur constante  $v = 180 \text{ km.h}^{-1}$  par rapport au référentiel terrestre (son parachute n'est pas encore ouvert).

En admettant qu'il est soumis à une force de frottement proportionnelle à la valeur de sa vitesse, déterminer la valeur du coefficient de proportionnalité en question en USI.

3. Une voiture est lancée en ligne droite sur un sol horizontal avec une quantité de mouvement constante de valeur

$p = 1,08.10^5 \text{ kg.km.h}^{-1}$  par rapport au référentiel terrestre, supposé galiléen. Elle freine brutalement, et la résultante des forces de frottement qu'elle subit est constante et a pour valeur  $f = 20 \text{ kN}$ .

Déterminer au bout de quelle durée elle est arrêtée.

4. On considère un satellite décrivant une orbite circulaire centrée sur le centre de la Terre, à une altitude  $z_s = 2,6.10^3 \text{ km}$ .

Montrer que son mouvement est uniforme, et déterminer un ordre de grandeur de sa période de révolution par rapport au référentiel géocentrique.

On donne : le rayon de la Terre  $R_T = 6,4.10^3 \text{ km}$ , sa masse  $M_T = 6,0.10^{24} \text{ kg}$ , la constante de gravitation universelle  $G = 6,7.10^{-11} \text{ USI}$ .

5. Deux satellites, repérés par leurs positions  $S_1$  et  $S_2$ , se trouvent en orbites circulaires autour d'une planète. La période de révolution de  $S_2$  est 27 fois supérieure à celle de  $S_1$ .

Déterminer le rapport des rayons de leurs orbites.

## Vers la prépa

On considère un solide représenté par un point matériel  $M$ , de masse  $m$ , projeté depuis l'origine  $O$  d'un repère orthonormé  $Oxyz$  ( $Oz$  vertical ascendant) solidaire du référentiel terrestre, supposé galiléen. La vitesse initiale  $\vec{v}_0$  communiquée à l'origine des dates ( $t = 0$ ) est contenue dans le plan  $xOz$ . On notera  $\theta$  l'angle qu'elle forme avec l'axe horizontal  $Ox$ . On suppose que ce solide est soumis, tout au long de son mouvement, à une force de frottement fluide colinéaire et opposée à sa vitesse :  $\vec{f}(t) = -\alpha \vec{v}(t)$ .

Montrer que les lois horaires suivantes sont solutions des équations du mouvement :

$$\overrightarrow{OM}(t) = \begin{pmatrix} x(t) = v_0 \tau \cos \theta \left(1 - e^{-\frac{t}{\tau}}\right) \\ y(t) = 0 \\ z(t) = -g\tau t + \tau(v_0 \sin \theta + g\tau) \left(1 - e^{-\frac{t}{\tau}}\right) \end{pmatrix}$$

où  $\tau$  est une constante dont on déterminera l'expression en fonction des paramètres du problème.

En déduire l'équation  $z(x)$  de la trajectoire de ce projectile.

## Halte aux idées reçues

**1.** Le principe d'inertie postule l'existence d'une catégorie particulière de référentiels par rapport auxquels il **énonce une équivalence** entre deux propositions :

- **L'éventuelle conservation de la quantité de mouvement de ce système.** Il s'agit d'une hypothèse cinématique, c'est-à-dire décrivant le mouvement du système. Dans le cas d'un système de masse constante, ceci revient à dire que son vecteur vitesse  $\vec{v}$  est constant, et donc que ce système persévère dans un état de mouvement rectiligne uniforme (direction et valeur de vitesse conservées), voire de repos (cas particulier où  $\vec{v} = \vec{0}$ ).
- **L'éventuelle nullité de la résultante des forces extérieures s'exerçant sur lui.** Il s'agit cette fois d'une hypothèse dynamique, c'est-à-dire décrivant une condition particulière sur les contraintes s'exerçant sur le système.

Or le fait que deux propositions soient équivalentes entre elles ne signifie pas que l'une ou l'autre soit forcément vérifiée, mais seulement qu'elles ont toujours même valeur de vérité (toutes les deux vraies ou toutes les deux fausses, mais toujours en même temps). Si deux propositions sont équivalentes entre elles, alors leurs négations le sont aussi.

Si par exemple nous devons donner le résultat de l'opération  $2 \times 3$ , alors nous pouvons énoncer l'équivalence entre les propositions « Donner la bonne réponse » et « Donner pour résultat 6 ». Cette équivalence ne signifie pas que nous donnerons la bonne réponse, non plus que nous ne donnerons le résultat 6. Nous pouvons très bien donner un résultat qui n'est pas 6, mais dans ce cas nous n'aurons pas donné la bonne réponse.

Ainsi considérons un système ponctuel qui **ne persévère pas** dans un état de repos ou de mouvement rectiligne uniforme par rapport à un référentiel galiléen. Nous pouvons parfaitement lui appliquer le principe d'inertie, qui nous permet d'affirmer que la résultante des forces s'exerçant sur ce système n'est alors PAS nulle.

Réciproquement, si nous considérons un système auquel nous appliquons un ensemble de forces dont la résultante n'est pas nulle, le principe d'inertie nous permet d'affirmer que la quantité de mouvement de ce système n'est PAS conservée.

Il est cependant intéressant de noter que si le principe d'inertie est vérifié même lorsque la quantité de mouvement n'est pas conservée, il n'apporte dans ce cas aucune information détaillée. En effet, si la quantité de mouvement d'un système

n'est pas conservée, la seule chose que nous puissions affirmer, c'est que la résultante des forces est non nulle.

Réciproquement, même si nous connaissons les détails d'une résultante de force non nulle exercée sur un système, la seule prédiction que nous puissions énoncer sur la base du principe d'inertie concernant son mouvement à venir est que sa quantité de mouvement ne sera pas conservée.

Pour aller plus loin et obtenir des informations précises même hors conservation de la quantité de mouvement/compensation des forces, nous avons besoin d'une loi plus puissante, à savoir la deuxième loi de Newton.

**2.** Le principe d'inertie ne permet de conclure à la compensation des forces s'exerçant sur un système, que si la quantité de mouvement de celui-ci est conservée, c'est-à-dire si, à masse constante, il persévère dans un état de repos ou de mouvement rectiligne uniforme par rapport à un référentiel galiléen. Or passer par un état de vitesse nulle n'implique en aucun cas une persévérance dans cet état. On peut par exemple citer le cas d'un projectile tiré verticalement vers le haut : sa vitesse décroît, s'annule, puis il repart vers le bas. À l'instant où sa vitesse est nulle, il est toujours soumis à son poids. Si en revanche il demeurerait suspendu en l'air, nous serions en droit de supposer qu'une nouvelle force est à l'œuvre, qui le retient en compensant les effets du poids que la Terre exerce sur lui.

Notons également qu'à l'inverse, un système peut très bien être soumis à une résultante de force nulle à un certain instant, sans que ceci n'entraîne la nullité de sa vitesse. On peut cette fois citer l'exemple d'un ressort auquel est suspendu une masselotte : si l'on tend le ressort et qu'on le relâche, la masselotte va se mettre à osciller de part et d'autre de sa position d'équilibre. Lorsqu'elle passe par cette position, la résultante des forces qu'elle subit est nulle (c'est pour cela que cette position est une position d'équilibre). Et pourtant, sous l'effet de son inertie, elle passe avec une vitesse non nulle.

**3.** Cette phrase est un non-sens : le poids, par définition, n'est pas une force exercée par un objet quelconque, mais par la Terre sur un objet positionné au voisinage de sa surface.

Cependant, si cette affirmation est fautive en toute rigueur, l'égalité à laquelle elle aboutit est vraie. En effet, si l'on considère un objet de masse  $m$  posé sur un plateau et y demeurant au repos par rapport à un référentiel galiléen, nous pouvons affirmer d'après le principe d'inertie, que les forces auxquelles est soumis cet objet se compensent entre elles.

Or ces forces sont son poids  $\vec{P}_{\text{Terre/objet}}$ , et la réaction normale du pèse-personne (PP) sur cet objet  $\vec{R}_{\text{n,PP/objet}}$ . Leur compensation permet d'écrire :

$$\vec{P}_{\text{Terre/objet}} + \vec{R}_{\text{n,PP/objet}} = \vec{0} \Leftrightarrow \vec{R}_{\text{n,PP/objet}} = -\vec{P}_{\text{Terre/objet}}$$

Par ailleurs, la troisième loi de Newton (principe des actions réciproques) nous permet d'affirmer que la force exercée par le pèse-personne sur l'objet est égale en valeur et direction, et opposée en sens à la force exercée par l'objet sur le pèse-personne, d'où :

$$\vec{F}_{\text{objet/PP}} = -\vec{R}_{\text{n,PP/objet}} = \vec{P}_{\text{Terre/objet}}$$

en utilisant l'égalité fournie précédemment par la première loi de Newton.

Ainsi, un objet persévérant dans un état de repos par rapport à un référentiel galiléen alors qu'il repose sur un pèse-personne exerce sur celui-ci une force qui se trouve être égale à la force exercée par la Terre sur cet objet.

4. Le fait d'être attiré désigne la soumission à l'exercice d'une force attractive. Mais l'exercice d'une telle force ne s'accompagne pas nécessairement d'un déplacement vers son auteur, comme en témoigne toute la variété des mouvements décrits par les systèmes soumis à l'attraction gravitationnelle : depuis le projectile dans sa phase ascendante, aux comètes lorsqu'elles s'éloignent du Soleil.

5. Lorsque l'on examine la troisième loi de Kepler, on voit que le rapport  $\frac{T^2}{a^3}$  est le même pour tous les satellites gravitant autour d'un même objet, où  $T$  représente la période de révolution d'un satellite sur son orbite, et  $a$  le demi-grand axe de cette orbite (soit son rayon, dans le cas d'une orbite circulaire). Cette propriété peut se réécrire sous la forme :

$$\left(\frac{T}{a}\right)^2 \times \frac{1}{a} = \text{cte}$$

Or nous savons que, dans le cas d'une orbite circulaire, le mouvement d'un satellite est uniforme. Nous pouvons encore préciser que la vitesse s'exprime alors comme le rapport de la circonférence de l'orbite, à la période de révolution, soit  $v = \frac{2\pi a}{T}$ . L'égalité vue plus haut nous permet alors d'écrire :

$$\frac{1}{v^2} \times \frac{1}{a} = \text{cte} \Leftrightarrow v \text{ est proportionnelle à } \frac{1}{\sqrt{a}}$$

Nous constatons donc que plus le rayon d'une orbite est important, plus la vitesse d'évolution d'une planète sur cette orbite est faible. La plus grande durée de révolution des planètes lointaines n'est donc pas une affaire d'orbite plus grande à parcourir : non seulement la distance à parcourir est plus importante mais la vitesse de progression sur cette orbite diminue à mesure que la planète considérée orbite plus loin du Soleil.

Qualitativement, ceci vient simplement du fait que la troisième loi de Kepler nous donne une relation de proportionnalité entre le **carré** de la période, et le **cube** du rayon orbital. Ainsi, si celui-ci augmente par exemple d'un facteur 4,  $a^3$  augmente d'un facteur  $4^3 = 64$ , donc  $T^2$  aussi (proportionnalité oblige), ce qui signifie que  $T$  augmente d'un facteur 8 :

l'augmentation de la période est deux fois plus importante que celle du rayon (et donc de la circonférence), du coup la vitesse sur cette orbite est divisée par 2.

## Du Tac au Tac

1. Le problème est considéré par rapport à un référentiel supposé galiléen, donc nous pouvons appliquer entre autres les deux premières lois de Newton. Le système décrit se trouve initialement au repos, et rien n'indique l'action d'une force extérieure (la mise à feu de la cartouche peut être réalisée par un retardateur, par exemple). Nous pouvons donc considérer que la résultante des forces extérieures au système est nulle avant que le coup de feu ne soit tiré, ainsi qu'après puisque les frottements sont négligés devant les autres forces.

D'après le principe d'inertie, nous pouvons donc affirmer que la quantité de mouvement du système {fusil + balle}, évaluée par rapport au référentiel terrestre, est conservée (donc égale au vecteur nul dans le cas présent) puisque l'ensemble est décrit comme initialement au repos. La quantité de mouvement totale d'un système étant égale à la somme des quantités de mouvement de ses diverses parties, si la balle part dans un sens, le reste du système doit partir dans l'autre. Soumis à une résultante de force nulle, le bloc porteur du fusil poursuivra donc sa route en mouvement rectiligne uniforme dans le sens opposé à celui dans lequel est partie la balle, animé de la vitesse  $\vec{v}_2$  qu'il aura acquise juste après le tir. Pour déterminer cette valeur, exprimons la nullité de la quantité de mouvement du système {fusil + balle} juste après le tir :

$$m_1 \vec{v}_1 + m_2 \vec{v}_2 = \vec{0} \Rightarrow v_1 = \frac{m_2}{m_1} v_2 = 3,0 \text{ cm.s}^{-1} = 11 \text{ m.h}^{-1}$$

en passant aux normes des vecteurs vitesses qui sont colinéaires entre eux.

2. Le parachutiste étant animé d'un mouvement rectiligne uniforme par rapport au référentiel terrestre supposé galiléen, nous pouvons affirmer, d'après le principe d'inertie, que la résultante des forces auxquelles il est soumis est égale au vecteur nul.

D'après les informations fournies, la force de frottement peut s'exprimer  $\vec{f} = -\alpha \vec{v}$  où  $\alpha$  est une constante positive (on choisit cette écriture pour mettre en évidence l'opposition de sens entre cette force et le vecteur vitesse du système). La compensation des forces s'écrit alors, en considérant une projection des forces sur un axe vertical orienté vers le bas :

$$\vec{P} + \vec{f} = \vec{0} \Rightarrow mg - \alpha v = 0 \Leftrightarrow \alpha = \frac{mg}{v} = 20 \text{ N.s.m}^{-1}$$

avec  $v = 50,0 \text{ m.s}^{-1}$  après conversion des  $\text{km.h}^{-1}$  en  $\text{m.s}^{-1}$ .

3. La voiture se déplaçant initialement en mouvement rectiligne uniforme par rapport au référentiel terrestre supposé galiléen, nous pouvons affirmer d'après le principe d'inertie que la résultante des forces extérieures auxquelles elle était soumise avant freinage était nulle. Le poids et la réaction normale de la route étant verticaux, et toute autre force étant *a priori* horizontale (forces de frottement permettant à la voiture d'avancer par réaction de la route sur ses roues, et forces

de frottement de l'air), nous pouvons affirmer que le poids et la réaction normale de la route se compensent entre elles.

Lorsque la voiture freine, la route ne sert plus à avancer, mais concourt au contraire au freinage de la voiture, à travers des frottements compris dans les 20 kN donnés par l'énoncé. Le poids et la réaction normale se compensant, tout se passe donc comme si la voiture n'était plus soumise qu'à cette force de frottements, décrite comme constante.

Nous pouvons en déduire, grâce à la deuxième loi de Newton, que le mouvement sera uniformément décéléré. Pour déterminer à quelle date la voiture s'arrête, il suffit d'explicitier cette loi :

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

puis d'en déduire,  $\vec{f}$  étant constante, une loi horaire de variation de la quantité de mouvement par intégration de la relation précédente par rapport au temps :

$$\vec{p}(t) = \vec{f} \times t + \vec{p}(t=0)$$

et pour finir de chercher la date  $t_{\text{ar}}$  à laquelle cette quantité de mouvement est égale à 0, soit :

$$\vec{0} = \vec{p}(t_{\text{ar}}) = \vec{f} \times t_{\text{ar}} + \vec{p}(t=0) \Rightarrow t_{\text{ar}} = \frac{p(t=0)}{f} = 1,5 \text{ s}$$

par passage aux normes des vecteurs (qui sont colinéaires), avec  $p(t=0) = 1,08.10^5 \text{ kg.km.h}^{-1}$  (l'équivalent d'une tonne lancée à  $108 \text{ km.h}^{-1}$ ), soit encore  $3,00.10^4 \text{ kg.m.s}^{-1}$ .

**4.** On considère le système constitué par le satellite, étudié par rapport au référentiel géocentrique supposé galiléen, et soumis à la seule force gravitationnelle exercée par la Terre. Pour démontrer l'uniformité du mouvement circulaire d'un système soumis au seul champ de gravité, on a recours à une base de vecteurs orthonormée mobile avec le satellite. Cette base est appelée base de Frenet, et comprend deux vecteurs. Le premier,  $\vec{T}$ , est dit vecteur tangentiel, et est en permanence tangent à la trajectoire du système, et orienté dans le sens du mouvement de celui-ci (on le définit comme le rapport du vecteur vitesse, à la valeur de la vitesse). Le second, dit vecteur normal  $\vec{N}$ , est perpendiculaire au précédent, en sorte que  $(\vec{T}, \vec{N}) = +\frac{\pi}{2}$ . Notons que selon qu'on observe le système par le dessus ou par le dessous, on peut encore donner deux orientations à  $\vec{N}$ . On prend en général celle qui lui permet de pointer vers la concavité de la trajectoire.

Dans cette base de vecteurs, on peut montrer que l'accélération du système suivi a pour expression :

$$\vec{a} = \frac{dv}{dt} \vec{T} + \frac{v^2}{r} \vec{N}$$

$v$  désigne la valeur de la vitesse, et  $r$  le rayon de courbure de la trajectoire. Il s'agit d'un concept un peu subtil, qui avant le Baccalauréat n'est abordé que dans le cadre du mouvement circulaire, où il se contente alors d'être égal au rayon du cercle décrit par le système au cours de son mouvement.

Le référentiel géocentrique étant supposé galiléen, nous pouvons y appliquer la deuxième de Newton et écrire alors, en notant  $m_s$  la masse du satellite :

$$m_s \left[ \frac{dv}{dt} \vec{T} + \frac{v^2}{r} \vec{N} \right] = -G \frac{M_T m_s}{(TS)^2} \vec{e}_{TS}$$

où  $T$  et  $S$  désignent respectivement le centre de la Terre et le satellite (modélisé par un point matériel à l'échelle de sa trajectoire), et  $\vec{e}_{TS}$  le vecteur unitaire orienté du premier vers le second.

Dans le cas d'un satellite décrivant une orbite circulaire centrée sur le centre de la Terre (hypothèse proposée par l'énoncé), nous pouvons alors établir les relations suivantes :

- $TS = r$ , rayon de l'orbite circulaire décrite par le satellite, égale à la distance séparant celui-ci du centre de la Terre, soit encore  $R_T + z_s = 9,00.10^3 \text{ km}$ , ou  $9,00.10^6 \text{ m}$ .
- $\vec{e}_{TS} = -\vec{N}$ , puisque le vecteur normal, perpendiculaire à la tangente à l'orbite circulaire, pointe nécessairement depuis le satellite (auquel il est rattaché), vers le centre du cercle en question (soit donc ici  $T$ ).

En injectant ces deux relations dans l'égalité précédente, celle-ci devient alors, après simplification par  $m_s$  :

$$\frac{dv}{dt} \vec{T} + \frac{v^2}{r} \vec{N} = G \frac{M_T}{r^2} \vec{N}$$

Les projections de cette égalité vectorielle donnent respectivement :

- Selon  $\vec{T}$  :  $\frac{dv}{dt} = 0$ .
- Selon  $\vec{N}$  :  $\frac{v^2}{r} = G \frac{M_T}{r^2}$ .

Le premier résultat nous permet de conclure immédiatement à l'uniformité du mouvement, la nullité de  $\frac{dv}{dt}$  entraînant la constance de  $v$ .

Le second résultat permet de retrouver la troisième de Kepler (l'une des victoires qui permirent à Newton d'entériner ses modèles). En effet, le mouvement étant uniforme, nous pouvons exprimer la vitesse du satellite comme rapport d'une quelconque distance parcourue par le satellite, à la durée de parcours correspondante. En particulier, dans le cas d'un tour complet, la distance parcourue est  $2\pi r$ . Si l'on note alors  $T_s$  la période de révolution du satellite (attention :  $T$  désigne consécutivement 3 choses dans cet exercice : le centre de la Terre (point), le vecteur tangentiel (vecteur, mais il est alors doté d'une flèche) et la période de révolution (durée, que nous dotons d'un indice  $S$  pour indiquer qu'elle concerne le satellite)), la valeur de la vitesse peut donc s'exprimer  $v = \frac{2\pi r}{T_s}$ .

En injectant cette relation dans l'égalité issue de la projection de la deuxième loi de Newton selon le vecteur  $\vec{N}$ , nous obtenons :

$$\frac{1}{r} \frac{4\pi^2 r^2}{T_s^2} = G \frac{M_T}{r^2} \Leftrightarrow \frac{T_s^2}{r^3} = \frac{4\pi^2}{GM_T}$$

Notons que ce résultat, démontré ici dans le cas particulier d'une orbite circulaire, est encore valable dans le cas général d'une orbite elliptique (on remplace alors  $r$  par  $a$ , longueur du demi-grand axe de l'ellipse).

Il nous reste à calculer la période de révolution  $T_S$ , qui d'après ce qui précède peut s'exprimer :

$$T_S = 2\pi r \sqrt{\frac{r}{GM_T}} \\ = \underbrace{2 \times 3,14 \times 9,10^6}_{\approx 6} \sqrt{\frac{9,10^6}{6,7 \cdot 10^{-11} \times 6,0 \cdot 10^{24}}}$$

Or  $\sqrt{9,10^6} = 3,10^3$ , tandis que  $6,7 \times 6,0 = 40$ , d'où

$$\sqrt{\frac{1}{6,7 \cdot 10^{-11} \times 6,0 \cdot 10^{24}}} = \frac{1}{\sqrt{4,0 \cdot 10^{14}}} = \frac{1}{2,0 \cdot 10^7} \\ = 5,0 \cdot 10^{-8}$$

On obtient ainsi, en effectuant la synthèse de tous ces résultats :

$$T_S = 6 \times 9,10^6 \times 3,10^3 \times 5,0 \cdot 10^{-8} \approx 8,10^3 \text{ s}$$

soit environ 2 h et quart.

5. L'énoncé nous fournit le rapport des périodes  $T_1$  et  $T_2$  des deux satellites :  $\frac{T_2}{T_1} = 27$ .

D'après la troisième loi de Kepler, nous savons que le rapport  $\frac{T^2}{a^3}$  est le même pour tout satellite d'un même astre, soit ici :

$$\frac{T_1^2}{a_1^3} = \frac{T_2^2}{a_2^3} \Leftrightarrow \frac{a_2}{a_1} = \left( \frac{T_2}{T_1} \right)^{\frac{2}{3}} = 9,0$$

en repérant que  $27 = 3,0^3$ .

## Vers la prépa

Introduisons le problème :

- **Système** : le projectile.
- **Référentiel** : terrestre supposé galiléen.
- **Bilan des forces extérieures s'exerçant sur le système** : son poids  $\vec{P} = -mg\vec{e}_z$ , et la force de frottement  $\vec{f} = -\alpha\vec{v}$ .

Le référentiel d'étude étant galiléen, nous pouvons y appliquer la deuxième loi de Newton ( $\vec{a}$  est l'accélération du point  $M$ ) :

$$m\vec{a} = \vec{P} + \vec{f} \Leftrightarrow \begin{pmatrix} \ddot{x} + \frac{\alpha}{m} \dot{x} = 0 \\ \ddot{y} + \frac{\alpha}{m} \dot{y} = 0 \\ \ddot{z} + \frac{\alpha}{m} \dot{z} = -g \end{pmatrix}$$

Le calcul des dérivées première et deuxième des fonctions  $x(t)$ ,  $y(t)$  et  $z(t)$  fournies dans l'énoncé nous donne, pour  $x(t)$  :

$$\dot{x}(t) = v_0 \cos \theta e^{-\frac{t}{\tau}} \Rightarrow \ddot{x}(t) = -\frac{v_0 \cos \theta}{\tau} e^{-\frac{t}{\tau}}$$

En injectant ces résultats dans l'équation portant sur  $x(t)$ , on trouve :

$$-\frac{v_0 \cos \theta}{\tau} e^{-\frac{t}{\tau}} + \frac{\alpha}{m} v_0 \cos \theta e^{-\frac{t}{\tau}} = 0 \\ \Rightarrow v_0 \cos \theta e^{-\frac{t}{\tau}} \left( \frac{\alpha}{m} - \frac{1}{\tau} \right) = 0$$

quelle que soit la date  $t$ .

La solution proposée convient donc, à condition que  $\tau = \frac{m}{\alpha}$ .

Pour  $y(t)$  :

$$\dot{y}(t) = 0 \Rightarrow \ddot{y}(t) = 0$$

La combinaison des deux dernières satisfait évidemment à :

$$\ddot{y} + \frac{\alpha}{m} \dot{y} = 0$$

Pour  $z(t)$  :

$$\dot{z}(t) = v_0 \sin \theta e^{-\frac{t}{\tau}} - g\tau \left( 1 - e^{-\frac{t}{\tau}} \right) \\ \Rightarrow \ddot{z}(t) = -\left( g + \frac{v_0 \sin \theta}{\tau} \right) e^{-\frac{t}{\tau}}$$

En injectant ces résultats dans l'équation portant sur  $z(t)$ , on trouve :

$$-\left( g + \frac{v_0 \sin \theta}{\tau} \right) e^{-\frac{t}{\tau}} + \frac{\alpha}{m} \left[ v_0 \sin \theta e^{-\frac{t}{\tau}} - g\tau \left( 1 - e^{-\frac{t}{\tau}} \right) \right] = -g \\ \Leftrightarrow \left( g + \frac{v_0 \sin \theta}{\tau} \right) e^{-\frac{t}{\tau}} \left( \frac{\alpha\tau}{m} - 1 \right) = g \left( \frac{\alpha\tau}{m} - 1 \right)$$

Le membre de gauche dépendant du temps tandis que celui de droite est constant, cette égalité est vérifiée quelle que soit la date  $t$ , si et seulement si  $\tau = \frac{m}{\alpha}$  (même condition que celle obtenue avec l'équation portant sur  $x(t)$ ).

Les solutions proposées conviennent donc, avec  $\tau = \frac{m}{\alpha}$ .

On remarque que :

- Le mouvement est plan ( $y = \text{cte} = 0$ ).
- Pour  $t \rightarrow \infty$  :
  - on a  $x \rightarrow v_0 \tau \cos \theta = \text{cte}$  : le mouvement horizontal présente une asymptote verticale et corrélativement, on a bien  $\dot{x}(t) \rightarrow 0$  ;

- on a  $\dot{z} \rightarrow -\frac{g}{\tau} = \text{cte}$  : le mouvement vertical tend vers l'uniformité, et est bien orienté vers le bas.

Le mouvement s'achève donc par une chute verticale décrite à vitesse constante, soit un mouvement rectiligne uniforme : la force de frottement compense alors le poids.

Déterminons enfin l'équation de la trajectoire  $z(x)$  ; on commence par isoler  $t$  en fonction de  $x$  ce qui donne, tous

calculs faits :  $t = \tau \ln \left( \frac{1}{1 - \frac{x}{v_0 \tau \cos \theta}} \right)$ . En injectant cette

expression de  $t$  dans l'équation horaire  $z(t)$ , on trouve enfin  $z(x)$  :

$$z(x) = \left( \tan \theta + \frac{g\tau}{v_0 \cos \theta} \right) x - g\tau^2 \ln \left( \frac{1}{1 - \frac{x}{v_0 \tau \cos \theta}} \right)$$

$$\text{avec } \tau = \frac{m}{\alpha}$$

**Remarque :** pour la partie exponentielle de  $z(t)$ , il suffit d'exprimer ladite exponentielle en fonction de  $x$  (sans aller jusqu'à isoler  $t$  proprement dit), et de la remplacer dans cette expression.

## 7.1 Généralités sur les ondes mécaniques progressives (OMP)

### — Qu'est-ce qu'une onde ?

On appelle **onde** le phénomène de propagation de la perturbation d'une grandeur physique, sans transport de matière. Le terme provient d'une manifestation élémentaire du phénomène : la perturbation de la surface de l'eau (l'onde pure dans laquelle se désaltérait l'agneau de Jean de La Fontaine ne répondait à l'origine à aucune équation).

Les ondes peuvent se manifester sous des formes très diverses, dont l'étude exhaustive motive des ouvrages entiers. Nous retiendrons dans cet ouvrage 2 domaines principaux :

- **les ondes mécaniques progressives (OMP)** : sont issues de l'ébranlement d'un milieu matériel ; diverses grandeurs peuvent alors être perturbées (position d'une portion du milieu en question, vitesse de cette portion, pression, masse volumique, etc.), dont la propagation constitue autant d'ondes ;

**Remarque** : le qualificatif « progressives », signifiant que l'onde se propage depuis la source de la perturbation vers le reste du milieu, peut sembler superflue. Vous verrez cependant plus tard qu'une onde peut se réfléchir sur un obstacle ; l'onde ainsi réfléchie, formellement, se déplace de l'extérieur vers la source (on parle alors d'onde régressive), et diffère de l'onde progressive.

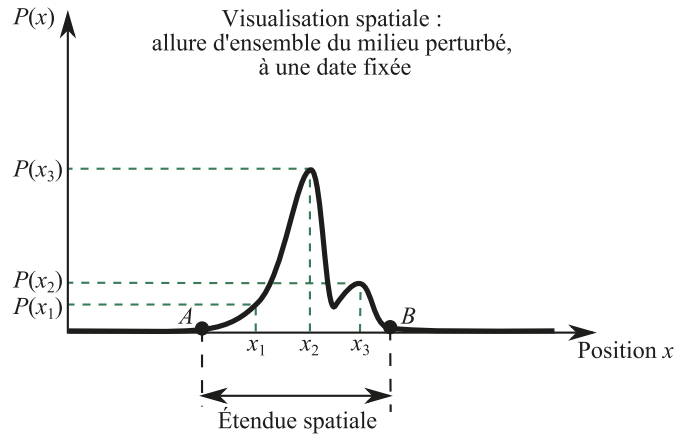
- **les ondes électromagnétiques (OEM)** : sont issues de la perturbation du champ électromagnétique généré par une particule électriquement chargée ; la lumière visible, en particulier, appartient à cette catégorie, dont elle ne constitue cependant qu'une infime partie. Ces ondes peuvent se propager dans le vide (pas besoin de support matériel).

Nous allons dans les pages qui suivent nous efforcer de rester aussi généraux que possible dans la description des phénomènes ondulatoires, et préciserons au cas par cas les spécificités propres aux divers types d'onde.

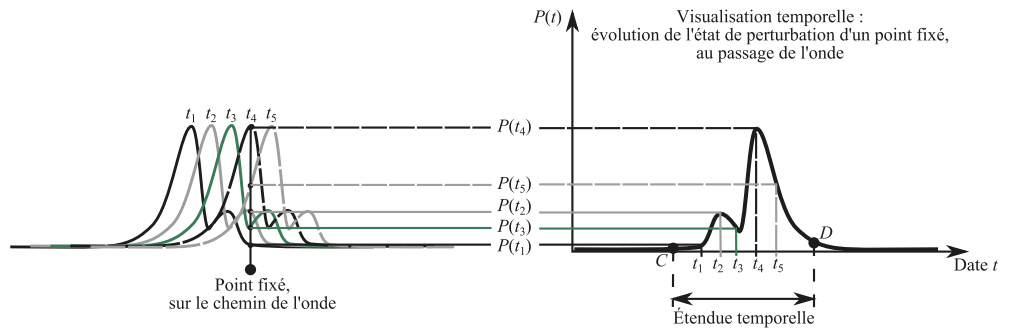
### — Comment peut-on visualiser une onde ?

On peut visualiser une onde de deux façons différentes :

- **Visualisation dans l'espace** : on choisit un instant précis, auquel on observe l'état de la grandeur perturbée, au niveau de chacune des positions situées sur le chemin de l'onde.



- **Visualisation dans le temps** : on choisit une position spécifique du milieu de propagation située sur le chemin de l'onde, et l'on observe l'évolution de son état de perturbation au cours du temps.



**Remarque** : il est aisé de localiser le début et la fin de l'onde en visualisation spatiale. En visualisation temporelle, en revanche, le début du passage de l'onde au niveau du point d'observation correspond au premier instant où ce point est perturbé (sur la gauche du graphe, donc) tandis que la fin correspond au dernier instant où ce point est perturbé (sur la droite).

On appelle alors :

- **Amplitude de l'onde**, la valeur maximale adoptée par la grandeur perturbée.
- **Étendue spatiale de l'onde**, la distance séparant, à une date donnée, la position qui commence à être perturbée, de celle qui finit de l'être ( $x_B - x_A$  sur le premier schéma).
- **Étendue temporelle de l'onde**, la durée pendant laquelle une position donnée est perturbée par le passage de l'onde ( $t_D - t_C$  sur le second graphe).

**Remarque** : notons que A et B sur le premier graphe, correspondent à des positions du milieu de propagation, tandis que C et D sur le second correspondent à des instants. Il n'y a donc pas lieu, à ce stade, de chercher une correspondance entre les deux couples.

## — Comment peut-on caractériser le rythme de déplacement d'une onde ?

En mécanique on caractérise le rythme de déplacement d'un point matériel par sa vitesse. Dans le cas d'une onde, c'est le rythme auquel se déplace la perturbation qui nous intéresse. On appelle alors **célérité**  $v$  d'une onde, le rapport de la distance  $AB$  séparant deux positions

A et B par lesquelles passe l'onde, à la durée  $\tau_{AB} = t_B - t_A$  séparant le passage de la perturbation à l'identique, par ces deux positions.  $\tau_{AB}$  est appelé **retard de l'onde au point B, par rapport au point A**.

$$v = \frac{AB}{\tau_{AB}}$$

|             |                      |
|-------------|----------------------|
| $v$         | en $\text{m.s}^{-1}$ |
| $AB$        | en m                 |
| $\tau_{AB}$ | en s                 |

En particulier, si  $AB$  est égale à l'étendue spatiale de l'onde,  $\tau_{AB}$  correspond à l'étendue temporelle de l'onde. En effet, l'onde commence à passer en A à  $t = t_A$ , et finit d'y passer lorsqu'elle commence à passer en B, c'est-à-dire en  $t = t_B$ , donc  $t_B - t_A$  correspond à la durée  $\tau$  qu'a mis l'onde pour parcourir sa propre étendue spatiale. La célérité d'une onde peut donc s'écrire comme le rapport de l'étendue spatiale  $l$  de l'onde, à son étendue temporelle  $\tau$ .

$$v = \frac{l}{\tau}$$

|        |                      |
|--------|----------------------|
| $v$    | en $\text{m.s}^{-1}$ |
| $l$    | en m                 |
| $\tau$ | en s                 |

## — Une onde se déplace-t-elle dans la même direction que la perturbation qu'elle véhicule ?

Avant de chercher à répondre à cette question, il importe déjà de se demander si parler de la direction de la perturbation a un sens ou non.

Dans le cas des ondes mécaniques progressives (OMP), la perturbation est elle-même de nature mécanique. Elle consiste donc toujours en un ébranlement local de la matière constitutive du milieu de propagation. La propagation résulte alors de la transmission de proche en proche de cet ébranlement : un point, perturbé par l'arrivée de l'onde, perturbe un point voisin, qui va à son tour perturber le suivant, etc.

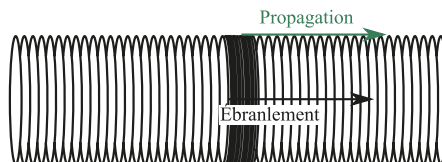
On doit alors distinguer :

- **La direction de perturbation** : c'est la direction selon laquelle est ébranlée la matière au passage de l'onde.
- **La direction de propagation** : c'est la direction selon laquelle se propage l'onde, c'est-à-dire la direction définie par la succession des points perturbés.

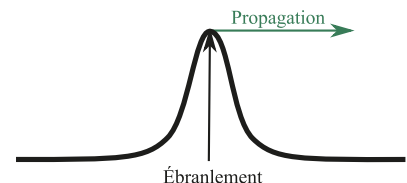
On trouve alors deux cas de figure essentiels :

- Si la direction de perturbation est **parallèle** à la direction de l'onde, l'onde est dite **longitudinale**.
- Si la direction de perturbation est **perpendiculaire** à la direction de l'onde, l'onde est dite **transversale** ou transverse.

Onde longitudinale :  
compression d'un ressort



Onde transverse :  
secousse d'une corde

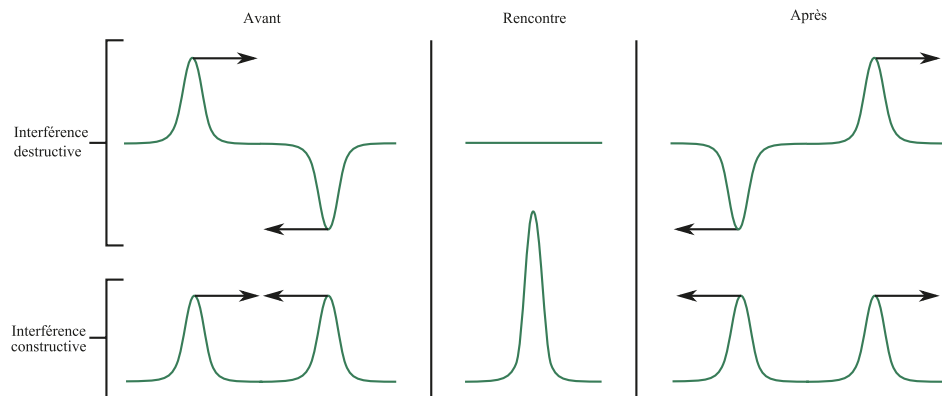


Par ailleurs, dans le cas où la grandeur perturbée est vectorielle (cas du champ électromagnétique, par exemple), on peut de même comparer la direction de la grandeur en question et celle selon laquelle se propage la perturbation. Il s'agit d'un domaine vaste et complexe ; cependant nous limiterons pour le moment ces considérations aux OMP, et concernant celles-ci nous nous contenterons de la différenciation longitudinale/transverse.

## — Deux ondes peuvent-elles coexister en une même position au même instant ?

Les ondes possèdent diverses propriétés fondamentalement différentes de celles des objets matériels. En particulier, là où deux objets ne peuvent se trouver précisément au même endroit au même instant, deux ondes le peuvent. En effet, n'étant par nature que des perturbations, on peut parfaitement envisager qu'une position soit le siège de deux perturbations simultanément. On dit dans ce cas qu'elles se **superposent**, et la perturbation totale s'obtient simplement comme la somme des grandeurs perturbant une même position.

En particulier, lorsque deux ondes de même nature évoluent selon la même direction mais en des sens opposés, elles sont amenées à se croiser. La perturbation d'un point subissant le passage des deux ondes simultanément résulte alors simplement de l'addition des perturbations véhiculées par chacune des deux ondes.



## — Qu'est-ce qu'une onde périodique ?

Un phénomène périodique est un phénomène qui se répète sans cesse à l'identique, à intervalles de temps réguliers. On appelle alors **période  $T$**  d'un tel phénomène, la plus courte durée au bout de laquelle il se reproduit à l'identique.

Ces phénomènes ne constituent pas des raretés : vous verrez en CPGE que toute situation d'équilibre stable peut être modélisée par ce que l'on appelle un **oscillateur harmonique**, dont le comportement est périodique. Des phénomènes de marées aux atomes se désexcitant, ces objets omniprésents en physique perturbent souvent leur environnement, y donnant naissance à des **ondes périodiques** de même période  $T$  que leurs oscillations. On peut ainsi associer à l'onde une période, durée séparant deux répétitions à l'identique de la perturbation qu'elle véhicule, en une position quelconque de son milieu de propagation.

On peut également caractériser une onde par sa fréquence  $\nu$ , représentative du nombre de fois où le phénomène se répète à l'identique par unité de temps et définie comme :

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

$$\begin{array}{ll} v & \text{en Hz} \\ T & \text{en s} \end{array}$$

**Remarque :** la lettre utilisée ci-dessus est la lettre grecque  $\nu$  (prononcer « nu »), à ne pas confondre avec la lettre latine « v » (comme dans « Vitesse »), que vous trouverez souvent dans le même domaine pour désigner la célérité d'une onde (oui, on aurait pu prendre « c » mais elle était déjà réservée à la célérité des OEM dans le vide). Tout ceci peut mener à une certaine confusion, d'autant que parfois la fréquence sera simplement notée  $f$ . Pour vous y retrouver, souvenez-vous du vieil adage : « Surveille tes calculs. Raisonne juste. Conserve ton sens physique. Et surtout, surtout, ne traite jamais avec une équation inhomogène. »

## — Quel intérêt particulier les ondes sinusoïdales présentent-elles ?

Une onde périodique est dite **sinusoïdale** lorsque les variations au cours du temps de la perturbation qu'elle véhicule peuvent être modélisées par une fonction sinusoïdale, autrement dit lorsqu'elles répondent à une expression de la forme :

$$P_A(t) = P_{\max} \sin \left[ 2\pi \left( \frac{t - t_A}{T} \right) \right]$$

$$t, t_A, T \quad \text{en s}$$

- $P_A(t)$  et  $P_{\max}$  ont la même unité : celle de la grandeur dont la perturbation se propageant constitue l'onde (altitude d'une corde horizontale, surpression d'une onde acoustique, température...).
- $P_{\max}$  est l'amplitude de cette onde ; dans le cas d'une propagation sans déformation, elle ne dépend pas de la position  $A$  au niveau de laquelle le passage de l'onde est étudié.
- $t_A$  est le retard de l'onde au niveau de la position  $A$ , par rapport à la source de l'onde : lorsque  $t = t_0 + t_A$ , la perturbation se retrouve en  $A$ , telle que la source l'a générée à la date  $t_0$ .
- Grâce au facteur  $2\pi$ , on retrouve bien la perturbation à l'identique entre deux dates séparées d'un nombre entier de fois la période  $T$ .

On note souvent  $\omega = \frac{2\pi}{T}$  la **pulsation** de l'onde, et  $\varphi_A = -2\pi \frac{t_A}{T}$  la **phase à l'origine des dates** en  $A$ . La perturbation peut alors s'exprimer :

$$P_A(t) = P_{\max} \sin(\omega t + \varphi_A)$$

$$\begin{array}{ll} \omega & \text{en rad.s}^{-1} \\ t & \text{en s} \\ \varphi_A & \text{en rad} \end{array}$$

L'argument du sinus,  $\omega t + \varphi_A$  est appelé **phase** de la perturbation. On remarque alors que si deux positions  $A$  et  $B$  possèdent des phases à l'origine des dates séparées d'un nombre entier de fois  $2\pi$ , la perturbation y prend la même valeur en tout instant :

$$\varphi_B = \varphi_A + 2k\pi \Rightarrow P_B(t) = P_A(t).$$

On dit alors des positions  $A$  et  $B$  qu'elles **vibrent en phase**.

Les ondes sinusoïdales présentent un double intérêt :

- d'une part elles se présentent d'elles-mêmes fréquemment dans la nature : coordonnées cartésiennes d'un système en mouvement circulaire uniforme, solutions des équations différentielles régissant le comportement de ces fameux oscillateurs harmoniques que nous évoquons plus haut... ;
- d'autre part, Joseph Fourier (1768-1830) a montré que toute fonction périodique pouvait se décomposer comme une somme de fonctions sinusoïdales dont les fréquences étaient des multiples entiers de la fréquence de la fonction de départ (fréquence dite *fondamentale*) ;
- ainsi, les fonctions sinusoïdales constituent une **base de fonctions** sur laquelle toute fonction périodique peut être décomposée, et auxquelles toute fonction périodique peut se ramener.

C'est pourquoi, sauf spécification contraire, nous nous limiterons dans la suite à des ondes de ce type.

**Remarque :** vous verrez plus tard qu'il est même possible de décomposer des fonctions non périodiques en somme de sinusoïdes. Il nous faudra cependant pour cela considérer un continuum de fréquences, et passer d'une somme discrète de sinusoïdes à une somme continue, autrement dit une intégrale. Marquez ce jour d'une pierre blanche : vous venez de faire votre premier pas dans le monde merveilleux des transformées de Fourier, outil mathématique extraordinaire dont vous n'avez pas fini d'entendre parler.

## — Comment rendre simultanément compte des dépendances spatiale et temporelle d'une onde sinusoïdale ?

La phase à l'origine des dates en une position  $A$  peut s'exprimer  $\varphi_A = -\omega t_A$ . Or, dans le cas d'une propagation à une dimension, en notant  $x_A$  la position du point  $A$  par rapport à la source de la perturbation, nous pouvons écrire  $t_A = \frac{x_A}{v}$ . Nous en déduisons que la perturbation en un point d'abscisse  $x$ , à la date  $t$ , peut s'exprimer :

$$P(x, t) = P_{\max} \sin \left[ \omega \left( t - \frac{x}{v} \right) \right]$$

|          |                        |
|----------|------------------------|
| $\omega$ | en $\text{rad.s}^{-1}$ |
| $t$      | en s                   |
| $x$      | en m                   |
| $v$      | en $\text{m.s}^{-1}$   |

## — Comment caractériser l'étendue spatiale d'une onde sinusoïdale ?

Une fonction périodique n'ayant à strictement parler ni début, ni fin (le contraire signifierait qu'elle cesse à un moment de se reproduire à l'identique, et perdrait du même coup sa nature périodique et un certain nombre de propriétés accompagnant cette nature), il n'est pas simple de définir son étendue, tant temporelle que spatiale. C'est pourquoi on préfère caractériser son évolution temporelle par sa période.

Or nous avons vu que les représentations spatiale et temporelle avaient par nature la même structure (avec une inversion gauche/droite). Nous pouvons donc anticiper le fait que la repré-

sensation spatiale d'une onde sinusoïdale sera également une sinusoïde, constituée de motifs chacun généré en une durée  $T$  et s'étalant sur une certaine longueur, que l'on pourra utiliser pour caractériser la périodicité spatiale de cette nouvelle sinusoïde.

On appelle ainsi **longueur d'onde  $\lambda$**  d'une onde, la plus courte distance entre deux positions par lesquelles passe cette onde, et qui vibrent en phase.

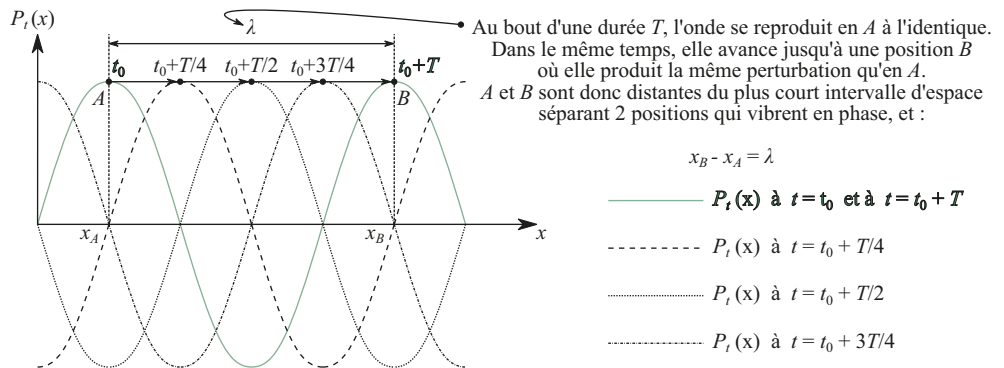
## — Comment les périodicités spatiale et temporelle d'une onde sinusoïdale sont-elles liées ?

Considérons une position  $A$  située sur le chemin d'une onde sinusoïdale, à une date  $t_1$ . À la date  $t_2 = t_1 + T$ , l'onde s'est propagée de façon à reproduire en  $A$  la perturbation telle qu'elle était à la date  $t_1$ . L'onde se reproduisant pour la première fois à l'identique en une position située plus loin, elle a donc avancé de  $\lambda$ .

$$v = \frac{\lambda}{T}$$

|           |                      |
|-----------|----------------------|
| $v$       | en $\text{m.s}^{-1}$ |
| $\lambda$ | en $\text{m}$        |
| $T$       | en $\text{s}$        |

**Remarque :** la célérité d'une onde dépend du milieu de propagation. Or un milieu ne saurait faire varier la durée entre deux répétitions consécutives, laquelle se décide en amont, au seul niveau de la source produisant la perturbation. Ainsi la période d'une onde ne change jamais lors d'un changement de milieu, et par suite sa fréquence non plus. La longueur d'onde, en revanche, est la longueur dont avance l'onde durant une période. Selon la célérité à laquelle elle se propage (et qui rappelons-le est une propriété du milieu de propagation),  $\lambda$  est quant à elle susceptible de changer si l'onde change de milieu.



## — Peut-on définir la forme d'une onde ?

À la base, une onde est un phénomène, et un phénomène n'a pas de forme, seuls les objets qu'il engage peuvent éventuellement en avoir une. Dans notre cas, l'objet affecté par l'onde est le milieu de propagation, qu'il soit matérialisé par un milieu matériel (cas des OMP) ou simplement par la géométrie de l'espace dans lequel l'onde se propage (cas des OEM).

Dans un cas comme dans l'autre, on peut visualiser la propagation de l'onde en observant les positions perturbées de la même manière au même instant. On appelle alors **front d'onde**

d'une onde, le lieu géométrique des positions qui vibrent en phase à un instant donné. La forme de ce front d'onde permet ainsi d'affecter une forme à l'onde. Vous rencontrerez notamment :

- Des ondes dites **planes** : les fronts d'onde sont des plans ; elles constituent un modèle simple, qui notamment décrit bien les ondes sphériques à une distance de la source très grande devant leur longueur d'onde.
- Des ondes dites **circulaires** : ce sont celles spontanément générées par une source ponctuelle dans un milieu isotrope (propagation identique dans toutes les directions) à 2 dimensions.
- Des ondes **sphériques** : même chose, en 3 dimensions.

**Remarque** : dans le cas d'une onde périodique, l'ensemble des positions qui vibrent en phase couvre en réalité plusieurs régions, distantes les unes des autres d'une longueur d'onde. Le front d'onde désigne alors une seule de ces régions, et l'on aura plusieurs fronts d'onde identiques se succédant au cours du temps.

## 7.2 Les phénomènes liés aux ondes

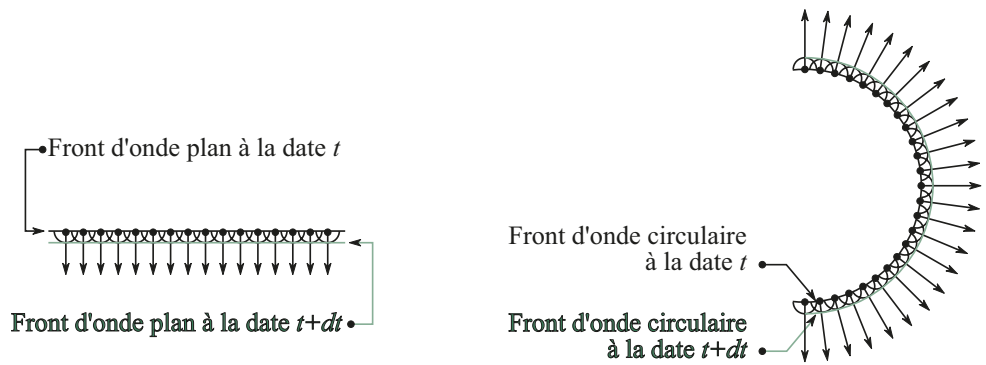
### Quels sont les phénomènes propres aux ondes et comment les interpréter ?

Nous l'avons déjà dit, les ondes constituent des objets physiques particuliers, et extrêmement riches tant le fait de n'être pas assujettis à un transport de matière les rend formellement maléables. Nous ne pourrions malheureusement pas traiter tout le bestiaire des phénomènes auxquels elles peuvent donner lieu, mais aborderons tout de même la liste suivante, dont vous devez être capable de définir chaque item (attention, beaucoup de noms se ressemblent) :

- **Réfraction** : changement de direction de propagation à la faveur d'un changement de milieu.
- **Dispersion** : dépendance de la célérité de propagation vis-à-vis de la fréquence.
- **Interférences** : renforcement/destruction mutuelle de deux ondes se superposant constamment avec la même différence de phase.
- **Diffraction** : modification de la forme du front d'onde et/ou de la répartition d'amplitude d'une onde lorsqu'elle franchit un obstacle.

Beaucoup de ces phénomènes peuvent s'expliquer à partir du principe de Huygens-Fresnel, interprétant la propagation d'une onde comme résultant du fait que chaque point d'un front d'onde constitue une source secondaire qui va son tour émettre une onde sphérique, la superposition des ondes sphériques émises par tous les points du front d'onde générant un nouveau front d'onde, qui va à son tour se comporter comme un ensemble de sources secondaires, etc.

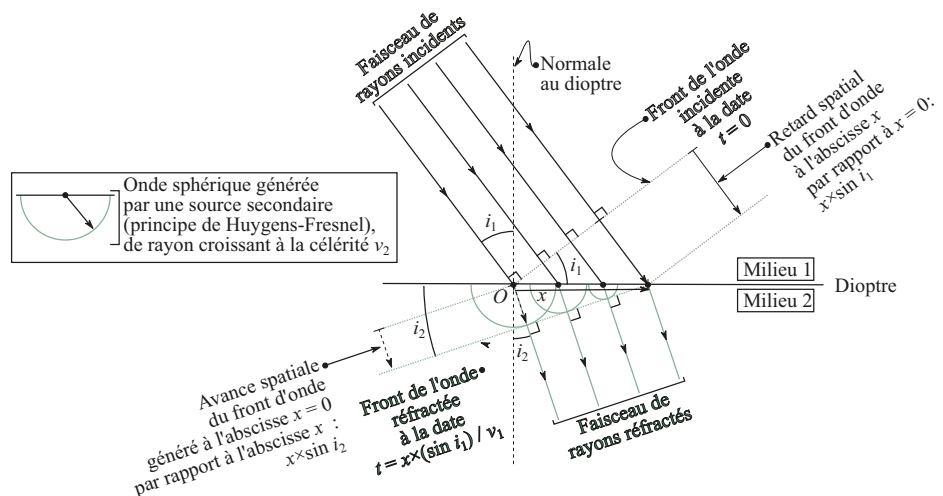
On peut en outre démontrer ce qui apparaît intuitivement sur le schéma ci-dessous, à savoir que la direction locale de propagation de l'onde est perpendiculaire aux fronts d'onde :



## — Comment le phénomène de réfraction peut-il s'interpréter en termes ondulatoires ?

Vous avez probablement vu le phénomène de réfraction en optique géométrique, en vous fondant sur le modèle des rayons lumineux. Il est cependant intéressant de noter qu'il ne concerne pas uniquement la lumière, mais l'ensemble des OEM, ainsi que les ondes mécaniques. Il s'agit donc d'un phénomène propre aux ondes, et que nous pouvons interpréter en termes ondulatoires, à partir du principe de Huygens-Fresnel énoncé ci-dessus.

Si des rayons lumineux parviennent sur un dioptre avec un angle d'incidence  $i_1$ , alors lorsque le rayon de plus proche du dioptre atteint celui-ci, les autres ont encore un peu de chemin à faire. Concrètement, si l'on affecte le dioptre d'un axe  $Ox$  dont l'origine est fixée au niveau du point d'incidence du premier rayon à l'atteindre, on constate que tout rayon atteignant le dioptre en une abscisse  $x$  devra parcourir un chemin supplémentaire  $(\Delta l)_x = x \times \sin i_1$ , subissant un retard  $(\Delta t)_x = \frac{(\Delta l)_x}{v_1}$  par rapport au point situé en  $x = 0$ , (avec  $v_1$  célérité de l'onde dans le premier milieu).



Le principe de Huygens-Fresnel indique donc que le point du dioptre d'abscisse  $x$  se comportera comme une source secondaire, qui commencera à émettre  $(\Delta t)_x$  après le premier point atteint par l'onde (celui d'abscisse  $x = 0$ ). À cette date, tous les points précédents ont déjà émis une onde sphérique de rayon :

$$r_x = v_2 \times (\Delta t)_x = x \times \sin i_1 \times \frac{v_2}{v_1}$$

Or si l'on représente le front d'onde généré par l'ensemble des ondes sphériques émises par les différents points du dioptré et sachant que la direction de propagation est perpendiculaire au front en question, on constate que ces rayons présentent une inclinaison  $i_2$  par rapport à la normale au dioptré, différente de  $i_1$ .

On peut même ajouter que cette inclinaison vérifie l'égalité :

$$\sin i_2 = \frac{r_x}{x} = \sin i_1 \times \frac{v_2}{v_1} \Leftrightarrow \frac{\sin i_1}{v_1} = \frac{\sin i_2}{v_2}$$

En se souvenant que l'indice de réfraction  $n$  d'un milieu est défini comme le rapport de la célérité de la lumière dans le vide à celle dans ce milieu, les célérités dans les deux milieux

s'expriment respectivement  $v_1 = \frac{c}{n_1}$  et  $v_2 = \frac{c}{n_2}$ , et la relation précédente devient simplement :

$$n_1 \times \sin i_1 = n_2 \times \sin i_2$$

Nous constatons ainsi que la loi de Snell-Descartes pour la réfraction peut s'interpréter comme résultant de la différence de célérité des ondes dans les deux milieux :

- si l'onde est moins rapide dans le second milieu que dans le premier ( $n_2 > n_1$ ), l'écart de longueur dû au retard d'un rayon par rapport à un autre va s'amoindrir. Ce moindre écart va redresser le front d'onde, et ce faisant rapprocher les rayons de la normale au dioptré ;
- à l'inverse, si l'onde est plus rapide dans le second milieu que dans le premier ( $n_2 < n_1$ ), l'écart de longueur va s'accroître, et avec lui l'inclinaison des rayons par rapport à la normale au dioptré.

## — Quels peuvent être les effets d'un milieu dispersif sur une onde ?

La célérité d'une onde dépend avant tout de la nature du milieu dans lequel elle se propage. Il arrive cependant que la réponse d'un milieu à une perturbation sinusoïdale soit plus ou moins rapide selon la fréquence de la perturbation en question. La célérité d'une onde dans ce milieu dépend alors de la **fréquence** de cette onde. Toutes les ondes sinusoïdales ne sont alors pas traitées à la même enseigne : celles oscillant à certaines fréquences se propagent plus vite que celles oscillant à d'autres. Ce phénomène de dépendance de la célérité vis-à-vis de la fréquence est appelé **dispersion**, et le milieu de propagation est alors dit **dispersif**. Si la courbe représentative des variations de la longueur d'onde  $\lambda$  d'une onde en fonction de celles de sa période  $T$  est une droite passant par l'origine du repère, alors le rapport  $\frac{\lambda}{T} = v$  est une constante, et le milieu n'est pas dispersif.

Dans le cas contraire,  $v$  dépend de  $T$ , donc de  $\nu$ , et le milieu est dispersif.

Manifestations expérimentales :

- dans le cas d'une OMP, cette dépendance fait que les différentes composantes sinusoïdales vont se déplacer à des célérités différentes, ce qui va inévitablement entraîner une **déformation du front d'onde** ;
- dans le cas d'une onde lumineuse, une telle dépendance signifie automatiquement que l'indice de réfraction du milieu considéré n'est pas le même pour toutes les radiations :

$$v = v(\nu) \Leftrightarrow n = \frac{c}{v} = \frac{c}{v(\nu)} \Leftrightarrow n = n(\nu)$$

Il s'ensuit que si une onde lumineuse polychromatique pénètre un milieu dispersif, l'angle selon lequel est réfractée une radiation dépendra de la fréquence de celle-ci :

$$n_1 \times \sin i_1 = n_2(v) \times \sin i_2 \Rightarrow i_2 = \text{Arcsin} \left[ \frac{n_1}{n_2(v)} \times \sin i_1 \right] \Rightarrow i_2 = i_2(v)$$

En d'autres termes, si un faisceau de lumière polychromatique rencontre un matériau dispersif, les diverses radiations vont poursuivre leur chemin selon des directions différentes, séparant ainsi les unes des autres les diverses composantes de la lumière de départ. Sans doute avez-vous déjà entendu parler de la **dispersion de la lumière blanche** (par un prisme de verre, par exemple) : ce terme est bel et bien utilisé dans le même sens que celui traité ici.

## — Comment peut-on faire interférer deux ondes sinusoïdales entre elles ?

Nous avons vu que lorsque deux ondes se croisent, elles se superposent et peuvent, selon que leurs amplitudes ont le même signe ou non, se renforcer ou au contraire se détruire mutuellement. On parle dans le premier cas d'interférence **constructive**, et dans le second d'interférence **destructive**.

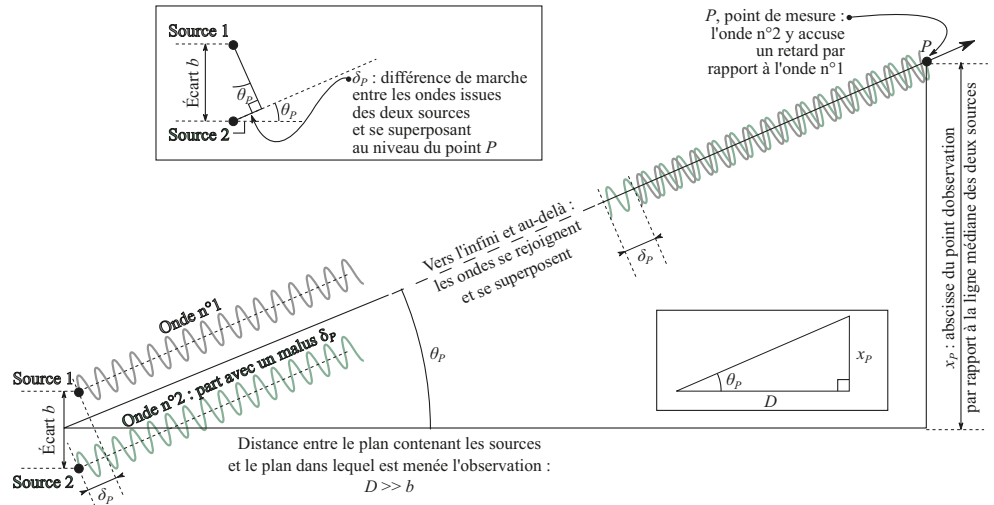
Cependant si les ondes ne font que se croiser, le phénomène est trop fugace pour être analysé. On peut donc proposer d'alimenter constamment le milieu en ondes qui vont se superposer et ce faisant interférer en permanence. Nous étudierons ici le cas de **2 sources sinusoïdales ponctuelles et synchrones**, c'est-à-dire :

- générant des ondes sinusoïdales de même fréquence ;
- non déphasées l'une par rapport à l'autre.

Pour des OMP, ces deux hypothèses suffisent à produire des interférences. Le cas d'OEM est en revanche plus exigeant ; les sources lumineuses émettent en effet des ondes sous forme de **trains d'onde** dont la durée, amplitude, direction de polarisation, etc. varient aléatoirement. Des critères supplémentaires doivent donc être ajoutés à ceux ci-dessus. Nous n'en donnerons pas le détail (complexe et sans utilité à ce stade de votre cursus) ici, mais précisons que lorsqu'ils sont vérifiés, les deux sources sont dites **cohérentes** entre elles.

Dans les faits, obtenir deux sources cohérentes est en général impossible, c'est pourquoi on procède en formant une image d'une source lumineuse, et en faisant interférer la lumière issue de cette image avec celle issue de la source proprement dite (vous découvrirez mille et une astuces dans ce domaine en CPGE). Pour les trous d'Young, on se contente d'éclairer deux trous percés dans un écran au moyen d'une même source lumineuse. Ici encore, la source en question doit répondre à certaines exigences, que satisfont notamment les sources laser ; c'est pourquoi vous verrez souvent celles-ci dans les TP portant sur ce sujet.

**Remarque :** le dispositif des trous d'Young ne laisse passer que très peu de lumière, et les figures d'interférences sont alors très difficiles à observer et à exploiter. On utilise donc en général plutôt les **fentes d'Young**, reposant cette fois sur 2 fentes taillées parallèlement l'une à l'autre. Le résultat est sensiblement équivalent, à quelques nuances près qui vous seront détaillées en CPGE.



Qu'il s'agisse d'OMP ou d'OEM, la mécanique est alors la même :

- par définition, tout point de la médiatrice des sources est équidistant des 2 sources. En conséquence, si les ondes issues de ces deux sources se superposent en un point de cette médiatrice, elles auront parcouru exactement le même chemin, et les sources étant synchrones, elles feront en ce point systématiquement la même chose au même instant : on dit alors d'elles qu'elles oscillent **en phase** l'une avec l'autre au niveau du point d'observation. Les 2 ondes se **renforceront** donc systématiquement sur la médiatrice, qui constituera en permanence un lieu d'**interférences constructives** ;
- si l'on fixe une abscisse sur cette médiatrice et que l'on s'en éloigne transversalement, nous introduisons une **différence de marche** dans les parcours suivis par les deux ondes ;
- à mesure que l'on s'éloigne, la différence de marche se creuse. On constate alors, si l'on se place en représentation spatiale, que les sinusoides se décalent et qu'inévitablement :
  - en certaines positions, l'une est maximale lorsque l'autre est minimale et réciproquement, et l'on dit alors d'elles qu'elles oscillent **en opposition de phase** ; ainsi apparaissent des positions d'**interférences destructives**,
  - après ces positions destructives, l'écart se creusera encore mais, les ondes étant sinusoidales, elles finiront forcément par se retrouver à nouveau en phase, générant une nouvelle position de zone constructive.

On obtient ainsi, dans tout plan parallèle au plan des sources, une série de positions où les ondes interfèrent constructivement, et qui alternent avec d'autres positions où elles interfèrent destructivement.

## — Que se passe-t-il lorsque deux ondes sinusoidales se rencontrent en une position fixée ?

Concrètement, vous apprendrez à montrer que l'énergie moyenne fournie par une onde sinusoidale d'amplitude  $A$  en un point recevant cette onde, est proportionnelle au demi-carré de son amplitude. Pour simplifier, nous allons ici considérer qu'elle y est égale :

$$\langle E \rangle = \frac{1}{2} A^2$$

**Remarque :** dans le cas d'une onde lumineuse, cette énergie est en fait ce que vous avez probablement appris à nommer *intensité lumineuse*. Cependant la même idée se retrouve pour toutes les ondes, puisqu'elles consistent toutes à la base en un déplacement d'énergie de proche en proche. L'idée est donc tout à fait transférable à une OMP, même si elle réclame évidemment quelques adaptations, dont vous apprendrez les finesses en CPGE.

Vous verrez également que si deux sources cohérentes  $S_1$  et  $S_2$  séparées d'une distance  $b$  émettent des ondes sinusoïdales d'amplitudes respectives  $A_1$  et  $A_2$ , alors à grande distance ( $D \gg b \gg \lambda$ ) l'énergie moyenne véhiculée par l'onde issue de leur superposition au niveau d'une position  $P$  a pour expression :

$$\langle E_{1+2} \rangle_t (P) = \frac{1}{2} \left[ A_1^2 + A_2^2 + 2A_1A_2 \times \cos \left( 2\pi \times \frac{\delta_P}{\lambda} \right) \right]$$

où  $\delta_P = (S_2P - S_1P)$  représente la **différence de marche** entre les deux ondes parvenant au niveau de  $P$ , et  $\lambda$  la longueur d'onde des ondes considérées, **dans le milieu de propagation**.

Ce point est particulièrement important concernant les OEM, qu'il est d'usage de caractériser non par leur fréquence mais par leur **longueur d'onde dans le vide**  $\lambda_0$ . Or ces deux longueurs d'onde diffèrent *a priori* l'une de l'autre, et pour un milieu d'indice de réfraction  $n$ , sont liées par l'égalité :

$$\lambda = v \times T = \frac{c}{n} \times T = \frac{\lambda_0}{n}$$

L'expression vue plus haut devient dans ce cas :

$$\langle E_{1+2} \rangle_t (P) = \frac{1}{2} \left[ A_1^2 + A_2^2 + 2A_1A_2 \times \cos \left( 2\pi \times \frac{n \times \delta_P}{\lambda_0} \right) \right]$$

C'est pourquoi, dans le cas d'interférences lumineuses, il est d'usage de travailler avec le **chemin optique** parcouru par l'onde, plutôt que simplement avec la longueur parcourue.

$$L = n \times SP$$

**Remarque :** si une onde traverse successivement plusieurs milieux, on fera alors la somme des produits des longueurs parcourues dans ceux-ci, chacune multipliée par l'indice de réfraction correspondant.

De même, on travaillera dans ce cas avec une **différence de chemin optique**  $\Delta L$  plutôt qu'avec une simple différence de marche. Dans le cas où les deux ondes se propagent dans un même milieu, on a simplement  $\Delta L(P) = n \times \delta_P$ . Dans des situations plus complexes, les chemins optiques respectivement suivis par chacune des deux ondes doivent être calculés séparément, puis différenciés.

Nous retrouvons alors ce que nous avons anticipé plus haut :

- si la différence de marche est égale à un nombre entier de fois la longueur d'onde ( $\delta = k \times \lambda$ ), les deux ondes parvenant en  $P$  vibrent en phase et interfèrent constructivement ; la valeur du cosinus est alors de 1 et l'énergie moyenne reçue en cette position désormais notée  $P_c$  vaut alors :

$$\langle E_{1+2} \rangle_t (P_c) = \frac{1}{2} [A_1^2 + A_2^2 + 2A_1A_2] = \frac{1}{2} (A_1 + A_2)^2$$

Notons qu'en particulier, si les ondes issues deux sources ont même amplitude  $A_1 = A_2 = A$ , alors ce résultat devient :

$$\langle E_{1+2} \rangle_t (P_c) = 4 \times \frac{A^2}{2} = 4 \times \langle E_{1 \text{ ou } 2} \rangle_t (P_c)$$

L'énergie moyenne n'est donc pas simplement le double de l'énergie véhiculée par chacune des ondes comme on aurait pu s'y attendre, mais le quadruple ;

- si la différence de marche est égale à un nombre entier plus une demi-fois la longueur d'onde ( $\delta = \left(k + \frac{1}{2}\right) \times \lambda$ ), les deux ondes parvenant en  $P$  vibrent en opposition de phase et interfèrent destructivement ; la valeur du cosinus est alors de  $-1$  et l'énergie moyenne reçue en cette position désormais notée  $P_d$  vaut alors :

$$\langle E_{l+2} \rangle_t (P_d) = \frac{1}{2} [A_1^2 + A_2^2 - 2A_1A_2] = \frac{1}{2} (A_1 - A_2)^2$$

Notons qu'en particulier, si  $A_1 = A_2 = A$ , alors l'énergie moyenne est cette fois nulle.

Chacun de ces deux résultats, pris isolément, pourrait suggérer une violation de la conservation de l'énergie : plus que la somme dans les zones constructives, moins que la somme dans les zones destructives. Cependant en moyennant ces énergies non plus seulement sur le temps mais sur l'espace, nous retrouvons bien simplement la somme des énergies dispensées respectivement par chacune des deux sources. Simplement, le phénomène d'interférence entraîne une répartition de cette énergie fonction de la différence de marche.

Finalement, deux sources émettant des ondes sinusoïdales de même longueur d'onde  $\lambda$ , et dispensant chacune la même énergie individuelle moyenne  $\langle E \rangle_t$ , génèrent à grande distance une figure d'interférence décrite par :

$$\langle E_{\text{tot}} \rangle_t (P) = \langle E \rangle_t \times \left[ 1 + \cos \left( 2\pi \times \frac{\delta_P}{\lambda} \right) \right]$$

On retiendra également que dans le cas d'ondes lumineuses, on peut remplacer  $\lambda$  par  $\lambda_0$ , à condition de remplacer la différence de marche par la différence de chemin optique entre les deux ondes.

## — Quelle est l'allure de la figure d'interférence complète dans le plan d'observation ?

L'expression ci-dessus permet certes d'apprécier l'évolution de la figure d'interférence en une certaine position  $P$  selon la différence de marche entre les deux ondes parvenant en cette position. Elle réclame cependant le calcul de cette différence de marche, or expérimentalement cette mesure est délicate et il est plus commode de mesurer par exemple la position dans le plan d'observation. Il peut donc être intéressant d'exprimer la relation vue plus haut en fonction de  $x_P$ .

Nous constatons, sur le schéma vu plus haut, que l'angle  $\theta_P$  formé par la direction d'observation (droite passant par le centre du segment liant les 2 sources et le point  $P$  d'observation) avec la médiatrice du segment liant les deux sources intervient dans 2 relations trigonométriques :

$$\sin \theta_P = \frac{\delta_P}{b} \quad \text{et} \quad \tan \theta_P = \frac{x_P}{D}$$

En supposant que l'observation se limite à des positions proches de l'axe ( $x_P \ll D$ ), alors l'angle  $\theta_P$  est très inférieur à 1 rad, nous savons que les valeurs de son sinus et de sa tangente peuvent être assimilés à sa propre valeur en radians, ce qui nous donne :

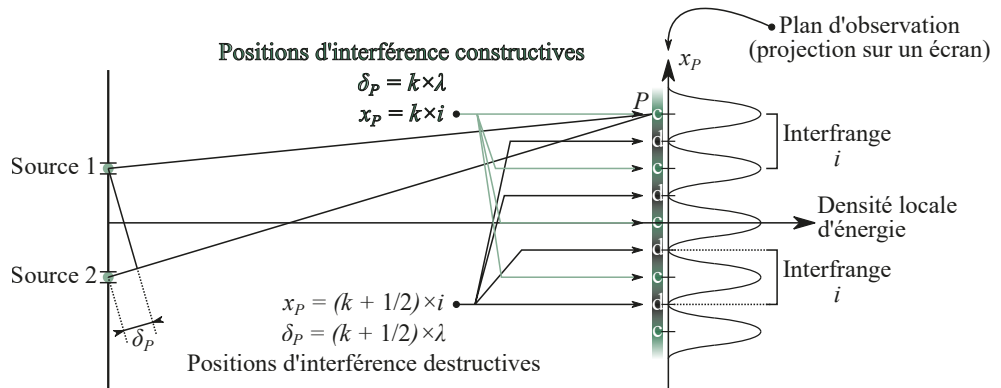
$$\theta_P \approx \frac{x_P}{D} \quad \text{et} \quad \delta_P \approx \theta_P \times b \Rightarrow \delta_P = \frac{b \times x_P}{D}$$

En injectant ce résultat dans la relation décrivant l'interférogramme, nous obtenons alors :

$$\langle E_{\text{tot}} \rangle_i(P) = \langle E \rangle_i \times \left[ 1 + \cos \left( 2\pi \times \frac{b \times x_P}{\lambda \times D} \right) \right]$$

Nous trouvons ainsi les abscisses, dans le plan d'observation :

| Positions d'interférences constructives         | Positions d'interférences destructives                                       |
|-------------------------------------------------|------------------------------------------------------------------------------|
| $\delta_{P,c} = k \times \lambda$               | $\delta_{P,d} = \left( k + \frac{1}{2} \right) \times \lambda$               |
| $x_{P,c} = k \times \frac{\lambda \times D}{b}$ | $x_{P,d} = \left( k + \frac{1}{2} \right) \times \frac{\lambda \times D}{b}$ |



L'interférogramme est donc constitué d'une alternance régulière de positions où les ondes interfèrent constructivement et destructivement. Cette alternance est caractérisée par l'écart entre deux zones consécutives où les ondes interfèrent de la même manière, appelé **interfrange**  $i$ , et qui s'exprime :

$$i = \frac{\lambda \times D}{b}$$

## — Comment les interférences entre deux ondes se manifestent-elles expérimentalement ?

Les interférences se manifestent différemment selon le type d'ondes que l'on fait interférer :

- dans le cas d'ondes à la surface d'un liquide (cuve à ondes), on observera que dans certaines directions, les vagues se propageant à la surface du liquide seront particulièrement creusées, tandis que dans d'autres ce sera le calme plat ;
- dans le cas d'ondes acoustiques (un GBF connecté deux haut-parleurs forment un pendan tout à fait acceptable des trous d'Young en acoustique), on observera que l'intensité sonore est maximale dans certaines directions, et minimale dans d'autres. Le fait que nous percevions un son continu (pour une fréquence située entre 20Hz et 20kHz ne doit pas nous faire perdre de vue (ou d'oreille) que l'air vibre et passe donc bien par une succession de minimum et de maximum (le sinus fait son travail). Cependant le fait que nous entendions ce son particulièrement fort dans certaines directions témoigne d'une amplitude particulièrement élevée, du fait d'un renforcement mutuel des ondes ;

- dans le cas d'OEM, même chose : nous n'observerons pas de zones clignotantes avec plus ou moins d'intensité. Les fréquences des OEM dans le visible sont de l'ordre de  $6.10^{14}$  Hz. Voyez la difficulté que vous pouvez avoir à suivre le défilement des centièmes de secondes sur un chronomètre, et considérez que les oscillations de la lumière se font 6000 milliards de fois plus vite. Nous n'observons donc qu'une intensité lumineuse moyenne, mais encore une fois : le renforcement ou la destruction mutuelle se faisant constamment de la même manière en une position donnée, la valeur moyenne s'en ressentira également de la même façon.

Les valeurs des paramètres que vous rencontrerez dans ce domaine seront typiquement de l'ordre de :

| Type d'onde                   | Fréquence<br>$f = \frac{v}{\lambda}$ | Célérité $v$                | Longueur d'onde<br>$\lambda = \frac{v}{f}$ | Interfrange mesurable $i$ | Rapport<br>$\frac{D}{b} = \frac{i}{\lambda}$ | Valeur réaliste pour $D$ | Valeur réaliste pour $b$ |
|-------------------------------|--------------------------------------|-----------------------------|--------------------------------------------|---------------------------|----------------------------------------------|--------------------------|--------------------------|
| OMP à la surface d'un liquide | $10^2$ Hz                            | $10^{-1}$ m.s <sup>-1</sup> | $10^{-3}$ m                                | $10^{-2}$ m               | 10                                           | 10 cm                    | 1 cm                     |
| Ondes acoustiques             | $10^3$ Hz                            | $10^2$ m.s <sup>-1</sup>    | $10^{-1}$ m                                | 1 m                       | 10                                           | 10 m                     | 1 m                      |
| Ondes lumineuses              | $10^{14}$ Hz                         | $10^8$ m.s <sup>-1</sup>    | $10^{-6}$ m                                | $10^{-3}$ m               | $10^3$                                       | 1 m                      | 1 mm                     |

## — En quoi le phénomène de diffraction consiste-t-il ?

On a coutume de différencier les phénomènes d'interférences et de diffraction l'un de l'autre. S'il est vrai que leurs conditions d'apparition et leurs manifestations diffèrent, il est cependant bon de conserver à l'esprit que, fondamentalement, ils résultent du même processus : la superposition de plusieurs ondes donne, selon la/les différence(s) de marche, une onde résultante tantôt augmentée, tantôt diminuée.

Une nouvelle fois, ce phénomène peut s'interpréter à l'aide du principe de Huygens-Fresnel. Nous avons vu précédemment qu'il était possible d'interpréter la propagation d'une onde en supposant que chaque point du front d'onde associé agissait comme une source secondaire émettant une onde sphérique. La superposition de ces ondes sphériques donnait alors un nouveau front d'onde, l'onde plane poursuivant sa course sous forme d'une onde plane, l'onde sphérique sous forme d'une onde sphérique.

Cependant, que se passe-t-il si l'on interpose sur le chemin d'une onde plane un objet limitant son passage ? Loin des bords, nous allons retrouver une onde localement plane, mais près des bords, l'absence de sources secondaires au-delà d'une certaine limite fait qu'un certain nombre de composantes sphériques vont rester sphériques. L'onde ainsi altérée va donc déborder sur les côtés et former après l'objet un front d'onde dont la forme diffèrera de celle de l'onde originelle.

On devine que ce débordement sera d'autant plus manifeste que les dimensions de l'objet diffractant seront plus petites ; expérimentalement, on observe en effet que la figure de diffraction s'étale d'autant plus que l'objet est plus petit. Cependant pour aller plus loin, il serait nécessaire de calculer précisément les caractéristiques de l'onde ainsi altérée.

Nous pouvons alors envisager de calculer l'onde générée au niveau d'un point d'observation, en effectuant la somme des ondes produites par chacune des sources secondaires évoquées par le principe de Huygens-Fresnel. Il s'agit en substance ce que nous avons fait pour les interférences d'Young ; mais nous devrions cette fois superposer les ondes issues non plus

seulement de 2 sources ponctuelles, mais d'une infinité de sources secondaires, infiniment proches les unes des autres.

On imagine aisément que ce calcul est un brin plus complexe que celui mené dans le cas des fentes d'Young. Il est cependant réalisable (moyennant certaines hypothèses de travail, en particulier une observation à très grande distance de l'objet diffractant). Vous aurez la joie, l'honneur et le privilège de le mener en CPGE (*spoiler-alert* : avec de l'analyse de Fourier dedans), mais au niveau qui nous occupe présentement, cependant, nous devons nous contenter d'en donner les résultats sous forme d'une description de la figure obtenue (figure dite *de diffraction*). Nous traiterons 2 cas :

- **la fente fine rectangulaire** : on entend par *fine* le fait que l'une de ses dimensions est très supérieure à l'autre. Nous pourrions ainsi limiter notre étude à la diffraction produite par les grands bords (faiblement écartés l'un de l'autre), en supposant que la diffraction par les petits bords (très écartés) n'aura que peu d'effet ;
- **le trou circulaire** : les résultats seront similaires, mais la géométrie, plus complexe que celle d'un simple rectangle, entraîne malgré tout des différences que vous devrez connaître.

Il est d'usage de caractériser cette figure par l'une et/ou l'autre des grandeurs suivantes :

- l'angle  $\theta_{\text{dif}}$  formé par la droite passant par le centre de l'objet diffractant et le **premier point d'extinction** de la figure de diffraction, avec l'axe passant par le centre de la fente et perpendiculaire au plan contenant cette dernière (cf. schéma plus bas) ;
- la largeur  $L$  de la **tache centrale** de la figure de diffraction.

On observe alors les résultats qualitatifs suivants :

| Objet diffractant        | Paramètre caractéristique | Figure de diffraction                        |                              |                                                | Influence d'une augmentation de... |                                |                               |
|--------------------------|---------------------------|----------------------------------------------|------------------------------|------------------------------------------------|------------------------------------|--------------------------------|-------------------------------|
|                          |                           | Allure                                       | Tache centrale               | Taches secondaires                             | Longueur d'onde $\lambda$          | Distance $D$                   | Ouverture ( $a$ ou $d$ )      |
| Fente fine rectangulaire | Largeur $a$               | Rectiligne, parallèle au côté de largeur $a$ | Largeur $L$ , très lumineuse | Largeur $L/2$ , luminosité moindre             | La largeur des taches augmente     | La largeur des taches augmente | La largeur des taches diminue |
| Trou circulaire          | Diamètre $d$              | Cercles concentriques                        |                              | Épaisseur $< \frac{L}{2}$ , luminosité moindre |                                    |                                |                               |

On peut alors montrer que pour une observation menée à grande distance de la fente,  $\theta_{\text{dif}}$  vérifie l'égalité :

| Fente rectangulaire de largeur $a$             | Trou cercle circulaire de diamètre $d$                     |
|------------------------------------------------|------------------------------------------------------------|
| $\sin \theta_{\text{dif}} = \frac{\lambda}{a}$ | $\sin \theta_{\text{dif}} = 1,22 \times \frac{\lambda}{d}$ |

Nous constatons par ailleurs que dans les deux cas, cet angle vérifie par construction

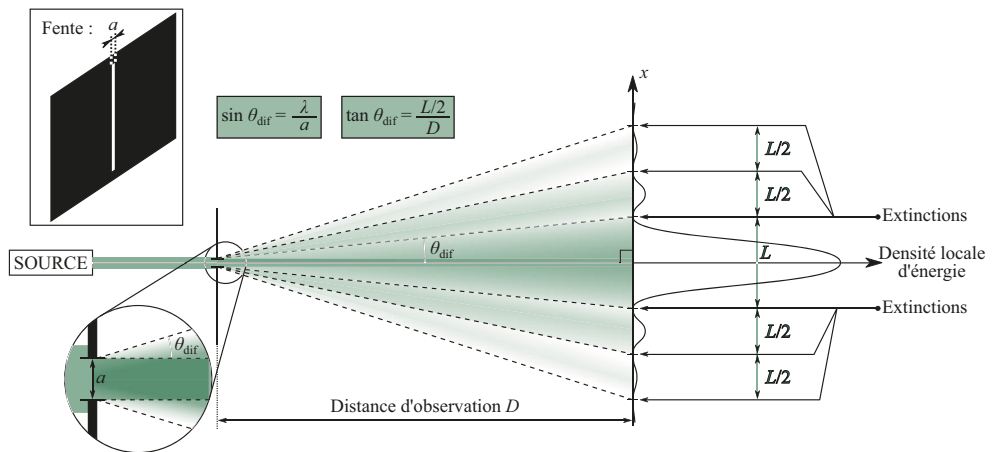
(cf. schéma plus bas) :  $\tan \theta_{\text{dif}} = \frac{L/2}{D}$

Travaillant la plupart du temps dans l'hypothèse des petits angles, nous pouvons alors assimiler le sinus et l'angle  $\theta_{\text{dif}}$  à sa valeur en radians, et égaliser ainsi les deux expressions ci-dessus, ce qui nous donne :

| Fente rectangulaire de largeur $a$         | Trou cercle circulaire de diamètre $d$                 |
|--------------------------------------------|--------------------------------------------------------|
| $\frac{L}{2} = \frac{\lambda \times D}{a}$ | $\frac{L}{2} = 1,22 \times \frac{\lambda \times D}{d}$ |

Nous retrouvons ainsi l'ensemble des appréciations qualitatives portées précédemment concernant l'allure de la figure de diffraction.

**Remarque :** attention, le membre de droite ressemble à s'y méprendre à l'expression de l'interfrange dans le cas des interférences d'Young. Nous avons ici fait le choix d'appeler  $a$  la largeur de la fente dans le cas d'une étude de diffraction, et  $b$  l'écartement entre les fentes dans le cas des interférences d'Young. Cependant rien ne garantit que vous retrouverez ces notations dans tel ou tel exercice. Il importe donc une nouvelle fois de lutter contre la formulette aigüe et de bien réfléchir à la physique de la situation que vous êtes en train de traiter avant de dégainer du calcul à tout va.



**Remarque :** dans le cas du trou circulaire, la première extinction de la fonction de diffraction se calcule à l'aide de fonctions (dites fonctions de Bessel). Il s'agit d'objets mathématiques sophistiqués dont seul le résultat nous importe, et qui aboutissent à ce facteur 1,22. Cette valeur (il s'agit d'ailleurs d'un arrondi), est donc à admettre et à utiliser, sans chercher pour le moment le détail de son origine.

## Comment prendre en compte simultanément les phénomènes d'interférence et de diffraction ?

Lorsque nous avons étudié les interférences d'Young, nous avons supposé que nos sources étaient ponctuelles. Dans les faits, nous savons bien qu'une source ne peut être rigoureusement ponctuelle. En d'autres termes, les sources que nous faisons interférer possèdent une extension spatiale, et vont, en toute rigueur, générer de la diffraction en sus des interférences. Cette situation, complexe de prime abord, aboutit en fait à un résultat assez simple (bien qu'issu de calculs gentiment sophistiqués) : la figure obtenue est en fait constituée de l'interférogramme que donneraient deux sources ponctuelles, modulé par la figure de diffraction individuelle de chacune des deux sources (supposées identiques).

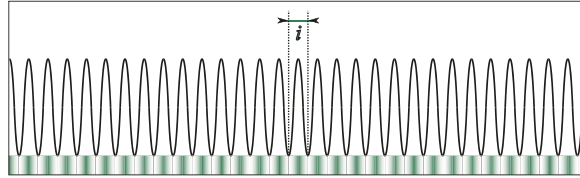
Nous pouvons donc déduire 2 informations d'une telle figure :

- à partir de la largeur  $L$  de la tache centrale de la figure de diffraction, on obtient la taille de chacune des sources :

| Sources rectangulaires                                                                             | Sources circulaires                                                                                               |
|----------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------|
| $\frac{L}{2} = \frac{\lambda \times D}{a} \Leftrightarrow a = 2 \times \frac{\lambda \times D}{L}$ | $\frac{L}{2} = 1,22 \times \frac{\lambda \times D}{d} \Leftrightarrow d = 2,44 \times \frac{\lambda \times D}{L}$ |

- à partir de l'interfrange  $i$  de la figure d'interférence, on obtient l'écartement  $b$  entre les centres des deux sources :  $i = \frac{\lambda \times D}{b} \Leftrightarrow b = \frac{\lambda \times D}{i}$ .

**Remarque :** dans le cas de fentes d'Young, la largeur des fentes est généralement très inférieure à l'écartement entre les fentes. On a donc  $a \ll b$ , et par suite  $L \gg i$ , ce qui permet d'observer plusieurs interfranges de la figure d'interférences à l'intérieur de la tache centrale de la figure de diffraction.



Interférogramme de 2 fentes idéales (infiniment fines), écartées de  $b$

$$i = \frac{\lambda \times D}{b}$$

×

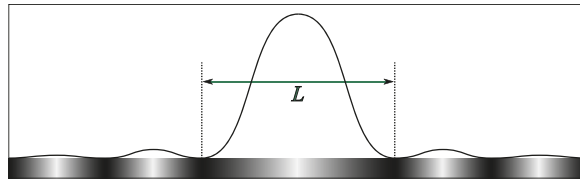
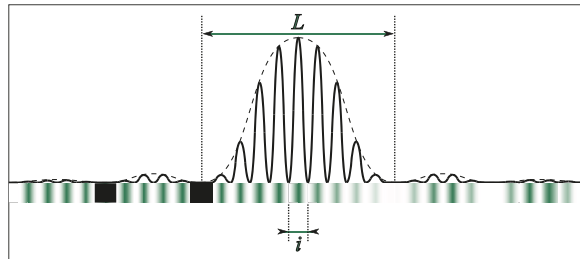


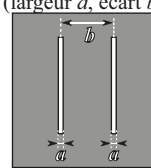
Figure de diffraction d'une fente unique de largeur  $a$

$$\frac{L}{2} = \frac{\lambda \times D}{a}$$

=



Interférogramme de 2 fentes réelles (largeur  $a$ , écart  $b$ ) :



- la mesure de  $L$  permet de remonter à  $a$  ;
- la mesure de  $i$  permet de remonter à  $b$ .

## Halte aux idées reçues

Expliquer, pour chacune des affirmations suivantes, pourquoi elle est fausse :

1. Une onde mécanique progressive est le phénomène de propagation d'une perturbation dans un milieu matériel, de sa source vers l'extérieur, sans déplacement de matière.
2. Lorsque deux ondes interfèrent, les positions constructives sont les lieux où l'onde résultante est constamment à sa valeur maximale.

3. L'étalement d'un faisceau laser lorsqu'il franchit une fente est dû au phénomène de diffraction, non au phénomène d'interférences.

4. Une onde ne peut être diffractée que par un obstacle dont les dimensions sont inférieures ou comparables à sa longueur d'onde.

5. Il n'est possible d'observer la figure lumineuse générée par un objet diffractant, qu'à très grande distance de cet objet diffractant.

## Du Tac au Tac

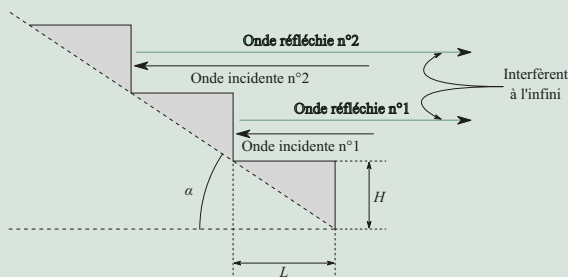
(Exercices sans calculatrice) On prendra pour célérité des ondes acoustiques  $v_a = 3,3 \cdot 10^2 \text{ m.s}^{-1}$  dans l'air,  $v_e = 1,53 \text{ km.s}^{-1}$  dans l'eau, pour la constante de Planck  $h = 6,6 \cdot 10^{-34} \text{ kg.m}^2.\text{s}^{-1}$  et pour la célérité de la lumière dans l'air et dans le vide  $c = 3,0 \cdot 10^8 \text{ m.s}^{-1}$ .

1. Pour éviter qu'un bateau à la dérive n'aille s'échouer n'importe où, une équipe d'artificiers le fait exploser à la surface de la mer. Dans les environs, des plongeurs subaquatiques, qui viennent de terminer un palier de décompression et sont en train de remonter vers la surface, entendent la déflagration alors qu'ils sont sous l'eau. Ils atteignent la surface au bout d'une durée  $\Delta t = 12 \text{ s}$  et entendent de nouveau le bruit de l'explosion.

Déterminer à quelle distance de l'épave se trouvent les plongeurs.

2. Un amphithéâtre de plein air destiné à donner des concerts est organisé en gradins sur une pente inclinée d'un angle  $\alpha = 30^\circ$ , qui réfléchissent les ondes. Pour des questions de qualité acoustique, on souhaite éviter la présence d'ondes réfléchies à certaines fréquences.

Déterminer les dimensions  $L$  et  $H$  que doivent avoir les marches, pour que les ondes de fréquence  $f = 400 \text{ Hz}$  réfléchies par deux marches successives interfèrent destructivement tout en respectant les contraintes ergonomiques des gradins.



Données :

- célérité du son dans l'air :  $v_s = 340 \text{ m.s}^{-1}$  ;
- contraintes ergonomiques : la hauteur des marches doit être comprise entre les valeurs  $H_{\min} = 50 \text{ cm}$  et  $H_{\max} = 70 \text{ cm}$  ;
- aides au calcul : on fera l'approximation  $\sqrt{3} = 1,7$  et on rappelle que  $\frac{1}{16} = 0,0625$ .

3. On considère une pièce de théâtre jouée sur une scène large de  $a = 6,8 \text{ m}$  par des comédiens dont le spectre vocal se situe pour l'essentiel entre 100 et 1 000 Hz.

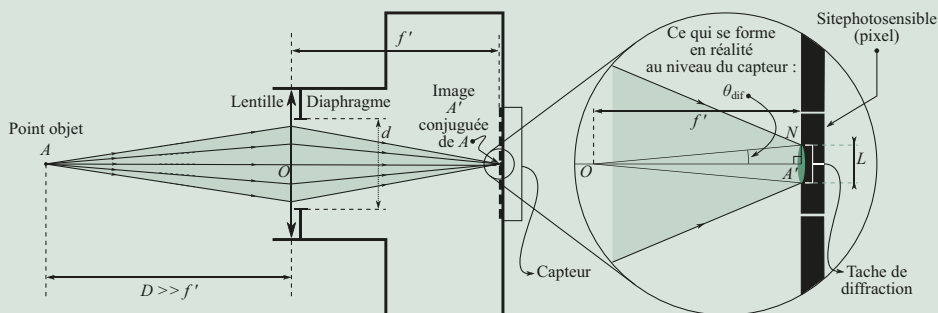
Déterminer de quel angle il est possible d'écarter les sièges par rapport aux côtés de la scène, sans perte significative du discours des comédiens (on supposera un peu abusivement que les ondes acoustiques produites par les comédiens parviennent au bord de la scène sous forme d'ondes planes).

4. On considère deux haut-parleurs, chacun constitué d'une membrane en forme de disque de diamètre  $d = 10 \text{ cm}$ . On note  $A$  et  $B$  leurs centres respectifs, distants de  $AB = 2,0 \text{ m}$ . Les haut-parleurs sont branchés à un GBF délivrant un signal sinusoïdal de fréquence  $f = 1,0 \text{ kHz}$ . Un observateur, se déplaçant selon un axe parallèle à la droite  $(AB)$  et distant de celle-ci d'une distance  $D = 20 \text{ m}$ , constate une alternance entre des zones silencieuses (où il ne perçoit pratiquement aucun son) et d'autres où au contraire il perçoit un niveau sonore maximal.

Expliquer ce phénomène et déterminer l'écart entre deux zones silencieuses.

5. La prise de vue d'une photographie par un appareil photographique numérique (APN) repose essentiellement sur 3 éléments :

- un **capteur**, constitué d'une matrice de récepteurs photosensibles. Chacun de ceux-ci va traduire en un signal électrique la quantité de lumière qu'il a reçue, et le cap-



teur va enregistrer ainsi l'image qui l'a impressionné sous forme d'une carte de valeurs (le fichier image) ;

- un **objectif** formant l'image de l'objet photographié sur le capteur. On peut modéliser cet objectif par une lentille mince convergente de longueur focale  $f'$  associant à chaque point  $A$  de l'objet une image  $A'$ , qui en l'absence de diffraction devrait être ponctuelle ;
- un **diaphragme**, ouverture de diamètre  $d$  réglable destinée à contrôler la quantité de lumière pénétrant l'appareil au moment de la prise de vue. L'ouverture du diaphragme est souvent exprimée par le nombre d'ouverture  $N$ , défini comme :  $N = \frac{f'}{d}$ .
- Déterminer à partir de quelle valeur du nombre d'ouverture  $N$  la taille des taches de diffraction associées aux points

images dépasse les dimensions d'un pixel et entraîne ce faisant une altération de la définition de l'image.

Données :

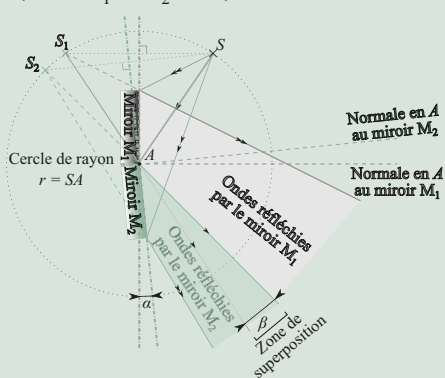
- longueur focale de l'objectif :  $f' = 50 \text{ mm}$  ;
- sujet photographié situé à une distance  $\gg f'$  ;
- longueur d'onde  $\lambda = 500 \text{ nm}$  ;
- dimensions du capteur :  
Côté long (6 000 pixels) :  $L = 36 \text{ mm}$   
Côté court (4 000 pixels) :  $H = 24 \text{ mm}$
- valeurs disponibles du nombre d'ouverture pour le diaphragme :

|     |     |     |     |     |     |     |    |    |    |
|-----|-----|-----|-----|-----|-----|-----|----|----|----|
| $N$ | 1,4 | 2,0 | 2,8 | 4,0 | 5,6 | 8,0 | 11 | 16 | 22 |
|-----|-----|-----|-----|-----|-----|-----|----|----|----|

- approximation :  $1,22 \approx 1,2$ .

## Vers la prépa

Dans la même période où Young inventait l'expérience qui porte son nom en dédoublant une source lumineuse à l'aide de 2 fentes, Fresnel en inventait une autre dédoublant une source (notée  $S$  sur le schéma ci-dessous) à l'aide de 2 miroirs (notés  $M_1$  et  $M_2$  autre) :



$M_1$  et  $M_2$  possèdent une arête commune (dont l'interception avec le plan de représentation est figurée par le point  $A$  sur le schéma ci-dessus) et leurs plans forment l'un avec l'autre un angle  $\alpha$ . Chacun d'eux réfléchit la lumière issue de  $S$ , et tout se passe comme si deux sources lumineuses  $S_1$

et  $S_2$ , images de la source  $S$  respectivement par  $M_1$  et  $M_2$ , génèrent deux faisceaux de rayons lumineux.

1. Justifier que les ondes lumineuses issues de  $S_1$  et de  $S_2$  puissent produire des interférences.

2. Déterminer l'expression de l'angle  $\beta$  de la zone de recouvrement en fonction de l'angle  $\alpha$ .

3. Déterminer comment doit être positionné l'écran d'observation pour obtenir un interfrange  $i = 0,50 \text{ mm}$ , et préciser la largeur  $L$  de la zone du plan d'observation ainsi fixé, sur laquelle on observera des franges d'interférences.

Données :

- distance de la source à l'arête :  $r = SA = 4,0 \text{ cm}$  ;
- longueur d'onde de la source  $S$  :  $\lambda = 633 \text{ nm}$  ;
- angle entre les plans des miroirs :  $\alpha = 1,0^\circ$  ;
- on rappelle qu'en optique géométrique :
  - l'image d'un point lumineux par un miroir plan est le symétrique de ce point par rapport au plan de ce miroir ;
  - le rayon réfléchi par un dioptré est le symétrique du rayon incident par rapport à la normale au dioptré au niveau du point d'incidence.

# Corrigés

## Halte aux idées reçues

**1.** L'erreur ici provient du terme *déplacement*, qui ne doit pas être confondu avec *transport*. En effet, la propagation d'une onde mécanique repose sur la transmission d'un ébranlement de proche en proche, donc à la base de ce phénomène se trouve de toute façon un mouvement de matière, et donc un déplacement. Celui-ci peut entraîner des perturbations de diverses sortes (position des tranches de matière, vitesse de celles-ci, pression à l'intérieur de celles-ci...), mais à la base se trouve toujours une secousse, un choc, ou toute autre forme de mise en mouvement. Il y a donc bien déplacement de matière, autour d'une position d'équilibre.

En revanche il n'y a pas transport de matière, puisque celle-ci n'accompagne pas l'ébranlement tout le long de son chemin, et finit par retrouver sa position de repos lorsque l'excitation initiale prend fin.

**2.** L'erreur commise ici consiste à confondre maximum d'amplitude et maximum tout court. Lorsque deux ondes interfèrent constructivement en une position, cela signifie qu'elles parviennent en phase en cette position précise, autrement dit qu'elles y font toutes deux systématiquement la même chose au même instant : si en cette position l'une est maximale, nulle, minimale... à une certaine date, alors l'autre l'est également à la même date. Cependant le fait pour chacune d'être maximale, nulle minimale... change à chaque instant.

Leur superposition va donc donner une perturbation double de ce que donnait chacune, pour le meilleur comme pour le pire. Si par exemple on considère une position en laquelle deux ondes sinusoïdales de même période  $T$  et de même amplitude  $A$  sont toutes les deux maximales à une date  $t_0$ , alors la perturbation résultant de leur superposition vaudra :

| Date                                        | $t_0$                    | $t_0 + \frac{T}{4}$  | $t_0 + \frac{T}{2}$          | $t_0 + \frac{3T}{4}$ |
|---------------------------------------------|--------------------------|----------------------|------------------------------|----------------------|
| Perturbation véhiculée par la première onde | $A$<br>→ maximale        | $0$<br>→ nulle       | $-A$<br>→ minimale           | $0$<br>→ nulle       |
| Perturbation véhiculée par la deuxième onde | $A$<br>→ maximale        | $0$<br>→ nulle       | $-A$<br>→ minimale           | $0$<br>→ nulle       |
| Perturbation totale                         | $A+A = 2A$<br>→ maximale | $0+0 = 0$<br>→ nulle | $-A - A = -2A$<br>→ minimale | $0+0 = 0$<br>→ nulle |

Le fait pour une position d'être le siège d'interférences constructives ne signifie donc aucunement que la valeur de l'onde y est constamment maximale. Les deux ondes étant sinusoïdales et oscillant à la même fréquence, leur superposition le sera également et ces oscillations la mèneront à alterner entre une valeur minimale et une valeur maximale (c'est un peu le propre d'une sinusoïde). C'est uniquement l'**amplitude** de l'onde sinusoïdale résultante, qui sera maximale.

Les seules positions où la perturbation conservera la même valeur en tout instant seront les positions destructives, à supposer que les ondes s'y superposant aient la même amplitude. Dans le cas contraire, l'onde résultante oscillera en ces lieux également, avec une amplitude  $|A_1 - A_2|$ .

**3.** L'usage fait que l'on étudie séparément les phénomènes de diffraction et d'interférences. Dans le cas des interférences, on considère usuellement deux ondes se superposant en une position de l'espace.

Cependant, rien n'interdit, bien au contraire, d'envisager une superposition à plus de deux ondes : on peut additionner trois, quatre... et même une infinité d'ondes. Or lorsqu'une onde est interceptée par un obstacle (par exemple une fente rectangulaire), chaque point de cet obstacle qui va laisser passer l'onde va se comporter comme une source ponctuelle délivrant une onde sphérique de même fréquence et de même amplitude que l'onde qu'elle a reçue (principe de Huygens-Fresnel).

Le phénomène de diffraction n'est en fait rien d'autre qu'un phénomène d'interférences entre les ondes issues de la multitude de sources ponctuelles générées au niveau de l'objet diffractant, par l'onde parvenant à cet objet. Si cet objet présente un plan de symétrie, alors chaque source ponctuelle d'un côté de ce plan trouvera son symétrique de l'autre côté. En conséquence, la différence de marche entre cette source et son symétrique sera nulle en tout point de ce plan de symétrie. Il s'ensuit qu'en tout point de celui-ci l'onde obtenue résultera de la superposition d'une multitude d'interférences constructives et aboutira au pic central bien connu de la figure de diffraction.

tion. L'existence de pics secondaires est plus complexe à démontrer et réclame des développements mathématiques qui dépassent le cadre de cet ouvrage.

Au final, donc, on ne peut mettre dans deux cases radicalement séparées les phénomènes d'interférences et de diffraction, la seconde étant cousine de la première. On peut en revanche distinguer interférences à deux ondes, et diffraction (interférences à une infinité d'ondes).

4. Comme dit ci-dessus, la diffraction repose sur un phénomène d'interférences, dont on voit mal pourquoi il cesserait d'œuvrer lorsque les dimensions de l'obstacle deviennent grandes devant la longueur d'onde de l'onde considérée. Notons à ce sujet le flou de l'expression « *inférieures ou comparables à sa longueur d'onde* ». Cherchons alors les limites de ce phénomène. Nous savons que pour une fente rectangulaire, le demi-angle au sommet de la zone centrale de diffraction est de l'ordre (en radians) de  $\theta = \frac{\lambda}{a}$ . Nous

constatons donc que si  $a \gg \lambda$  alors  $\theta \ll 1$  rad : les bords du cône de diffraction tendent à être perpendiculaires au plan de l'objet diffractant, c'est-à-dire que l'onde, au lieu de s'étaler sur un cône, suit simplement le couloir délimité par les bords de l'objet diffractant.

Le phénomène est donc moins sensible, mais il ne cesse pas pour autant d'œuvrer. Ici encore, l'énoncé confond deux choses, à savoir le caractère observable d'un phénomène et la validité d'un modèle supposé décrire ledit phénomène.

Répétons-le : une onde est toujours diffractée. Même en l'absence d'objet diffractant, sa seule propagation peut déjà s'interpréter en termes d'interférences grâce au principe de Huygens-Fresnel. Il se trouve simplement qu'en l'absence d'obstacle, ces interférences n'entraînent pas de modification de forme du front d'onde.

Ce qui diffère selon les situations, c'est le modèle à même de décrire proprement la figure de diffraction obtenue. Les relations que nous avons vues concernant le demi-angle au sommet de la tache centrale de la figure de diffraction, sont valables dans le cadre de la diffraction dite de Fraunhofer, et reposent pour être précis sur l'hypothèse du même nom :

$D \gg \frac{d^2}{\lambda}$  en reprenant les notations du cours ( $d$  dimension caractéristique de l'objet diffractant,  $\lambda$  longueur d'onde et  $D$  distance d'observation).

Cette hypothèse, plus exigeante que celle vue dans le cas des interférences d'Young, apparaît spontanément comme nécessaire à la résolution d'une intégrale (connue sous le nom de *formule de Fresnel-Kirchhoff*, que vous verrez en CPGE) dont découlent les expressions de  $\sin \theta_{\text{dif}}$  vues en cours. Donc moins cette hypothèse de Fraunhofer est vérifiée, moins les valeurs de  $\sin \theta_{\text{dif}}$  fournies par les relations vues plus haut seront en adéquation avec les résultats expérimentaux.

Pour résumer : oui, nous observerons une figure de diffraction même à proximité de l'objet diffractant. Mais les caractéristiques de cette figure (directions d'extinctions, largeur de la tache centrale...) ne seront alors pas convenablement décrites par les formules que vous connaissez.

Par ailleurs, l'une des applications essentielles de la diffraction est l'étude de la façon dont elle altère les images fournies par les instruments d'optique. Ces dispositifs posent question, en ce sens que la distance d'observation, si on la résume par exemple à la distance entre un objectif et un écran de projection, ne satisfait en général pas à la relation ci-dessus.

Pourtant la diffraction de Fraunhofer (à laquelle est attachée par exemple, la relation  $\sin \theta_{\text{dif}} = 1,22 \times \frac{\lambda}{d}$ ) est également

valable au voisinage des images formées en optique géométrique. Ceci provient de ce que l'instrument d'optique comporte des lentilles, miroirs, etc. qui vont déformer le front d'onde et avoir un retentissement sur la façon dont la diffraction va l'impacter.

Considérons par exemple un appareil photographique constitué de :

- une lentille mince convergente (l'objectif, longueur focale typique  $f' = 50$  mm) ;
- un diaphragme accolé à celle-ci (pour limiter la quantité de lumière impressionnant le capteur, diamètre courant typiquement 2 à 10 fois plus petit que la longueur focale de l'objectif, nous prendrons  $d = 10$  mm) ;
- un capteur photosensible (sur lequel l'objectif doit former une image nette).

Dans des conditions ordinaires, la distance de l'objet photographié à l'objectif est très supérieure la longueur focale de l'objectif, et l'image doit donc se former à peu de choses près dans le plan focal image de celui-ci. Travaillant à des longueurs d'onde de l'ordre de  $\lambda = 5.10^{-7}$  m, on a alors :  $\frac{d^2}{\lambda} \approx 5.10^3$  m.

La diffraction serait alors convenablement décrite par le modèle de Fraunhofer, uniquement pour un capteur positionné à une distance de l'objectif très supérieure à 5 km.

Pourtant le modèle fonctionne également dans ce contexte, grâce à la présence de la lentille. En effet, on peut modéliser l'ensemble objectif/diaphragme par une succession lentille/diaphragme/lentille accolés :

- la première lentille ( $L_1$ ) ayant une longueur focale égale à la distance séparant l'objet photographié de l'objectif, et formant donc l'image de celui-ci à l'infini ;
- le diaphragme fait son travail de diaphragme, et diffracte la lumière qu'il reçoit d'une source dont tout se passe comme si elle était située à l'infini ;
- la seconde lentille ( $L_2$ ) ayant une longueur focale égale à la distance séparant l'objectif du plan où se forme l'image nette (soit environ  $f'$ ).

On constate ainsi que la présence d'une lentille fait que tout se passe comme si la lumière avait été diffractée par un diaphragme situé à l'infini, puis dans l'autre sens était arrivée de l'infini jusqu'au plan de projection.

Le résultat est que la diffraction de la lumière est convenablement décrite par le modèle de Fraunhofer (donc les expressions de  $\sin \theta_{\text{dif}}$  que nous avons vues peuvent être considérées comme valables) non seulement à l'infini de l'objet

diffRACTANT, mais également lorsque cet objet est accolé à une lentille mince.

## Du Tac au Tac

1. Lorsque l'explosion se produit à la surface de l'eau, elle génère une onde acoustique se propageant dans l'air à la célérité  $v_a$ , et une autre se propageant dans l'eau à la célérité  $v_e$ . Si un auditeur se trouve à une distance  $D$  du lieu de cette explosion, il recevra donc l'onde acoustique au bout d'une durée  $T_e = \frac{D}{v_e}$  s'il se trouve dans l'eau, et au bout d'une durée  $T_a = \frac{D}{v_a}$  s'il se trouve dans l'air.

Le son se propageant environ cinq fois plus vite dans l'eau que dans l'air, le bruit de l'explosion est donc d'abord perçu par les plongeurs immergés, avant des personnes qui attendraient ces plongeurs sur la même verticale mais en surface, par exemple. Dans le cas présent, on nous dit que malgré un délai  $\Delta t = 12$  s de remontée, les plongeurs immergés au départ perçoivent également le bruit en surface. Nous pouvons donc écrire :

$$\Delta t = T_a - T_e = D \left( \frac{1}{v_a} - \frac{1}{v_e} \right) \Rightarrow D = \frac{v_a v_e \Delta t}{v_e - v_a} = 5,0 \text{ km}$$

Notons que l'application numérique demandait quelques astuces. On pouvait d'entrée de jeu voir que  $v_e - v_a = 1,20 \text{ km.s}^{-1}$ .

Numériquement, on obtenait donc  $\frac{\Delta t}{v_e - v_a} = 10 \text{ s}^2.\text{km}^{-1}$ . Restait

à effectuer le produit  $v_a v_e$ . On pouvait le faire à la main mais il était possible d'aller plus vite au prix d'une petite approximation, en remarquant que  $v_a = 0,33 \text{ km.s}^{-1}$ , soit numériquement environ  $1/3 \text{ km.s}^{-1}$ . Le produit donnait alors  $v_a v_e \approx 0,51 \text{ km}^2.\text{s}^{-2}$  qui, multiplié par le terme précédent, donnait 5,1 km. Le calcul rigoureux donne 5,049 qui, après arrondis conformément au nombre de chiffres significatifs autorisé, donne 5,0 km. Cette erreur de 2 % provient de l'amalgame  $0,33 = \frac{1}{3}$ .

2. Il importe comme toujours de commencer par comprendre la physique du problème : l'onde n° 2 parcourt à l'évidence un chemin supplémentaire par rapport à l'onde n° 1, ce qui introduit une différence de marche entre les deux ondes. Celles-ci vont donc interférer de manière constructive ou destructive, selon les valeurs comparées de cette différence de marche et de leur longueur d'onde.

Certaines valeurs de  $L$  vont ainsi permettre d'obtenir une onde retour n° 2 en opposition de phase avec l'onde retour n° 1. La superposition des deux va alors entraîner une interférence destructive. Cependant les valeurs de  $L$  étant directement corrélées à celles de  $H$  via l'inclinaison de la pente, les contraintes sur  $H$  vont se répercuter sur les valeurs autorisées à  $L$ . Reste à savoir si parmi ces valeurs de  $L$  autorisées, certaines permettront d'obtenir les interférences destructives souhaitées.

Commençons donc par le point clef de ce problème, à savoir la différence de marche entre les ondes interférant : en supposant les 2 ondes issues d'une même source, nous voyons que la largeur de la marche au-dessus de laquelle l'onde n° 2 doit passer alors que l'onde n° 1 a déjà été réfléchi, fait que l'onde n° 2 doit effectuer un aller-retour supplémentaire. Il s'ensuit que l'onde n° 2 aura finalement parcouru un chemin supplémentaire égal à  $2 \times L$ , d'où une différence de marche  $\delta$  entre les ondes réfléchies n° 1 et n° 2 :

$$\delta = 2L$$

Nous savons que pour que deux ondes interfèrent de manière destructive, il suffit que la différence de marche entre celles-ci soit égale à :

$$\delta = \left( k + \frac{1}{2} \right) \times \lambda$$

En remplaçant  $\delta$  par  $2L$  et en explicitant  $\lambda = \frac{v_s}{f}$ , il vient :

$$2L = \left( k + \frac{1}{2} \right) \times \frac{v_s}{f}$$

Par ailleurs, la pente étant connue on peut reporter la contrainte sur les valeurs de  $H$ , sur celles de  $L$  :

$$L = \frac{H}{\tan \alpha} \Rightarrow \frac{H_{\min}}{\tan \alpha} = L_{\min} \leq L < L_{\max} = \frac{H_{\max}}{\tan \alpha}$$

Nous en déduisons finalement que s'il est possible de faire interférer destructivement ces deux ondes, alors il doit exister un entier  $k$  qui vérifie les inégalités :

$$\frac{H_{\min}}{\tan \alpha} \leq \left( k + \frac{1}{2} \right) \times \frac{v_s}{2 \times f} < \frac{H_{\max}}{\tan \alpha} \Leftrightarrow$$

$$\left[ \frac{2 \times f}{v_s \times \tan \alpha} \times H_{\min} - \frac{1}{2} \right] \leq k < \left[ \frac{2 \times f}{v_s \times \tan \alpha} \times H_{\max} - \frac{1}{2} \right]$$

Calculons alors la valeur du terme :

$$\frac{2 \times f}{v_s \times \tan \alpha} = \frac{2 \times 400 \text{ Hz}}{340 \text{ m.s}^{-1} \times \frac{1}{1,7}} = 4,0 \text{ m}^{-1}$$

$$\text{avec } \tan 30^\circ = \frac{\sqrt{3}}{3} = \frac{1}{\sqrt{3}}.$$

Nous obtenons alors, avec les valeurs  $H_{\min} = 50 \text{ cm}$  et  $H_{\max} = 70 \text{ cm}$  :

$$\left[ 2,0 - \frac{1}{2} \right] = 1,5 \leq k < 2,3 = \left[ 2,8 - \frac{1}{2} \right]$$

Il existe donc bien une solution (unique) qui est  $k = 2$ , d'où nous déduisons :

$$\begin{aligned} 2L &= \left( k + \frac{1}{2} \right) \times \frac{v_s}{f} \Leftrightarrow L = \left( k + \frac{1}{2} \right) \times \frac{v_s}{2f} \\ &= 2,5 \times \frac{340 \text{ m.s}^{-1}}{2 \times 400 \text{ Hz}} = \frac{17}{16} \text{ m} = \left( 1 + \frac{1}{16} \right) \text{ m} = 1,1 \text{ m} \end{aligned}$$

Correspondant à la hauteur :

$$H = L \times \tan \alpha = \frac{17}{16} \text{ m} \times \frac{1}{1,7} = 10 \times \frac{1}{16} \text{ m} = 63 \text{ cm}$$

3. Nous savons que le cône de diffraction d'une onde par une ouverture rectangulaire de largeur  $a$  a pour demi-angle au sommet  $\theta = \frac{\lambda}{a} = \frac{v_a}{af}$ , où  $\lambda$  représente la longueur d'onde de l'onde considérée,  $f$  sa fréquence,  $v_a$  sa célérité (celle des ondes acoustiques dans l'air dans le cas présent) et où  $\theta$  est exprimé en radians.

L'objectif ici est de savoir si l'on peut ou non placer des sièges qui ne soient pas exactement face à la scène. Il importe de comprendre que cette possibilité est favorisée, et non empêchée par la diffraction. En l'absence de ce phénomène, les ondes acoustiques (en supposant des ondes planes) poursuivraient leur chemin en suivant l'allée délimitée au départ par les bords de la scène, sans qu'aucun son ne parte sur les côtés. Ainsi, plus la diffraction est manifeste, plus  $\theta$  est important et plus il est possible de placer des spectateurs loin sur les côtés.

D'après la relation donnée plus haut, nous voyons que les sons les moins diffractés sont ceux de haute fréquence, qui sont plus directifs. Nous allons donc limiter notre étude à ceux-ci, puisque s'ils passent dans une direction, les sons plus graves passent aussi.

Pour une fréquence  $f = 1\,000$  Hz, nous trouvons ainsi  $\theta = \frac{v_a}{af} = 5,0 \cdot 10^{-2}$  rad. En approximant  $\pi$  par la valeur 3, nous voyons

que 1 rad équivaut à peu près à  $60^\circ$ , et  $5,0 \cdot 10^{-2}$  rad correspondent donc à environ  $3^\circ$ . Dans les hypothèses de l'exercice, la marge angulaire est donc très faible.

Notons malgré tout que si un siège mesure par exemple 0,5 m, on peut envisager de placer une colonne de sièges supplémentaires par tranche de 10 m d'éloignement de la scène. En effet, l'angle étant petit devant  $2\pi$  rad, nous pouvons assimiler la valeur en radians de  $\theta$  et sa tangente. En notant alors  $D$  la distance à la scène et  $e$  la largeur du siège, nous pouvons écrire que  $\frac{e}{D} = \tan \theta \approx \theta_{\text{rad}} = 0,05$ . Avec  $e = 0,5$  m, on voit que  $D = 10$  m.

Précisons toutefois que les ondes acoustiques produites par les comédiens ne parviennent en général pas au bord de la scène sous forme d'ondes planes, ce qui explique qu'en pratique on place plus de spectateurs sur les côtés que ne le laisseraient penser les résultats ci-dessus. Notons enfin que dans les théâtres antiques, les spectateurs étaient répartis sur un hémicycle quasi complet mais, se situant en extérieur, la scène n'était pas limitée par des obstacles.

4. Nous savons que l'interfrange dans un interférogamme d'Young a pour expression  $i = \frac{\lambda D}{b}$ , où  $\lambda$  est la longueur d'onde des ondes considérées,  $D$  la distance séparant le plan d'observation et la droite  $(AB)$  et  $b$  l'écartement entre les centres des sources (soit ici  $AB = 2,0$  m).

Lorsque l'observateur se déplace parallèlement à  $AB$ , ses oreilles passent successivement par des zones où les ondes acoustiques se renforcent (interférences constructives), provoquant une sensation de niveau sonore maximal, et d'autres où

elles s'annulent (interférences destructives), provoquant une sensation de quasi-silence.

L'écart entre deux zones silencieuses n'est autre que l'interfrange cité précédemment, dont la valeur est ici, avec  $\lambda = \frac{v_a}{f} = 0,33$  m,  $D = 20$  m et  $b = 2,0$  m :

$$i = \frac{\lambda D}{b} = 3,3 \text{ m.}$$

Cette expérience peut également être menée avec des ondes ultrasonores (typiquement à la fréquence  $f = 40$  kHz, soit une longueur d'onde d'environ 8 cm), mais dans ce cas l'oreille humaine ne suffit plus (il s'agit d'ultrasons) et les alternances doivent être visualisées au moyen d'un récepteur dédié, relié à un oscilloscope.

Notons également que, dans ce dernier cas, les émetteurs étant usuellement plus petits (puisque la longueur d'onde l'est aussi, on peut réduire leurs dimensions sans trop avoir à souffrir de la diffraction), l'écartement entre les sources peut être réduit (pour des émetteurs de 1 cm de diamètre, un écartement de 10 cm suffit) et la distance d'observation également (elle doit être nettement supérieure à l'écartement, mais ici 1 m suffit). Cette alternative présente certes l'avantage d'un moindre encombrement, mais on perd alors le caractère audible de l'expérience.

5. Nous nous trouvons ici dans le cas de la diffraction au voisinage de l'image géométrique d'un objet, dont nous savons qu'elle est également décrite par la diffraction de Fraunhofer (cf. « Halte aux idées reçues » n° 5). L'angle  $\theta_{\text{dif}}$  répond donc à la relation :

$$\sin(\theta_{\text{dif}}) = 1,2 \times \frac{\lambda}{d}$$

Cette relation va évidemment nous servir, mais pour mener l'exercice sans calculatrice, le sinus risque de poser un problème. Nous savons qu'il peut être assimilé à la valeur en radians de son argument dès lors que celle-ci est très inférieure à 1 rad.

Commençons donc par vérifier ce point ; l'ouverture  $d$  du diaphragme peut s'exprimer en fonction du nombre d'ouverture  $N$  comme :

$$N = \frac{f'}{d} \Leftrightarrow d = \frac{f'}{N}$$

Donc même en prenant la plus grande valeur de  $N$  proposée ( $N = 22$ , que nous majorerons à 25 par précaution), la plus petite valeur de  $d$  serait de l'ordre de :

$$d = \frac{f'}{N} = 2 \text{ mm}$$

Comme par ailleurs l'énoncé évoque une longueur d'onde de  $5 \cdot 10^{-7}$  m (disons  $10^{-6}$  m par précaution), la valeur de  $\sin(\theta_{\text{dif}})$  serait donc au grand maximum de l'ordre de :

$$\sin(\theta_{\text{dif}}) = 1,2 \times \frac{\lambda}{d} \approx 1,2 \times \frac{10^{-6} \text{ m}}{2 \cdot 10^{-3} \text{ m}} = 6 \cdot 10^{-4} \ll 1$$

Donc même en exagérant dans le sens défavorable (grandes valeurs de  $N$  et de  $\lambda$  abondent dans le sens d'une grande

valeur de  $\sin(\theta_{\text{dif}})$ ), il apparaît tout à fait licite d'assimiler le sinus de  $\theta_{\text{dif}}$  (et dans la foulée, sa tangente) à la valeur en radians de l'angle lui-même :

$$\theta_{\text{dif}} \approx \sin \theta_{\text{dif}} \approx \tan \theta_{\text{dif}}$$

Ce préliminaire étant établi, revenons à la physique du problème : là où l'optique géométrique devrait nous fournir des points lumineux, la diffraction nous donne des taches lumineuses dont le demi-angle au sommet  $\theta_{\text{dif}}$  (en fixant ledit sommet au centre optique de la lentille) répond à l'expression :

$$\theta_{\text{dif}} = 1,2 \times \frac{\lambda}{d}$$

Projetées sur le capteur numérique, les taches ainsi obtenues risquent de s'étaler sur plusieurs pixels, altérant l'image d'un flou d'autant plus prononcé que les taches seront plus larges. Or elles seront d'autant plus larges que le diaphragme sera plus fermé.

Cependant, le capteur étant constitué de multiples sites photosensibles, l'exercice propose de considérer que ce flou ne sera pas significatif tant que les taches de diffraction resteront plus petites que les sites photosensibles eux-mêmes.

Poursuivons donc en exprimant la largeur de l'une de ces taches de diffraction sur le capteur. L'énoncé mentionne que le sujet photographié se situe à une distance de l'objectif, très supérieure à sa longueur focale. Nous pouvons donc approximer cette situation par un objet situé à l'infini, dont l'image géométrique se formera directement dans le plan focal de l'objectif. En d'autres termes, le capteur doit donc être positionné à une distance égale à  $f'$  (information que, dans leur générosité légendaire, les auteurs avaient laissé planer sur le schéma de l'énoncé).

Nous disposons donc de 2 relations simples :

- d'une part, l'expression de la largeur angulaire qui devient :

$$\theta_{\text{dif}} = 1,2 \times \frac{\lambda}{d}$$

- d'autre part, dans le triangle  $OA'N$  rectangle en  $A'$  :

$$\theta_{\text{dif}} \approx \tan \theta_{\text{dif}} = \frac{L/2}{f'}$$

En combinant ces deux expressions, nous obtenons une relation liant directement la largeur de la tache de diffraction au nombre d'ouverture :

$$\frac{L}{2f'} = 1,2 \times \frac{\lambda}{d} \Leftrightarrow L = 2,4 \times \frac{\lambda \times f'}{d} \Leftrightarrow L = 2,4 \times \lambda \times N$$

Nous constatons que si  $N$  est grand, alors  $d = \frac{f'}{N}$  est petit (le diaphragme est donc peu ouvert), entraînant bien une tache de diffraction large.

Il ne nous reste plus qu'à calculer les dimensions d'un pixel ; nous constatons qu'elles sont égales sur les deux côtés (les pixels sont donc carrés) :

$$e = \frac{L}{6000} = \frac{H}{4000} = 6,0 \mu\text{m}$$

Il ne nous reste plus alors qu'à exprimer la contrainte demandée sur la taille de la tache :

$$L < e \Leftrightarrow 2,4 \times \lambda \times N < e \Leftrightarrow N < \frac{e}{2,4 \times \lambda} = \frac{6,0 \cdot 10^{-6} \text{ m}}{2,4 \times 5,0 \cdot 10^{-7} \text{ m}} = 5,0$$

Nous constatons donc qu'à cette longueur d'onde, au regard du critère proposé, la diffraction sera négligeable jusqu'à un nombre d'ouverture de  $N_{\text{max}} = 4,0$ . À partir de  $N = 5,6$  (donc pour un diaphragme ouvert d'environ  $\frac{f'}{5} = 1 \text{ cm}$ ), les taches de diffraction seront plus larges que les pixels du capteur.

## Vers la prépa

Cet exercice n'est pas plus compliqué sur le plan de la physique qu'un bon exercice de terminale (« Il était une fois, deux sources cohérentes dont les ondes se superposaient d'un amour tendre, elles furent heureuses et eurent beaucoup d'interfranges... »). Après tout l'étudiant(e) en CPGE sort tout juste de terminale et sa progression sur les flancs de la montagne de la connaissance doit tout de même être progressive. La difficulté est ici plutôt formelle, le dispositif étant un peu plus compliqué à mettre en équation que celui des trous ou des fentes d'Young.

Prenons donc les questions dans l'ordre et commençons à échauffer gentiment l'aire cérébrale que nous dédions à la géométrie...

**1.** Produire des interférences sur la base d'ondes issues de deux sources réclame deux conditions :

- les sources doivent être cohérentes ;
- il doit exister une région de l'espace où les ondes qu'elles produisent se superposent.

Concernant la première condition, l'énoncé précise que tout se passe comme si les deux faisceaux de lumière obtenus étaient issus de deux sources  $S_1$  et  $S_2$ , images respectives de la source  $S$  par les miroirs  $M_1$  et  $M_2$ . Ces deux sources virtuelles reproduisent donc à chaque instant les ondes générées par  $S$  au même instant ; elles sont donc toutes deux cohérentes avec  $S$ , et par suite également cohérentes entre elles.

Concernant la seconde condition, nous voyons sur le schéma que le fait pour les miroirs d'être inclinés l'un par rapport à l'autre (avec une concavité tournée vers  $S$ ) amène une partie de chacun des deux faisceaux à recouvrir l'autre, donnant donc un lieu où les ondes issues des deux sources secondaires vont pouvoir se superposer.

**2.** Nous abordons ici la partie mathématique du problème. Les données de l'énoncé nous rappellent que les rayons réfléchis par un miroir sont symétriques des rayons incidents leur donnant naissance par rapport à la normale à ce miroir au niveau du point d'incidence.

Il faut un angle pour attraper un angle, c'est bien connu, et nous sentons ici pointer l'utilité de définir quelques angles d'incidence et de réflexion. L'angle  $\beta$  demandé étant formé

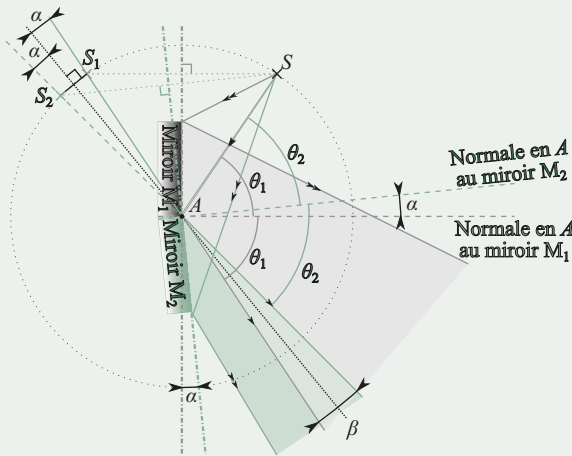
par les rayons issus de  $S$  et réfléchis au point  $A$ , respectivement par chacun des deux miroirs, il semble utile d'étudier les angles d'incidence et de réflexion associés à ces rayons. Nous les noterons respectivement  $\theta_1$  pour le couple incident/réfléchi au niveau du miroir  $M_1$ , et  $\theta_2$  pour le couple incident/réfléchi au niveau du miroir  $M_2$ . Nous constatons alors que :

- les normales aux deux miroirs étant inclinées d'un même angle (par nature  $90^\circ$ ) dans le même sens par rapport aux plans de leurs miroir respectifs, elles forment donc entre elles le même angle que forment entre eux les miroirs, autrement dit  $\alpha$ ;
- les rayons incidents issus de  $S$  et atteignant les deux miroirs au point  $A$  étant initialement confondus, nous pouvons alors écrire, d'après le schéma ci-dessus :  $\theta_2 = \theta_1 - \alpha$ ;
- en nous penchant sur les rayons réfléchis, nous constatons cette fois que  $\beta$  vérifie :

$$\beta + \theta_2 = \theta_1 + \alpha \Leftrightarrow \beta = \theta_1 - \theta_2 + \alpha$$

- en substituant alors  $\theta_2$  par l'expression vue plus haut, nous obtenons :

$$\beta = \theta_1 - (\theta_1 - \alpha) + \alpha \Leftrightarrow \beta = 2\alpha$$



3. Nous savons alors que l'interférogramme présentera un interfrange répondant à l'égalité :

$$i = \frac{\lambda \times D}{b}$$

avec :

- un plan d'observation parallèle au plan contenant les sources (ici  $S_1$  et  $S_2$ ), ou si l'on préfère, perpendiculaire à la médiatrice du segment liant ces sources ;
- des angles d'observation faibles (ce qui sera nécessairement vérifié ici, la zone de recouvrement où la figure d'interférences est susceptible de se former étant limitée par un angle au centre

$$\beta = 2\alpha = 2,0^\circ = 2,0^\circ \times \frac{\pi \text{ rad}}{180^\circ} \approx 0,03 \text{ rad} \ll 1 \text{ rad}$$

- $b$ , la distance séparant les deux sources, que nous pouvons

$$\text{ici exprimer comme : } \sin \alpha = \frac{S_1 S_2}{2r} \Leftrightarrow b = 2r \times \sin \alpha$$

Nous trouvons ainsi que la distance d'observation requise pour satisfaire à la condition imposée par l'énoncé sur  $i$  est donnée par l'égalité :

$$\begin{aligned} i &= \frac{\lambda \times D}{2r \times \sin \alpha} \Leftrightarrow D = \frac{2r \times i \times \sin \alpha}{\lambda} \\ &= \frac{2 \times 4,0 \cdot 10^{-2} \text{ m} \times 5,0 \cdot 10^{-4} \text{ m} \times \sin 1,0^\circ}{6,33 \cdot 10^{-7} \text{ m}} = 1,1 \text{ m} \end{aligned}$$

**Remarque :** d'ordinaire, face aux petits angles, nous avons souvent le réflexe d'utiliser l'approximation  $\sin \theta \approx \theta$ . Cependant dans le cas présent, l'énoncé fournit l'angle en degrés tandis que l'approximation ci-dessus fonctionne uniquement si la valeur de  $\theta$  est en radians. Autrement dit nous nous épargnerions certes un sinus au moment de saisir le calcul, mais au prix d'une conversion degré/radian, et le tout pour obtenir en fin de compte un résultat légèrement inexact. Attention donc : rien ne sert de calculer, il faut savoir approximer à point.

Le plan d'observation doit donc être situé à  $D = 1,1 \text{ m}$  du plan des sources. Notons que ce dernier n'étant matérialisé par rien, il est préférable de donner la distance  $D'$  séparant l'écran de l'arête des miroirs, laquelle diffère de  $D$  suivant l'égalité :

$$D' = D - r \times \cos \alpha = 1,1 \text{ m}$$

Nous constatons donc qu'à 2 chiffres significatifs, les deux distances peuvent être assimilées (effet de l'observation à grande distance).

L'énoncé demandait enfin de déterminer l'étendue sur laquelle s'étend l'interférogramme à cette distance. Il vient naturellement :

$$\tan \alpha = \frac{L}{2 \times D'} \Leftrightarrow L = 2 \times D' \times \tan \alpha = 3,7 \text{ cm}$$

Nous comprenons ici pourquoi, l'expérience d'Young est souvent préférée (en terminale notamment) à celle de Fresnel. Les franges d'interférences données par cette dernière ne seront en effet visibles qu'à plus d'un mètre de distance (c'était déjà le cas avec les fentes d'Young), mais surtout l'interférogramme sera moins étendu. Nous pouvons ajouter que le positionnement de l'écran d'observation parallèlement à l'axe des sources sera plus difficile à réaliser que dans le cas du dispositif d'Young, où la bifente permettait de matérialiser le plan des sources.

Le dispositif de Fresnel présente tout de même un avantage majeur : les surfaces de collection de la lumière (les miroirs, donc) étant beaucoup plus étendues que dans le cas des fentes d'Young, la figure d'interférence sera également plus lumineuse.

Cependant les caractéristiques de l'interférogramme (interfrange, étendue) présentant une très grande dépendance vis-à-vis d'un angle de valeur  $\alpha$  très petite (et donc par nature difficile à mesurer) rend malgré tout cette expérience plus délicate à mettre en œuvre, et surtout à exploiter quantitativement.



### 8.1 La physique des quantas

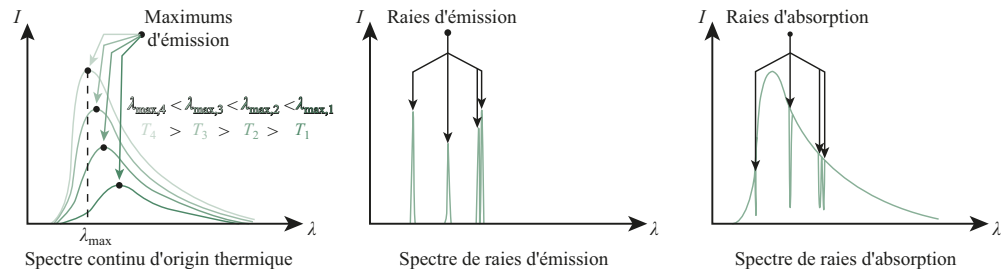
#### — Où en est-on concernant la lumière, à l'aube du $xx^e$ siècle ?

Nous avons vu dans le chapitre précédent que la lumière donnait lieu à l'ensemble des phénomènes propres aux ondes (réfraction, dispersion, interférences, diffraction pour ne citer que ceux-ci). L'extension aux ondes électromagnétiques, établie sur le plan théorique par James Clerk Maxwell en 1864 puis vérifiée expérimentalement en 1887 par Heinrich Hertz, parachève le tableau. À cette date, la nature ondulatoire de la lumière ne fait donc plus aucun doute.

Une question cependant résiste encore et toujours à l'envahisseur, et qui de simple singularité va progressivement devenir l'une des deux grandes pommes de discorde de la physique à la fin du  $xix^e$  siècle avant de la révolutionner de fond en comble : le spectre de rayonnement des corps denses chauffés.

#### — Quels sont les différents types de spectres lumineux connus ?

L'analyse des spectres lumineux des différents types de sources connus à l'époque met en lumière (si l'on peut dire) 3 types de spectres, associés à 3 contextes :



- **les corps denses** (solide, liquide ou gaz sous haute pression) chauffés, donnent un spectre dit **spectre continu d'origine thermique**, dans lequel toutes les radiations sont présentes, avec des intensités variables. Le profil spectral associé présente une forme de cloche, avec un maximum (dit maximum d'émission) dont nous noterons  $\lambda_{\max}$  la longueur d'onde. Ce profil spectral est sensiblement **le même quelle que soit la nature du corps considéré** : un morceau de braise incandescente, un fer chauffé, de la lave en fusion, etc. rayonnent tous, bon an, mal an, le même spectre, dont les caractéristiques quantitatives semblent dépendre uniquement de la température du corps en question. La longueur d'onde  $\lambda_{\max}$  du maximum d'émission, en particulier, est liée à la température absolue  $T$  selon la **loi du déplacement de Wien** :  $\lambda_{\max} = \frac{A}{T}$  avec  $A = 2,898.10^{-3} \text{ K.m}$ , constante de Wien ;

- **les gaz sous basse pression**, soumis à des décharges électriques, donnent des **spectres de raies d'émission**, dans lesquels seules quelques radiations sont présentes, et donnent ainsi une poignée de raies colorées. Les longueurs d'onde de ces raies sont cette fois **caractéristiques de la nature du gaz utilisé** ;
- **les gaz sous basse pression traversés par un spectre continu d'origine thermique**, enfin, donnent un **spectre de raies d'absorption** : on retrouve le spectre continu d'origine thermique de départ, mais parsemé de raies sombres correspondant donc à de brutales chutes d'intensité à certaines longueurs d'onde. Les radiations ainsi altérées sont précisément les mêmes que celles émises par ce même gaz lorsqu'il était excité par des décharges électriques.

La diversité de ces spectres, conjuguée à d'autres expériences telles que l'effet photoélectrique (existence d'une fréquence à partir de laquelle une onde lumineuse permet d'extraire un électron d'un métal) pose question. En particulier, l'échec des différentes tentatives d'expliquer le spectre continu d'origine thermique émis par les corps denses, sur la base de considérations microscopiques, va réclamer des études fines sur la structure des atomes qui vont finalement accoucher de la mécanique quantique.

## — Comment étudier le comportement d'un atome isolé ?

Le premier problème qui se pose, pour étudier le comportement d'un atome, est d'isoler celui-ci. En effet, s'il est inséré dans une masse compacte de matière (corps dense : solide, liquide ou gaz sous haute pression), chaque atome va voir son comportement modifié par la présence de ses voisins. Les collisions, l'interaction des particules élémentaires de l'un avec celles de l'autre, sont autant de facteurs qui vont empêcher l'observateur d'étudier le comportement d'un atome seul, et par suite de pouvoir tirer des conclusions concernant les lois auxquelles il obéit. Il semble donc nécessaire de travailler dans un milieu où la matière est raréfiée, par exemple un gaz sous basse pression.

En réalité, les raies émises par un ensemble d'atomes sont d'autant plus larges que ces atomes sont plus proches les uns des autres. Dans le cas d'un corps dense, cet élargissement est tel que les raies se recouvrent et donnent un spectre continu. Dans le cas d'un gaz sous basse pression, en revanche, tout se passe comme si chaque atome était isolé des autres. Les raies observées résultent donc de la superposition des spectres d'une population d'atomes isolés. Si tous les atomes sont identiques entre eux, chacun émettra les mêmes radiations, et présentera donc le même spectre.

Le spectre de raies d'émission apparaît donc comme une version collective du spectre d'un atome isolé. On ne peut pas tirer de conclusion sur l'intensité lumineuse du spectre observé, mais il est en revanche possible de considérer que la spécificité des longueurs d'onde qu'il révèle résulte de la seule interaction entre le noyau de cet atome et ses électrons, à l'exclusion notamment de toute interaction avec un autre atome.

## — Quelles informations les spectres de raies permettent-ils de déduire concernant la structure de l'atome ?

Pour obtenir un tel spectre, on produit des décharges électriques dans le gaz sous basse pression. Nous savons que la lumière véhicule de l'énergie, donc qu'il y a émission de spectre dit émission d'énergie.

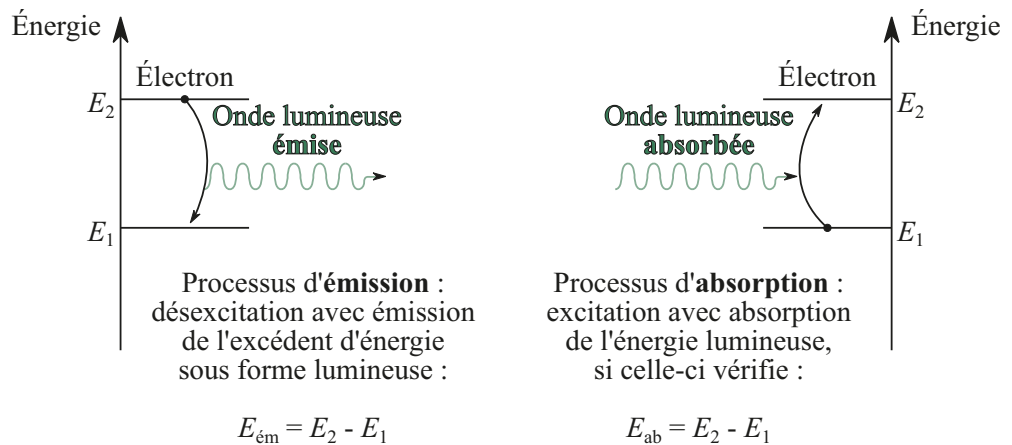
L'énergie étant une grandeur conservative, si les atomes en libèrent, c'est qu'ils en ont acquis. Les décharges électriques vont avoir pour effet d'exciter les atomes, qui vont absorber une partie de l'énergie véhiculée par ces décharges. Ce faisant, ils vont devenir instables et rapi-

dement revenir à leur état précédent, libérant ainsi l'énergie qu'ils avaient acquise. La lumière émise par ces atomes est la manifestation de cette libération d'énergie.

Cette idée est confirmée par les spectres de raies d'absorption, où les atomes du gaz sous basse pression absorbent précisément les radiations qu'ils sont capables d'émettre.

La structure d'un atome l'autoriserait donc à absorber l'énergie véhiculée par une radiation, si et seulement si cette énergie est également de celles que cette même structure l'autorise à émettre.

**Remarque :** rappelons que la lumière visible ne constitue qu'une toute petite partie du spectre du rayonnement électromagnétique (celle dont les longueurs d'onde dans le vide s'étendent de 384 à 768 nm). Les radiations ainsi émises peuvent donc être visibles, mais également appartenir aux domaines infrarouges, ultraviolet, etc.



Cependant, ce qui trouble les physicien(ne)s du début du XX<sup>e</sup> siècle est moins l'idée que les atomes puissent échanger de l'énergie, que le fait que ces échanges semblent ne se faire qu'à des fréquences spécifiques. Ceci mène à considérer l'idée que les atomes ne peuvent échanger de l'énergie qu'en quantités bien spécifiques, sous formes de paquets auxquels Max Planck donne le nom de **quanta**, terme qui vaudra son nom à la **mécanique quantique**.

Comme la seule différence que l'on note entre les différentes radiations est leur fréquence, et que la lumière véhicule de l'énergie, on peut imaginer qu'une différence de fréquence témoigne d'une différence d'énergie et que l'énergie  $E$  véhiculée par une radiation dépend donc de sa fréquence  $\nu$ .

C'est Max Planck qui, en cherchant sur la base de considération microscopiques (ce que l'on appelle de la *physique statistique*) un modèle capable de rendre compte du spectre continu d'origine thermique (modèle idéal, dit du *corps noir*) découvrir qu'une simple relation de proportionnalité entre ces deux grandeurs suffit à rendre compte de la loi en question :

$$E = h\nu$$

|       |        |
|-------|--------|
| $E$   | en J   |
| $h$   | en J.s |
| $\nu$ | en Hz  |

$h = 6,63 \cdot 10^{-34}$  J.s est appelée **constante de Planck**. Il est d'ailleurs intéressant de noter qu'à l'époque Planck reconnaît lui-même ne pas être en mesure d'expliquer la cause profonde de cette relation ; il la propose simplement parce qu'elle permet à son modèle de corps noir d'être conforme à l'expérience (ce qui reste tout de même le meilleur des arguments dans une science expérimentale).

S’ensuit l’idée que si les atomes ne peuvent échanger de l’énergie qu’en quantités particulières, c’est que les énergies dont ils sont capables d’être porteurs ne peuvent elles-mêmes adopter que des valeurs particulières : on dira de celles-ci qu’elles sont **quantifiées**.

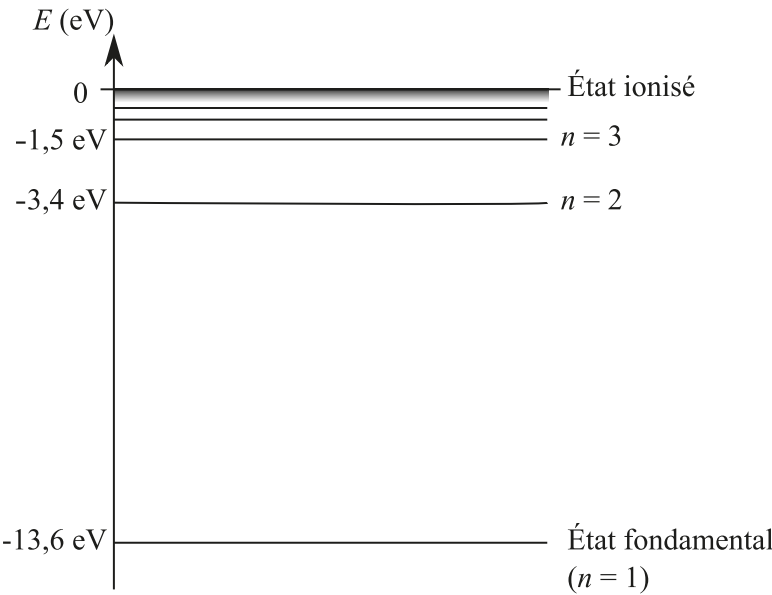
— **Quelles sont les règles régissant la quantification de l’énergie d’un atome ?**

Tous les atomes ne donnent pas les mêmes raies. Pour faire simple, prenons le plus simple des atomes : l’isotope 1 de l’hydrogène. En analysant en détail les énergies véhiculées par les radiations présentes dans le spectre de raies de cet atome, on finit par découvrir qu’elles peuvent toujours s’écrire sous la forme :

|                                                                     |                                 |                             |
|---------------------------------------------------------------------|---------------------------------|-----------------------------|
| $E_{\text{rad}} = R_H \left( \frac{1}{n^2} - \frac{1}{m^2} \right)$ | $E_{\text{rad}}, R_H$<br>$m, n$ | en eV<br>entiers sans unité |
|---------------------------------------------------------------------|---------------------------------|-----------------------------|

$R_H$  est appelée constante de Rydberg, et vaut 13,6 eV. La forme prise par ces échanges d’énergie incite à penser que les niveaux d’énergie auxquels peut accéder un atome sont de la forme :

|                          |                   |                             |
|--------------------------|-------------------|-----------------------------|
| $E_n = -\frac{R_H}{n^2}$ | $E_n, R_H$<br>$n$ | en eV<br>entiers sans unité |
|--------------------------|-------------------|-----------------------------|



**Remarque :** rappelons que l’électron-volt (eV) est une unité d’énergie particulièrement adaptée aux petites énergies rencontrées dans les atomes. Elle vaut :  $1 \text{ eV} = 1,602 \cdot 10^{-19} \text{ J}$ , ce qui correspond à l’énergie développée par un électron lorsque le potentiel électrique auquel il est soumis varie de 1 V (d’où son nom). Cette valeur correspond à l’ordre de grandeur typique du potentiel électrique généré par un proton (charge  $+e$ ) à une distance de l’ordre d’un rayon atomique (typiquement  $10^{-10} \text{ m}$ ). Les énergies engagées dans l’interaction noyau-électron sont de l’ordre de quelques eV à quelques dizaines d’eV.

Les spectres des autres atomes révèlent des règles de quantification de l’énergie analogues.

**Remarque :** à un niveau de complexité un peu plus élevé, on constate que les énergies des molécules sont également quantifiées. Elles peuvent être décrites à la manière de gros atomes, dont le noyau serait délocalisé sur plusieurs centres.

## — Existe-t-il d'autres grandeurs quantifiées dans un atome ?

La quantification de l'énergie n'était en réalité que la partie émergée de l'iceberg, et de nouvelles expériences vont rapidement prouver que d'autres grandeurs, jusqu'alors considérées comme variant continûment en mécanique classique, sont en réalité quantifiées. Le formalisme qu'il va falloir développer pour rendre compte de ces multiples quantifications va peu à peu dévoiler comment s'organisent les électrons au sein de l'atome. Vous connaissez en réalité déjà cette organisation, dont nous avons donné un aperçu dans le chapitre 1 :

| Grandeur quantifiée                              | Nombre quantique | Valeurs autorisées                  | Structure associée                                                                                   | Nombre possible                                          |
|--------------------------------------------------|------------------|-------------------------------------|------------------------------------------------------------------------------------------------------|----------------------------------------------------------|
| Énergie $E$                                      | Primaire : $n$   | 1, 2, 3...                          | Couche électronique ( $E_1, E_2, E_3, \dots$ )                                                       | Potentiellement infini                                   |
| Moment cinétique $\vec{L}$                       | Secondaire : $l$ | 0, 1, ..., $n-1$                    | Orbitale électronique ( $s, p, d, f, \dots$ )                                                        | $n$ orbitales pour la couche $n^\circ$                   |
| Projection de $\vec{L}$ dans un champ magnétique | Magnétique : $m$ | $-l, -l+1, \dots, 0, \dots, l-1, l$ | Case quantique (1 pour une orbitale $s$ , 3 pour une orbitale $p$ , 5 pour une orbitale $d, \dots$ ) | $2l+1$ cases quantiques pour une orbitale de $n^\circ l$ |
| Projection du spin                               | De spin : $s$    | $-\frac{1}{2}, +\frac{1}{2}$        | État quantique                                                                                       | 2 états quantiques par case quantique                    |

Ainsi, la structure électronique que vous avez appris à décrire au chapitre 1 découle-t-elle entièrement de cette série de quantifications.

## — Les atomes sont-ils les seuls systèmes au sein desquels les grandeurs physiques sont quantifiées ?

Nous avons jusqu'ici considéré le noyau de l'atome comme une grosse particule chargée. Or nous savons qu'il est lui-même constitué de particules : les protons et les neutrons. Les protons portent tous une charge électrique  $+e$ , donc ils se repoussent. Cette répulsion est en outre considérable : étant très proche les uns des autres dans le noyau, la force coulombienne (qui rappelons-le varie en  $1/r^2$ ) devient énorme (du moins à leur échelle).

Le fait que le noyau tienne malgré cette répulsion interne met en évidence l'existence d'une autre interaction, dite **interaction nucléaire forte**. C'est elle qui assure la cohésion du noyau en maintenant les nucléons ensemble.

Or ce système étant encore plus petit que l'atome, il n'échappe pas à la mécanique quantique et les énergies qu'il peut adopter sont elles aussi quantifiées. En d'autres termes, le noyau

ne peut exister que si son énergie adopte certaines valeurs. Lorsqu'il passe d'un niveau d'énergie à un autre d'énergie inférieure, il libère son excédent d'énergie sous forme de rayonnement électromagnétique, comme le font les électrons d'un atome. Mais les énergies engagées pour assurer la cohésion du noyau sont considérablement plus élevées que celles engagées dans la liaison noyau-électron. Typiquement, elles sont de l'ordre du mégaelectron-volt (MeV), c'est-à-dire un million de fois plus élevée que celles engagées dans la liaison noyau-électron.

Il en résulte que le rayonnement dégagé par un noyau qui se désexcite n'est pas simplement de la lumière visible, mais quelque chose de beaucoup plus énergétique : un rayonnement  $\gamma$ . Nous constatons ainsi que toute la physique nucléaire porte en fait sur la désexcitation de noyaux instables entre deux niveaux d'énergie, exactement comme la physique atomique. La différence étant que les particules liées sont cette fois des nucléons dans un noyau et non un noyau et des électrons dans un atome, et que la cohésion n'est plus assurée par l'interaction électromagnétique, mais par l'interaction nucléaire forte.

## 8.2 La structure de l'atome et l'interaction noyau-électron

### Comment le noyau et les électrons sont-ils liés dans l'atome ?

Nous voici donc dotés d'un fait auquel ne nous avait pas habitués la mécanique de Newton. Selon cette dernière, l'énergie d'un système peut varier continûment. L'énergie d'un système soumis à des frottements, par exemple, se dissipe progressivement. On n'envisage pas *a priori* de système qui ne pourrait perdre de l'énergie par frottements que si la dissipation ainsi provoquée lui permettait d'accéder à un état d'énergie particulière (dans les faits c'est le cas, puisque le système lui-même est constitué d'atomes, mais, les écarts en question étant de quelques eV, cette quantification est imperceptible à l'échelle d'un système macroscopique).

Pour aller plus loin, il devient nécessaire de considérer la constitution d'un atome. Pour commencer par un cas simple, nous allons de nouveau nous intéresser à l'atome d'hydrogène 1 : un noyau constitué d'un unique proton, auquel se trouve associé un électron. Voyons les caractéristiques de ces particules :

- Pour le proton : sa masse vaut  $1,67 \cdot 10^{-27}$  kg, et sa charge  $q_p = +e = 1,602 \cdot 10^{-19}$  C.
- Pour l'électron : sa masse vaut  $9,11 \cdot 10^{-31}$  kg, et sa charge  $q_e = -e = -1,602 \cdot 10^{-19}$  C.

Chacune de ces deux particules possède une masse, elles s'attirent donc suivant la loi d'attraction gravitationnelle, pour laquelle on trouve, avec  $r = 5,0 \cdot 10^{-11}$  m, une valeur  $F_g = 4,1 \cdot 10^{-47}$  N.

Chacune de ces deux particules possède une charge électrique de signe opposé à celui de l'autre, elles s'attirent donc suivant la loi d'attraction électrique, pour laquelle on trouve, avec  $r = 5,0 \cdot 10^{-11}$  m, une valeur  $F_e = 9,2 \cdot 10^{-8}$  N.

On constate donc qu'au sein d'un atome l'interaction gravitationnelle est considérablement moins élevée que l'interaction coulombienne (environ 39 ordres de grandeur d'écart, soit un facteur  $10^{39}$ ). Ainsi le système {proton + électron} qui constitue l'atome d'hydrogène est-il régi essentiellement par l'interaction coulombienne.

## — Peut-on considérer que les électrons tournent autour du noyau comme les satellites autour d'un astre ?

On savait depuis Rutherford que l'atome était en fait constitué d'un noyau de très petite taille, concentrant l'essentiel de la masse et autour duquel se trouvaient les électrons. Niels Bohr développa cette idée d'un modèle planétaire de l'atome sur le plan dynamique.

Du point de vue physique, l'interaction gravitationnelle et l'interaction coulombienne sont fondamentalement différentes. Cependant, si nous cherchons à connaître le mouvement de ce système, l'application du PFD nous mène à des équations en tout point analogues à celles

rencontrées lors de l'étude du mouvement d'un satellite autour d'un astre :  $\vec{a} = -\frac{K}{r^2} \vec{e}_r$ , et ne différant que par la valeur (et possiblement le signe) de  $K$ .

En effet, les forces gravitationnelle et coulombienne variant toutes deux en  $\frac{1}{r^2}$ , le traite-

ment mathématique est identique quelle que soit la nature de l'interaction. Les solutions de ces équations devraient donc être les mêmes, et par suite les mouvements décrits devraient être semblables. C'est-à-dire que si les lois dynamiques à l'œuvre dans un atome étaient les mêmes qu'à l'échelle humaine, alors le mouvement d'un électron autour du noyau serait exactement le même que celui d'un satellite autour d'un astre.

Or deux choses viennent infirmer cette idée :

- On constate expérimentalement qu'un électron doté d'une accélération non nulle rayonne de l'énergie. Donc il perd de l'énergie. Or, si un satellite en orbite autour d'un astre perd de l'énergie, il tombe vers celui-ci. Donc, si l'électron tournait autour du noyau comme un satellite autour d'un astre, il devrait s'écraser sur le noyau.
- En mécanique classique, un satellite peut posséder n'importe quelle énergie mécanique. En particulier, cette idée ne rend absolument pas compte de la quantification des énergies observées dans les spectres de raies.

## — D'où les singularités de comportement de l'atome sont-elles issues ?

Nous avons développé dans les pages qui précèdent le fait qu'aux échelles atomique et subatomique, les grandeurs physiques étaient quantifiées. Cependant la question de savoir à quoi cette quantification est due reste pour l'instant sans réponse. Ceci parce que nous avons l'habitude de décrire la matière en termes de particules, auxquelles sont associées un certain nombre d'idées : rigidité, position repérable sans limite de précision, vitesse calculable à l'envi, etc. Autant de concepts qui ne font manifestement plus recette pour décrire l'atome.

La quantification n'est cependant pas une idée toute neuve : elle est connue depuis longtemps, mais dans un autre domaine, à savoir celui des ondes (et plus particulièrement des ondes stationnaires, malheureusement pas au programme de terminale mais que vous aborderez en CPGE). Mais ondes et particules sont connues pour être des objets physiques très différents, et l'on se figure mal comment appliquer un formalisme ondulatoire à une particule...

## 8.3 La dualité onde-corpuscule

### — La lumière : onde ou corpuscule ?

L'un des très grands débats qui animent la Physique à la fin du XIX<sup>e</sup> siècle porte sur la nature de la lumière.

Les études menées par Maxwell permettent de décrire le rayonnement électromagnétique en général, et la lumière en particulier, comme une onde. Mais la physique des quanta, avec ses paquets d'énergie, évoque plutôt l'idée de particules. Einstein, partisan de cette dernière théorie, donne à ces hypothétiques particules de lumière le nom de **photons**.

Pour comprendre le problème, on peut mener l'expérience suivante : on dirige un faisceau de lumière issue d'une source de puissance réglable sur un objet diffractant. On fait ensuite décroître cette puissance jusqu'à ce qu'elle ne débite plus qu'un seul quanta d'énergie à la fois. On dispose alors un capteur photosensible (pellicule ou matrice de capteurs numériques) à la suite de l'objet diffractant :

- Si la lumière est une onde, alors la figure de diffraction va apparaître entièrement sur la plaque photographique dès le début, très faiblement puis de plus en plus visible à mesure que les ondes lumineuses viendront impressionner le capteur.
- Si la lumière est constituée de photons, alors chaque photon va marquer un point précis sur le capteur, et l'on verra de la lumière apparaître point par point. Mais qu'obtiendra-t-on à la fin ?

Au début de l'expérience, la victoire semble revenir aux partisans des photons : la plaque est bien impressionnée point après point. Mais au bout de quelques temps, on constate que ces points dessinent une figure bien particulière, qui n'est autre que la figure de diffraction attendue. Les photons se répartissent selon une disposition qui se trouve être précisément celle que serait supposée donner une onde.

Pour mettre les deux points de vue en accord, on associe alors aux particules que sont les photons une **onde de probabilité**, la longueur d'onde d'une radiation correspondant alors à la longueur d'onde de cette onde de probabilité. La lumière n'est plus une onde à proprement parler, mais un flux de particules (les photons) dont le comportement, en particulier la façon de se répartir, est décrit par une onde. Lorsque l'un d'eux franchit un obstacle, il part *a priori* dans une direction aléatoire. Mais sur un grand nombre de photons, certaines de ces directions se révèlent particulièrement probables (donnant les pics lumineux de la figure de diffraction), et d'autres particulièrement improbables (zones d'extinction).

### — Les autres particules peuvent-elles également être décrites par des ondes de probabilité ?

Face à cette découverte, se pose la question de savoir si cette onde de probabilité est l'apanage des particules de lumière, ou bien si le comportement de toute particule (électron, proton, neutron...) peut également être décrit par des ondes et serait donc de nature à subir les effets de la diffraction, par exemple.

C'est Louis de Broglie qui propose de généraliser ce concept. Il associe ainsi à tout système matériel doté d'une quantité de mouvement  $p$  une **longueur « d'onde de probabilité »  $\lambda$**  :

$$\lambda = \frac{h}{p}$$

|           |                         |
|-----------|-------------------------|
| $\lambda$ | en m                    |
| $h$       | en J.s                  |
| $p$       | en kg.m.s <sup>-1</sup> |

$h$  est la constante de Planck, déjà rencontrée dans la relation  $E = h\nu$ . Or  $h = 6,67 \cdot 10^{-34}$  J. Les valeurs de  $\lambda$  sont donc en général très faibles, et la diffraction des ondes de probabilité rarement manifeste (rappelons qu'une onde ne peut être diffractée que si elle franchit un obstacle dont la taille est au plus de l'ordre de sa longueur d'onde).

**Remarque :** réciproquement, il est ainsi possible d'associer à toute onde de longueur d'onde  $\lambda$  une quantité de mouvement  $p = h/\lambda$ . Ainsi des photons, en dépit de leur absence de masse, possèdent-ils une inertie. Ce point est vérifié expérimentalement.

On constate cependant que pour un électron, de masse  $m = 9,11 \cdot 10^{-31}$  kg, animé par exemple d'une vitesse  $v = 5,0 \cdot 10^6$  m.s<sup>-1</sup> (suffisante pour traverser un solide cristallin), on trouve  $\lambda = 1,5 \cdot 10^{-10}$  m. Dans le spectre des ondes électromagnétiques, cette longueur d'onde appartient au domaine des rayons X. Elle est de l'ordre de grandeur de la distance séparant les noyaux des atomes dans un réseau cristallin, lequel devrait donc constituer un obstacle de largeur suffisamment faible pour diffracter le faisceau d'électrons, tout comme les rayons X.

Ainsi, d'après cette théorie, un faisceau d'électrons porté à une vitesse de cinq mille kilomètres par seconde devrait se répartir suivant une figure sensiblement identique à celle obtenue par diffraction d'un faisceau de rayons X par ce même réseau. Et c'est précisément ce qui est observé.

## — Quelles sont les conséquences pour les électrons dans un atome ?

La description du comportement des particules en termes d'onde de probabilité va alors permettre de comprendre la structure des atomes. Pour décrire la présence d'une particule en termes de probabilité, Erwin Schrödinger est amené à développer, avec plusieurs autres physiciens, un formalisme analogue à celui de la mécanique newtonienne, et dans lequel l'électron n'est plus décrit en tant que particule, mais en tant qu'**onde de probabilité**. C'est-à-dire qu'au lieu de chercher à rendre compte du mouvement de l'électron en décrivant précisément où il se trouve, on décrit la **densité volumique de probabilité de présence** de l'électron autour du noyau. Il en résulte une délocalisation de l'électron sur l'ensemble de l'atome, ce qui amène à décrire la succession des positions qu'il occupe par un **nuage électronique** plus que par une quelconque orbite (même si les états accessibles aux électrons du cortège portent tout de même le nom d'*orbitales atomiques*).

Dans le cadre de ce formalisme, les calculs amènent spontanément la quantification des énergies. Le prix à payer est d'accepter qu'à l'échelle atomique, il est impossible de connaître précisément la position d'une particule.

## — Quelle est la forme du nuage électronique d'un atome ?

Aux différentes couches d'un atome sur lesquelles se répartissent les électrons, sont associées diverses fonctions décrivant la densité de probabilité de présence. Celles-ci sont définies en tout point de l'espace.

Ainsi par exemple, les orbitales  $s$  possèdent une symétrie sphérique, et les surfaces d'isodensité de probabilité de présence sont-elles des sphères, sur lesquelles la densité présente un ou plusieurs maximums (selon la valeur du nombre quantique principal) à des rayons spécifiques.

Les orbitales  $p$ , pour leur part, présentent 3 paires de lobes portées par 3 axes formant un trièdre rectangle.

À partir des orbitales  $d$ , la géométrie se complique significativement.

## Halte aux idées reçues

Expliquer, pour chacune des affirmations suivantes, pourquoi elle est fautive :

1. Un corps dense répondant parfaitement au modèle du corps noir de Planck apparaît forcément de couleur noire.
2. Puisque tous les atomes possèdent les mêmes orbitales, ils doivent également posséder les mêmes niveaux d'énergie et toujours répondre à la relation  $E_n = -\frac{R_H}{n^2}$ .
3. Si une particule électriquement chargée se rapproche d'une autre, alors les valeurs des forces électromagnétiques respectivement exercées par chacune sur l'autre augmentent instantanément.

4. La puissance lumineuse délivrée par un laser est d'autant plus importante que la longueur d'onde de la lumière qu'il émet est plus courte.

5. Puisqu'un photon possède une quantité de mouvement  $p = \frac{h}{\lambda}$ , et que par ailleurs la quantité de mouvement s'exprime  $p = m \times v$ , on peut en déduire que le photon (de vitesse  $v = c$  dans le vide) possède une masse  $m = \frac{h}{\lambda \times c} = \frac{h}{c^2} \times f$ , d'autant plus élevée que la fréquence de l'onde électromagnétique associée l'est aussi.

## Du Tac au Tac

(Exercices sans calculatrice)

1. Montrer que le nombre total d'états quantiques disponible sur la couche électronique de nombre quantique principal  $n$  est égal à  $2n^2$ .

Donnée : on rappelle que la somme des  $n$  premiers entiers est égale à :  $1 + 2 + 3 + \dots + n = \frac{n \times (n+1)}{2}$ .

2. Les panneaux solaires sont constitués de cellules photovoltaïques ayant la propriété d'absorber l'énergie véhiculée par les ondes électromagnétiques pour la convertir en énergie électrique. On considère un tel panneau, constitué de  $N_c = 50$  cellules carrées, de  $a = 400$  mm de côté. Éclairé d'une puissance surfacique  $\frac{P_{\text{lum}}}{S} = 1,00 \text{ kW.m}^{-2}$ , il débite dans un conducteur ohmique de résistance  $R = 250 \Omega$  un courant d'intensité constante  $I = 2,00$  A. Déterminer le rendement des cellules photovoltaïques utilisées.

3. On considère une diode laser délivrant dans l'air un faisceau de lumière de longueur d'onde  $\lambda = 663$  nm, et de puissance  $P = 120$  mW. Déterminer combien de photons il produit par seconde.

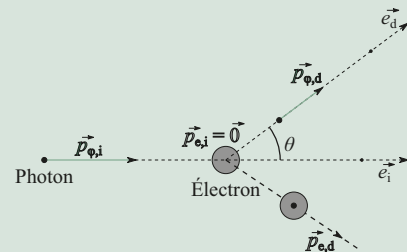
4. La voile solaire est une idée de système de propulsion spatiale reposant sur la quantité de mouvement transmise par les photons à un objet lorsqu'ils se réfléchissent à sa surface. En substance, elle permettrait de propulser une sonde à la manière d'un navire, en utilisant la lumière en guise de vent. Montrer qu'une surface parfaitement réfléchissante recevant une puissance lumineuse  $P_{\text{lum}}$  sous incidence normale va subir une force de valeur :

$$F_{\text{lum}} = 2 \times \frac{P_{\text{lum}}}{c}$$

En déduire une valeur approchée des dimensions que devrait posséder une voile parfaitement réfléchissante pour

développer une force de 1,0 N lorsqu'elle est éclairée d'une puissance surfacique  $\frac{P_{\text{lum}}}{S} = 1,00 \text{ kW.m}^{-2}$ .

5. Lorsqu'un photon rencontre un électron (masse  $m_e$ ) isolé et initialement immobile, il peut être diffusé par celui-ci : à la suite de leur rencontre, le photon repart dans une autre direction. On constate alors que la longueur d'onde  $\lambda_d$  du photon après diffusion diffère de sa longueur d'onde initiale  $\lambda_i$  (phénomène connu sous le nom d'effet Compton).



Montrer, en vous appuyant sur la conservation de l'énergie totale ainsi que sur celle de la quantité de mouvement totale du système {électron+photon}, que les longueurs d'onde associées respectivement aux photons incident et diffusé sont liées à l'angle de diffusion  $\theta$  selon l'égalité :

$$\lambda_d - \lambda_i = \frac{h}{m_e c} \times (1 - \cos \theta)$$

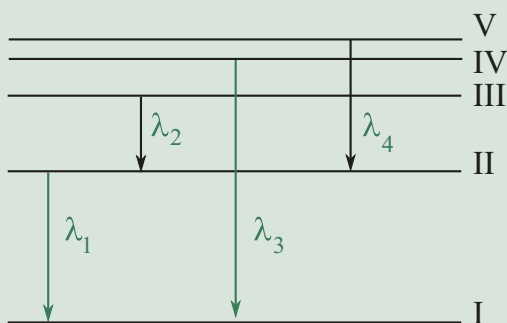
Données :

- quantité de mouvement d'un photon de longueur d'onde  $\lambda$  :  $p = \frac{h}{\lambda}$
- formule générale de l'énergie d'un système de quantité de mouvement  $p$  et de masse  $m$  :

$$E = \sqrt{m^2 c^4 + p^2 c^2}$$

## Vers la prépa

La figure suivante présente les niveaux d'énergie de l'atome neutre de lithium, de numéro atomique égal à 3, de structure électronique  $1s^2 2s^1$ .



On considère les quatre transitions représentées sur le diagramme. Les longueurs d'onde correspondantes sont :  $\lambda_1 = 671 \text{ nm}$ ,  $\lambda_2 = 812 \text{ nm}$ ,  $\lambda_3 = 323 \text{ nm}$  et  $\lambda_4 = 610 \text{ nm}$ .

**1.** La(les)quelle(s) de ces radiations apparten(en)nt au domaine du visible ?

**2.** Montrer que l'énergie  $E_\varphi$  d'un photon est liée à sa longueur d'onde  $\lambda$  par la relation :  $E_\varphi = \frac{1,24 \cdot 10^3}{\lambda}$ , où  $\lambda$

est exprimée en nm et  $E_\varphi$  en eV. Déterminer ensuite, en eV, l'énergie des photons émis lors de chacune des quatre transitions indiquées.

**3.** L'énergie du niveau I vaut  $E_I = -5,39 \text{ eV}$ . C'est l'énergie de l'électron externe de l'atome dans son état fondamental. Déterminer alors les valeurs des énergies des niveaux II, III, IV et V.

**4.** Pour quelle valeur de la longueur d'onde des radiations incidentes les atomes de lithium subiront-ils une ionisation à partir de l'état fondamental ?

On donne :

- la charge électrique élémentaire :  $e = 1,60 \cdot 10^{-19} \text{ C}$  ;
- la constante de Planck :  $h = 6,63 \cdot 10^{-34} \text{ J.s}$  ;
- la célérité de la lumière dans le vide :  $c = 3,00 \cdot 10^8 \text{ m.s}^{-1}$  ;
- la masse de l'électron :  $m = 9,11 \cdot 10^{-31} \text{ kg}$ .

## Halte aux idées reçues

1. Le modèle dit du corps noir établi par Planck est ainsi nommé parce que l'une de ses hypothèses fondatrices est que le corps en question peut uniquement émettre et absorber des ondes électromagnétiques, à l'exclusion notamment de toute forme de réflexion. Dans ces conditions, à température ambiante (où le maximum d'émission est faible et se situe autour de  $10\text{ }\mu\text{m}$ , cf. loi de Wien), l'œil humain ne percevra rien en provenance de ce corps : les radiations qu'il émet se situent pour l'essentiel hors de son domaine de perception ( $384\text{ à }768\text{ nm}$ ), et puisqu'il n'en réfléchit aucune de celles que son environnement pourrait lui envoyer, il paraîtra d'un noir absolu. Cette hypothèse a été formulée par Planck afin de pouvoir faire aboutir ses calculs (déjà bien musclés), et le spectre décrit par son modèle de corps noir rend donc d'autant mieux compte d'un corps réel, que celui-ci est moins réfléchissant. Cependant, si nous nous intéressons maintenant à un corps porté par exemple à une température de  $6000\text{ K}$ , comme la température de surface de certaines étoiles, le maximum d'émission sera cette fois beaucoup plus élevé et se situera autour de  $500\text{ nm}$ . Dans ces conditions, même si les propriétés de ce corps sont telles qu'il ne réfléchisse aucune radiation, il en émettra malgré tout beaucoup, et qui plus est dans le domaine visible. Il apparaîtra donc brillant (dans le cas présent le maximum d'émission serait dans le vert, mais les autres composantes visibles seraient tellement présentes que nous percevrions du blanc) ; l'expression *corps noir* n'est donc qu'une expression et doit être considérée comme telle.

2. Il est exact de dire que la structure du cortège électronique d'un atome comporte toujours les mêmes orbitales. En effet, les différentes grandeurs quantifiées dans un atome sont les mêmes, et la caractérisation d'un état quantique par un quadruplet  $(n, l, m, s)$  est donc valable quel que soit l'atome. Ceci entraîne également l'existence des orbitales que vous connaissez désormais bien :  $1s, 2s, 2p, 3s, 3p, \dots$ . Cependant le fait que ces grandeurs soient toujours quantifiées ne signifie pas qu'elles adoptent les mêmes jeux de valeurs. En particulier, la formule de Rydberg rappelée dans l'énoncé et décrivant les valeurs des énergies accessibles vous a été fournie dans un contexte bien spécifique, à savoir celui d'un atome neutre d'hydrogène 1. Pour un atome plus complexe, comportant par exemple ne serait-ce qu'un électron de plus, les interactions vont devenir plus complexes, avec cette fois deux interactions noyau/électrons, sans compter l'interaction entre les deux électrons. Les niveaux d'énergie resteront discrets, mais les valeurs qui leur seront associées seront différentes, et répon-

dront à des relations nettement plus élaborées (quand elles sont calculables).

3. Nous savons que la force électrique exercée par une particule électriquement chargée  $A$  sur une autre particule  $B$  peut être décrite comme le produit de la charge électrique  $q_B$  de la particule subissant cette force, par le champ électrique généré par la première particule, ce champ ayant lui-même une expression connue 
$$\vec{E}_A(B) = \frac{1}{4\pi\epsilon_0} \frac{q_A}{(AB)^2} \vec{e}_{AB}.$$

Or nous savons également que la lumière est une onde électromagnétique, c'est-à-dire le phénomène de propagation d'une perturbation du champ électromagnétique, cette perturbation ayant la faculté de se propager dans le vide à la célérité  $c = 299\,792\,458\text{ m.s}^{-1}$ . Ainsi, si la source d'un champ électrique se déplace, l'effet de ce déplacement n'est sensible sur les objets assujettis à son influence qu'un peu plus tard, le temps que la modification se propage.

On cite parfois l'exemple du Soleil, situé à un peu plus de 8 minutes-lumière de la Terre, et dont les informations lumineuses que nous recevons sont donc déjà vieilles de ces 8 minutes. Ces considérations sont le pain quotidien des personnes chargées des télécommunications avec les sondes (et peut-être prochainement les vaisseaux habités) évoluant dans l'espace interplanétaire, les ondes radio utilisées étant de même nature que la lumière visible ou tout autre onde électromagnétique.

4. La puissance échangée par un système est le rapport de l'énergie échangée par ce système avec le reste de l'univers, à la durée  $\Delta t$  de cet échange. Le fait qu'un laser émette à une longueur d'onde plus courte qu'un autre signifie certes que l'énergie individuelle de chacun des photons qu'il produit est plus élevée, puisque cette énergie s'exprime  $E = \frac{hc}{\lambda}$ .

Mais une fois ceci établi, la quantité totale d'énergie que débite ce laser par unité de temps, ne dépend pas uniquement de l'énergie véhiculée par chaque photon, mais également du nombre de ces photons que produit le laser par unité de temps.

Ainsi si l'on divise la longueur d'onde d'un laser par deux, les photons seront-ils deux fois plus énergétiques. Mais si corrélativement le laser en émet deux fois moins par unité de temps, l'énergie délivrée par unité de temps reste la même, voire est inférieure si l'on réduit encore le débit de photons.

5. L'erreur centrale de cet énoncé consiste à prendre deux relations issues de deux contextes différents et à les croiser sans états d'âme. S'agissant d'un photon, la relation de de Broglie est pertinente. En revanche  $p = m \times v$  est une relation

issue de la mécanique classique, à laquelle se sont substituées la mécanique quantique et la mécanique relativiste, justement pour pallier les insuffisances du modèle classique, respectivement pour les systèmes très petits et ceux très rapides. Le photon entrant dans ces deux catégories, le problème doit être traité avec beaucoup plus de précautions et ne pas se résumer (enfin, encore moins que d'habitude) à une simple salade de formules. Le calcul ci-dessus donnerait, pour un photon visible, une masse de l'ordre de  $10^{-36}$  kg, quand l'expérience montre que cette masse, si elle existait, sera nécessairement inférieure à  $10^{-54}$  kg (soit tout de même un petit facteur un milliard de milliards de fois moins).

## Du Tac au Tac

1. Dans cet énoncé, le nombre quantique principal est fixé et égal à  $n$ . Il nous suffit donc de calculer successivement :

- combien d'orbitales sont associées à la couche électronique de niveau  $n$  : de  $l = 0$  à  $l = n - 1$  par saut de 1, nous dénombrons  $n$  orbitales ;
- pour chacune de ces orbitales, de combien de cases quantiques elle dispose ; pour une orbitale de nombre quantique secondaire  $l$  : de  $m = -l + 1$  à  $m = l - 1$  par saut de 1, nous dénombrons  $(2l + 1)$  cases quantiques. Le nombre  $N_{CQ}$  de cases quantiques présentes sur la couche  $n$  (en cumulant celles offertes par les orbitales  $s, p, d$ ) serait donc de :

$$N_{CQ} = \sum_{l=0}^{n-1} (2l + 1) = 2 \times \left( \sum_{l=0}^{n-1} l \right) + n$$

le terme  $n$  provenant du « 1 » qui intervient dans chacun des termes  $(2l + 1)$ .

En utilisant la relation fournie par l'énoncé, la somme des  $n - 1$  premiers entiers se calcule comme :

$$\sum_{l=0}^{n-1} l = \frac{(n-1) \times (n-1+1)}{2} = \frac{(n-1) \times n}{2}$$

Et en injectant cette égalité dans la relation précédente, nous trouvons un nombre de cases quantiques égal à :

$$N_{CQ} = 2 \times \left[ \frac{(n-1) \times n}{2} \right] + n = (n-1) \times n + n$$

$$\Leftrightarrow N_{CQ} = (n-1+1) \times n = n^2$$

- enfin, à raison de deux états de spin possibles par case quantique, nous trouvons un nombre total d'états quantiques disponibles sur la couche  $n$  égal à :

$$N_{EQ} = 2 \times N_{CQ} = 2n^2$$

Nous retrouvons ainsi les valeurs déjà rencontrées en chimie :

- pour  $n = 1$ ,  $2 \times 1^2 = 2$  électrons (ceux de l'orbitale  $1s$ ) ;
- pour  $n = 2$ ,  $2 \times 2^2 = 8$  électrons (2 sur l'orbitale  $2s$ , plus 6 sur l'orbitale  $2p$ ) ;
- pour  $n = 3$ ,  $2 \times 3^2 = 18$  électrons (2 sur l'orbitale  $3s$ , plus 6 sur l'orbitale  $3p$ , plus encore 10 sur l'orbitale  $3d$ ) ;

- pour  $n = 4$ ,  $2 \times 4^2 = 32$  électrons (2 sur l'orbitale  $4s$ , plus 6 sur l'orbitale  $4p$ , plus encore 10 sur l'orbitale  $4d$  et enfin 14 sur l'orbitale  $4f$ ) ;

ainsi que la largeur totale du tableau de Mendeleiev : 18 colonnes pour les blocs  $s, d$  et  $p$ , plus un bloc de 14 colonnes pour le bloc  $f$  intervenant seulement à partir de  $n = 4$ .

2. Nous savons que le rendement d'un système de conversion énergétique est défini comme le rapport entre :

- la puissance utile, soit en l'occurrence la puissance électrique délivrée par le panneau et dissipée par le conducteur ohmique :

$$P_{el} = R \times I^2 = 250 \, \Omega \times (2,00 \, A)^2 = 1,00 \, kW$$

- la puissance fournie à l'entrée du convertisseur, soit ici la puissance lumineuse collectée par le panneau :

$$P_{lum,col} = \frac{P_{lum}}{S} \times S_{col} = \frac{P_{lum}}{S} \times N_c \times a^2$$

$$= 1,00 \, kW.m^{-2} \times 50 \times (4,00.10^{-1} m)^2 = 8,00 \, kW$$

Nous trouvons donc finalement un rendement :

$$\eta = \frac{P_{el}}{P_{lum,col}} = 0,125 = 12,5\%$$

3. Nous savons que la puissance émise par un système est le rapport de l'énergie qu'il produit pendant une durée  $\Delta t$ , à cette durée. Dans le cas présent, en notant  $\frac{N}{\Delta t}$  le nombre de photons produits par seconde par le laser, et  $E_\varphi$  l'énergie d'un de ces photons, nous pouvons écrire que :

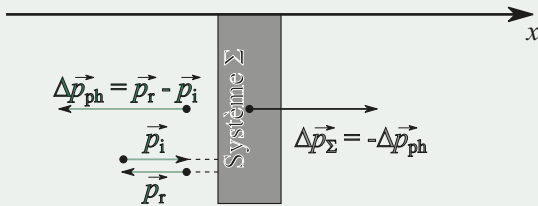
$$P = \frac{NE_\varphi}{\Delta t}$$

En exprimant l'énergie d'un de ces photons sous la forme

$E_\varphi = h\nu = \frac{hc}{\lambda} = 3,0.10^{-19} \, J$  (soit donc environ 2 eV : nous sommes dans les bons ordres de grandeur pour une énergie échangée au niveau du cortège électronique), nous trouvons ainsi que :

$$\frac{N}{\Delta t} = \frac{P}{E_\varphi} = 4,0.10^{17} \text{ photons par seconde}$$

4. Une erreur à ne pas commettre sur un exercice de ce type consisterait à chercher à tout prix une relation entre force et puissance fondée sur la mécanique newtonienne. S'agissant de photons, encore une fois, nous devons rester très prudent(e)s quant aux relations que nous pouvons utiliser. Concernant un système plus classique tel qu'une sonde ou quelque autre vaisseau spatial, en revanche, l'utilisation des lois de Newton reste tout à fait autorisée. Reste à trouver un dénominateur commun aux deux, que l'énoncé pointe clairement du doigt, à savoir la quantité de mouvement. Considérons donc ce qu'il advient d'un système matériel lorsqu'un photon se réfléchit à sa surface :



L'incidence étant normale, tout le problème se déroule uniquement selon la direction  $Ox$  et nous travaillerons donc directement sur les projections des vecteurs sur cet axe.

Nous constatons que lorsqu'il est réfléchi, le photon voit sa quantité de mouvement varier de :

$$\Delta p_{\text{ph}} = p_r - p_i = -2 \times p_{\text{ph}} = -2 \times \frac{h}{\lambda}$$

En supposant le système  $\Sigma$  isolé, sa quantité de mouvement globale ne peut pas varier et nous avons alors :

$$\Delta p_{\text{ph}} + \Delta p_{\Sigma} = 0 \Leftrightarrow \Delta p_{\Sigma} = -\Delta p_{\text{ph}} = 2 \times \frac{h}{\lambda}$$

Et pour une sonde recevant un nombre de  $N_{\text{ph}}$  de photons, cette relation devient :

$$\Delta p_{\Sigma} = 2N_{\text{ph}} \times \frac{h}{\lambda}$$

Or nous savons que si la quantité de mouvement de la sonde varie sur une durée  $\Delta t$ , alors elle subit une force :

$$F_{\text{lum}} = \frac{\Delta p_{\Sigma}}{\Delta t} = 2 \times \frac{N_{\text{ph}}}{\Delta t} \times \frac{h}{\lambda}$$

En notant que  $\frac{N_{\text{ph}} \times h}{\lambda} = \frac{E_{\text{lum}}}{c}$  nous trouvons :

$$F_{\text{lum}} = \frac{2}{c} \times \underbrace{\frac{E_{\text{lum}}}{\Delta t}}_{\substack{\text{Puissance} \\ \text{véhiculée} \\ \text{par } N_{\text{ph}} \\ \text{photons.s}^{-1}}} \Leftrightarrow F_{\text{lum}} = \frac{2}{c} \times P_{\text{lum}}$$

Il nous suffit pour finir d'exprimer la puissance lumineuse reçue comme :

$$P_{\text{lum,col}} = \frac{P_{\text{lum}}}{S} \times S_{\text{col}} \Leftrightarrow S_{\text{col}} = \frac{P_{\text{lum,col}}}{1,0 \text{ kW.m}^{-2}} = \frac{c}{2} \times \frac{F_{\text{lum}}}{1,0 \text{ kW.m}^{-2}}$$

L'application numérique donne alors :

$$S_{\text{col}} = 1,5 \cdot 10^8 \text{ m.s}^{-1} \times \frac{1,0 \text{ N}}{1,0 \text{ kW.m}^{-2}} = 1,5 \cdot 10^5 \text{ m}^2$$

En approximant cette valeur par  $1,6 \cdot 10^5 \text{ m}^2 = (400 \text{ m})^2$ , nous constatons qu'une voile carrée de 400 m de côté ferait l'affaire. La chose serait certes encombrante et la force obtenue modeste, mais l'espace ne manque pas de place et les missions spatiales sont généralement des projets à long terme.

Notons également qu'une force de valeur 1,0 N est synonyme d'une accélération de  $1,0 \text{ m.s}^{-2}$  pour un système de masse égale à 1,0 kg. Pour une sonde pesant par exemple 500 kg, elle passerait à  $2,0 \text{ mm.s}^{-2}$ , autrement dit la vitesse augmenterait

uniformément de  $2,0 \text{ mm.s}^{-1}$  à chaque seconde écoulée. La valeur peut sembler modeste, mais une journée comprenant de l'ordre de  $10^5 \text{ s}$ , nous constatons qu'en une journée elle augmenterait de presque  $200 \text{ m.s}^{-1}$  ; au bout d'un mois elle aurait gagné dans les  $6 \text{ km.s}^{-1}$  et au bout d'un an dans les  $70 \text{ km.s}^{-1}$ . Encore une fois, au prix de ce carburant et des délais dont nous disposons, l'option patience reste intéressante.

5. L'énoncé nous suggère d'utiliser deux lois de conservation. Notre système est constitué de 2 objets (le photon d'une part, l'électron d'autre part), chacun considéré à 2 dates (avant diffusion et après diffusion). Il peut donc être bon, pour commencer, d'inventorier les expressions des grandeurs considérées pour chacun de ces objets, à chacune de ces dates :

|          |                       | Initiale                                           | Après diffusion                                    |
|----------|-----------------------|----------------------------------------------------|----------------------------------------------------|
| Photon   | Quantité de mouvement | $\vec{p}_{\phi,i} = \frac{h}{\lambda_i} \vec{e}_i$ | $\vec{p}_{\phi,d} = \frac{h}{\lambda_d} \vec{e}_d$ |
|          | Énergie               | $E_{\phi,i} = \frac{hc}{\lambda_i}$                | $E_{\phi,d} = \frac{hc}{\lambda_d}$                |
| Électron | Quantité de mouvement | $\vec{p}_{e,i} = \vec{0}$                          | $\vec{p}_{e,d}$                                    |
|          | Énergie               | $E_{e,i} = m_e c^2$                                | $E_{e,d} = \sqrt{m_e^2 c^4 + p_{e,d}^2 c^2}$       |

Ceci fait, nous pouvons alors exprimer les lois de conservations évoquées :

- Conservation de la quantité de mouvement totale :

$$\vec{p}_{\phi,i} + \vec{p}_{e,i} = \vec{p}_{\phi,d} + \vec{p}_{e,d} \Rightarrow \vec{p}_{e,d} = \frac{h}{\lambda_i} \vec{e}_i - \frac{h}{\lambda_d} \vec{e}_d$$

- Conservation de l'énergie totale :

$$E_{\phi,i} + E_{e,i} = E_{\phi,d} + E_{e,d} \Rightarrow m_e c + \frac{h}{\lambda_i} = \sqrt{m_e^2 c^2 + p_{e,d}^2} + \frac{h}{\lambda_d}$$

La relation demandée engage les longueurs d'onde avant et après diffusion, ainsi que le cosinus de l'angle  $\theta$  formé par les vecteurs unitaires  $\vec{e}_d$  et  $\vec{e}_i$ . En revanche la quantité de mouvement de l'électron après diffusion n'intervient pas. Il semble donc avisé d'éliminer  $p_{e,d}$  entre les deux équations. Exprimons donc  $p_{e,d}^2$  à partir de la première équation :

$$p_{e,d}^2 = \vec{p}_{e,d} \cdot \vec{p}_{e,d} = \frac{h^2}{\lambda_i^2} + \frac{h^2}{\lambda_d^2} - 2 \times \frac{1}{\lambda_i} \times \frac{1}{\lambda_d} \times \underbrace{\vec{e}_i \cdot \vec{e}_d}_{\cos \theta} \Leftrightarrow$$

$$p_{e,d}^2 = h^2 \times \left( \frac{1}{\lambda_i^2} + \frac{1}{\lambda_d^2} - \frac{2 \cos \theta}{\lambda_i \lambda_d} \right)$$

Et faisons de même à partir de la seconde :

$$\sqrt{m_e^2 c^2 + p_{e,d}^2} = m_e c + h \times \left( \frac{1}{\lambda_i} - \frac{1}{\lambda_d} \right) \Leftrightarrow$$

$$p_{e,d}^2 = h^2 \times \left( \frac{1}{\lambda_i} - \frac{1}{\lambda_d} \right)^2 + 2 m_e h c \times \left( \frac{1}{\lambda_i} - \frac{1}{\lambda_d} \right)$$

Il ne nous reste plus alors qu'à mettre en vis-à-vis ces deux expressions :

$$h^2 \times \left( \frac{1}{\lambda_i^2} + \frac{1}{\lambda_d^2} - \frac{2 \cos \theta}{\lambda_i \times \lambda_d} \right) = h^2 \times \left( \frac{1}{\lambda_i} - \frac{1}{\lambda_d} \right)^2 + 2m_e h c \times \left( \frac{1}{\lambda_i} - \frac{1}{\lambda_d} \right)$$

En simplifiant des deux côtés par  $h^2$ , et en éliminant de part et d'autre les termes quadratiques redondants, nous trouvons alors :

$$-\frac{2 \cos \theta}{\lambda_i \times \lambda_d} = -\frac{2}{\lambda_i \times \lambda_d} + \frac{2m_e c}{h} \times \left( \frac{\lambda_d - \lambda_i}{\lambda_i \times \lambda_d} \right)$$

Il ne reste plus qu'à multiplier des deux côtés par  $\frac{\lambda_i \times \lambda_d}{2}$ , et nous obtenons :

$$-\cos \theta = -1 + \frac{m_e c}{h} \times (\lambda_d - \lambda_i) \Leftrightarrow \lambda_d - \lambda_i = \frac{h}{m_e c} \times (1 - \cos \theta) \quad (\text{CQFD})$$

## Vers la prépa

**1.** Le domaine du visible contient l'ensemble des radiations dont les longueurs d'onde sont comprises entre 384 et 768 nm approximativement. La première et la quatrième radiations sont donc visibles. On peut préciser que la deuxième appartient au domaine des infrarouges, tandis que la troisième appartient à celui de l'ultraviolet.

**2.** L'énergie d'un photon associé à une fréquence  $\nu$  est  $E_\varphi = h\nu$ . Puisque la longueur d'onde dans le vide de cette radiation vaut  $\lambda = \frac{c}{\nu}$ , on en déduit que  $E_\varphi(\text{J}) = \frac{hc}{\lambda(\text{m})}$ . En

utilisant le fait que  $\lambda(\text{m}) = 10^{-9} \lambda(\text{nm})$  et que  $1 \text{ eV} = e \text{ J}$ , on obtient finalement  $W(\text{eV}) = \frac{10^9 hc}{e \cdot \lambda(\text{nm})}$ , relation menant bien à la formule indiquée.

Calculons alors les énergies des quatre radiations du spectre :

- pour  $\lambda_1$ , on trouve 1,85 eV ;
- pour  $\lambda_2$ , on trouve 1,53 eV ;
- pour  $\lambda_3$ , on trouve 3,84 eV ;
- pour  $\lambda_4$ , on trouve 2,04 eV.

**3.** Il s'agit d'un spectre d'émission, on a donc  $E_\varphi > 0$  qui correspond à la différence entre le niveau d'énergie le plus élevé et celui d'énergie le moins élevé. Pour une transition d'un niveau d'énergie  $n$  à un niveau d'énergie inférieure  $p$  :

$$E_n = E_p + \frac{1,24 \cdot 10^3}{\lambda} \text{ avec les unités choisies (eV et nm).}$$

On calcule alors les énergies des niveaux du lithium :

- pour le niveau II :

$$n = 2 \text{ et } p = 1, E_{\text{II}} = E_{\text{I}} + \frac{1,24 \cdot 10^3}{671} = -3,54 \text{ eV ;}$$

- pour le niveau III : de manière analogue  $E_{\text{III}} = -2,01 \text{ eV}$  ;
- pour le niveau IV : de manière analogue  $E_{\text{IV}} = -1,55 \text{ eV}$  ;
- pour le niveau V : de manière analogue  $E_{\text{V}} = -1,51 \text{ eV}$ .

**4.** Par convention, l'énergie d'un atome ionisé est nulle. Il suffit donc de fournir à un électron pris au niveau fondamental une énergie égale à  $0 - E_{\text{I}} = +5,39 \text{ eV}$  pour l'ioniser. Ceci correspond à un photon de longueur d'onde  $\lambda = \frac{1,24 \cdot 10^3}{5,39} = 230 \text{ nm}$ , qui appartient au domaine des ultraviolets.



### 9.1 Généralités sur la lumière

#### — Quel est le cadre d'application de l'optique géométrique ?

L'optique géométrique est un modèle dans lequel la lumière est représentée par des **rayons lumineux**. Ceux-ci sont émis par des sources lumineuses et se propagent dans certains milieux (dits transparents), selon des lignes droites tant que le milieu est homogène.

**Remarque :** une manifestation bien connue d'inhomogénéité perturbant la propagation rectiligne des rayons lumineux est la fluctuation de température au voisinage d'une source de chaleur (route goudronnée un jour de grand soleil, air au-dessus d'un barbecue).

Il s'agit d'un modèle simple et purement descriptif, qui en particulier ne nécessite pas d'hypothèse sur la nature de la lumière (onde ou corpuscule, vieux débat). Il apparaît cependant comme une conséquence du modèle ondulatoire, dans le cas particulier où les dimensions des différents objets limitant la propagation de la lumière (lentilles, diaphragmes, etc.) sont très supérieures à la longueur d'onde des radiations utilisées.

**Remarque :** notons l'aspect purement conceptuel du rayon lumineux, qu'il est impossible d'isoler expérimentalement. La chose nécessiterait en effet de faire passer la lumière à travers un trou infiniment petit, ce qui, outre des difficultés matérielles, entraînerait inévitablement la diffraction du rayon.

#### — À quelles conditions un objet est-il visible ?

On appelle **source lumineuse** tout objet d'où partent des rayons lumineux. Cette source est dite primaire si l'objet génère lui-même sa propre lumière, secondaire s'il se contente de réfléchir la lumière d'une autre source.

Un objet lumineux est alors visible depuis un point d'observation, si et seulement si il existe au moins un rayon allant de cet objet au point d'observation en question.

Précisons ici le **principe du retour inverse de la lumière**, selon lequel toute trajectoire suivie par la lumière entre un point source et un point d'observation sera suivie en sens inverse si l'on intervertit les positions de la source et du point d'observation.

#### — Qu'est-ce que l'indice de réfraction d'un milieu transparent ?

La célérité de propagation  $v$  de la lumière dans un milieu matériel diffère de sa célérité dans le vide. Cette dernière, notée  $c$ , a pour valeur  $c = 299792458 \text{ m.s}^{-1}$ . Il s'agit d'une valeur exacte (elle est posée et ne se mesure pas), souvent approximée par la valeur de  $3,00.10^8 \text{ m.s}^{-1}$ .

L'indice de réfraction  $n$  d'un milieu transparent se définit alors comme le rapport :

|                   |                                               |
|-------------------|-----------------------------------------------|
| $n = \frac{c}{v}$ | $c, v$ en $\text{m.s}^{-1}$<br>$n$ sans unité |
|-------------------|-----------------------------------------------|

**Remarque :** la célérité de la lumière ne pouvant excéder sa valeur dans le vide, on a toujours  $n \geq 1$ , l'égalité étant vérifiée uniquement dans le vide.

Quelques valeurs pour fixer les idées :

| Milieu | Vide | Air  | Eau  | Verre | Diamant |
|--------|------|------|------|-------|---------|
| $n$    | 1    | 1,00 | 1,33 | 1,50  | 2,41    |

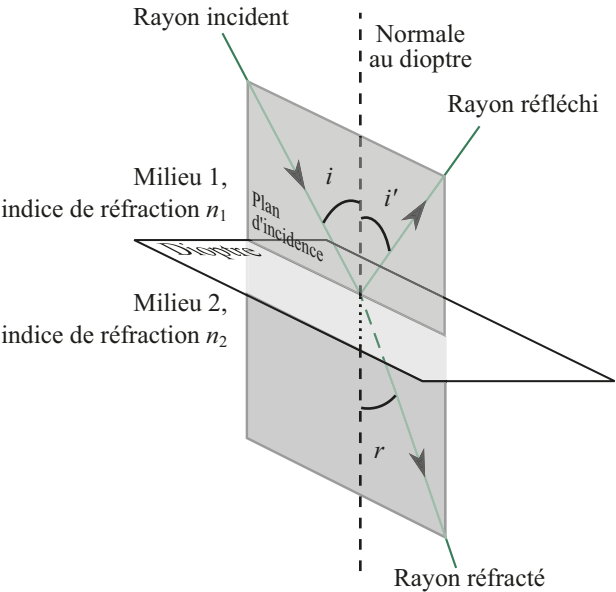
— **Qu'advient-il d'un rayon lumineux lorsqu'il rencontre un dioptre ?**

On appelle **dioptre** l'interface entre deux milieux de nature différente. Lorsqu'un rayon incident rencontre un dioptre, il peut donner naissance à deux nouveaux rayons : l'un qui repart dans son milieu d'origine (rayon dit **réfléchi**), et un autre qui poursuit sa course de l'autre côté du dioptre (rayon dit **réfracté**).

On constate alors que le rayon réfracté se propage dans une direction différente de celle du rayon incident. Ce changement d'orientation au franchissement d'un dioptre constitue le phénomène de réfraction.

— **Qu'énoncent les lois de Snell-Descartes ?**

Les **lois de Snell-Descartes** prévoient le comportement géométrique des rayons lumineux à l'interface entre deux milieux transparents. On utilisera les notations définies sur la figure ci-dessous:



On peut regrouper les lois de Snell-Descartes en trois énoncés :

- Les rayons réfléchi et réfracté appartiennent au plan d'incidence (plan défini par le rayon incident et la normale au dioptre au point d'incidence).
- L'angle de réflexion est égal à l'angle d'incidence :

|          |                        |
|----------|------------------------|
| $i' = i$ | $i, i'$ en ° ou en rad |
|----------|------------------------|

- Le produit de l'indice de réfraction du milieu (1) par le sinus de l'angle d'incidence est égal à celui de l'indice de réfraction du milieu (2) par le sinus de l'angle de réfraction :

|                           |                                                    |
|---------------------------|----------------------------------------------------|
| $n_1 \sin i = n_2 \sin r$ | $i, r$ en ° ou en rad<br>$n_1$ et $n_2$ sans unité |
|---------------------------|----------------------------------------------------|

**Remarque :** on montre que lorsque l'on passe d'un milieu à un autre plus réfringent ( $n_2 > n_1$ ), le rayon réfracté se rapproche de la normale : dans le cas contraire, il s'éloigne de la normale.

## — Qu'est-ce que le phénomène de réflexion totale ?

Lorsqu'un rayon lumineux passe d'un milieu à un autre moins réfringent, le rayon réfracté s'éloigne de la normale au dioptre par rapport au rayon incident. On va donc pouvoir trouver une valeur de l'angle d'incidence  $i_L$  correspondant à un angle de réfraction égal à  $\frac{\pi}{2}$  (situation dite d'*émergence rasante*). Cette valeur s'exprime :

|                                             |                      |
|---------------------------------------------|----------------------|
| $i_L = \arcsin\left(\frac{n_2}{n_1}\right)$ | $i_L$ en ° ou en rad |
|---------------------------------------------|----------------------|

Si on impose un angle d'incidence encore supérieur, l'angle du rayon réfracté avec la normale n'est plus défini (il supposerait  $\sin r > 1$ ). Il n'existe alors plus de rayon réfracté et seul demeure le rayon réfléchi. On appelle ce phénomène la **réflexion totale**. Elle est utilisée dans certains dispositifs tels que les **fibres optiques** à saut d'indice.

## — Qu'est-ce qu'un prisme ?

Un **prisme** est un milieu transparent (généralement du verre) délimité par deux dioptries plans formant entre eux un angle  $A$ , appelé **angle au sommet** du prisme. Les prismes usuels ont un angle au sommet de  $60^\circ$ . Ils permettent de faire passer la lumière par deux dioptries successifs (air/verre puis verre/air), et ce faisant de la réfracter deux fois. Leur principal intérêt est d'accentuer le phénomène de **dispersion** qui accompagne souvent (de manière discrète) celui de réfraction de la lumière, et sont donc couramment utilisés pour obtenir le spectre d'une lumière.

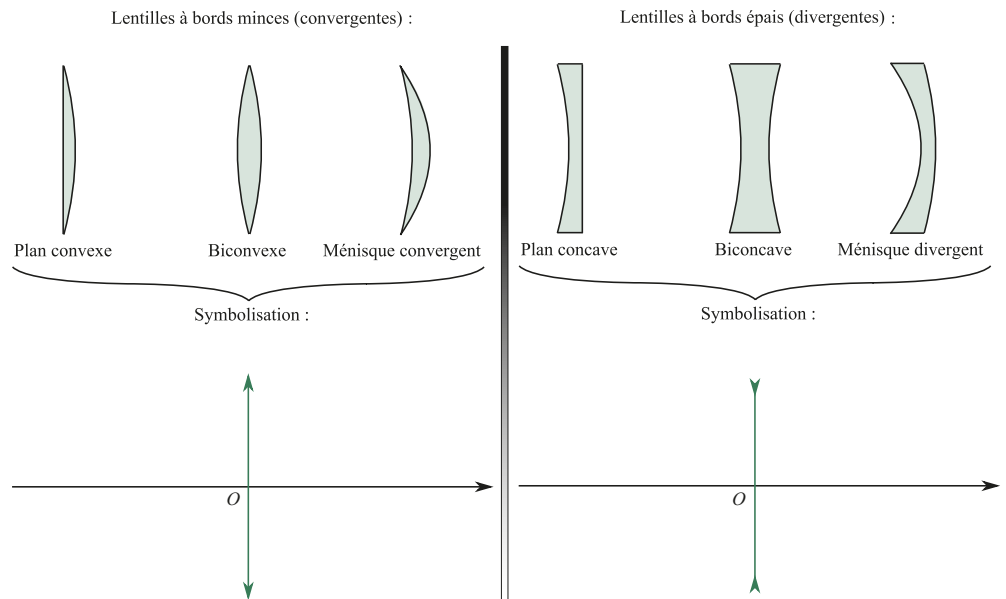
## 9.2 Les lentilles minces

### Qu'est-ce qu'une lentille ? Quels sont les différents types de lentilles ?

Une **lentille** est constituée par un milieu transparent délimité par deux dioptries sphériques. On distingue deux catégories de lentilles :

- les lentilles dites à **bords minces** (c'est-à-dire plus minces que le centre), qui constituent des systèmes convergents ;
- les lentilles dites à **bords épais** (c'est-à-dire plus épais que le centre), qui constituent des systèmes divergents.

**Remarque :** indépendamment du caractère mince ou épais des bords par comparaison à l'épaisseur du centre de la lentille, une lentille est dite mince dès lors que l'épaisseur de sa partie centrale est faible par rapport aux rayons de courbure des dioptries délimitant la lentille. Nous supposons toujours cette hypothèse satisfaite dans le cadre de cet ouvrage et envisagerons, outre les lentilles minces convergentes vues en terminale, également les lentilles minces divergentes.



### Quelles sont les caractéristiques géométriques d'une lentille mince ?

Une lentille mince possède en tout 7 caractéristiques géométriques :

- **un axe optique** ( $\Delta$ ) ; dans les faits il s'agit de l'axe de symétrie de révolution de la lentille. Dans le cadre du modèle des lentilles minces, il se résume à une droite passant par le milieu de la représentation symbolique de la lentille, et normale au plan de celle-ci ;
- **trois points** particuliers, présentant des propriétés qui nous permettront de lier un point objet et son image conjuguée par la lentille :

- le **centre optique**  $O$  de la lentille. Dans le cadre du modèle des lentilles minces, il se situe à l'intersection de la représentation symbolique de la lentille avec son axe optique.
- le **foyer objet**  $F$  de la lentille, situé à gauche de la lentille si elle est convergente, à droite si elle est divergente,
- le **foyer image**  $F'$  de la lentille, situé à droite de la lentille si elle est convergente, à gauche si elle est divergente.

Les foyers  $F$  et  $F'$  sont symétriques l'un de l'autre par rapport au centre optique  $O$ ;

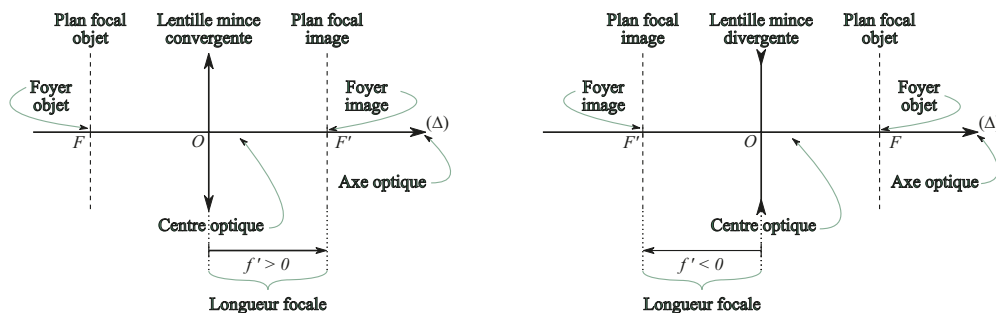
- **une longueur** appelée longueur focale :

$$f' = \overline{FO} = \overline{OF'}$$

Cette longueur focale est algébrique, c'est-à-dire qu'elle contient son propre signe. Ainsi la longueur focale d'une lentille mince convergente est-elle positive, tandis que celle d'une lentille mince divergente est négative ;

- **deux plans** :

- le **plan focal objet**, perpendiculaire à l'axe optique et contenant le foyer objet  $F$ ,
- le **plan focal image**, perpendiculaire à l'axe optique et contenant le foyer image  $F'$ .



Les exercices de CPGE vous fourniront parfois, au lieu de la longueur focale d'une lentille, une autre grandeur appelée vergence (notée  $V$  ou  $C$ , en général). Celle-ci est définie comme l'inverse de la longueur focale. Son USI est donc le  $\text{m}^{-1}$ , que l'on remplace usuellement par la dioptrie (symbole  $\delta$ , USI également :  $1\delta = 1\text{m}^{-1}$ ).

$$V = \frac{1}{f'}$$

|      |             |
|------|-------------|
| $V$  | en $\delta$ |
| $f'$ | en m        |

La vergence apparaît donc comme une mesure du pouvoir de convergence d'une lentille : si les rayons convergent rapidement après la lentille (focale courte), c'est que celle-ci est très convergente, et réciproquement.

**Remarque** : les valeurs des corrections des verres de lunettes ou des lentilles de contact sont, typiquement, fournies en dioptries. Par exemple un verre vergence  $V = 2,0\delta$  n'est ni plus ni moins qu'une lentille convergente de focale  $f' = 50\text{ cm}$ .

Nous traiterons par la suite uniquement les lentilles minces convergentes, cependant le fonctionnement des lentilles divergentes repose sur les mêmes principes, n'hésitez pas à vous y essayer dès à présent.

## Quelles sont les propriétés générales des lentilles minces ?

Une lentille mince est donc constituée de deux dioptries, qui vont permettre une double réfraction de la lumière, comme dans un prisme. Cependant la forme sphérique de ces dioptries fait que chaque rayon va rencontrer l'équivalent d'un dioptre d'angle au sommet plus ou moins prononcé : les faces sont d'autant moins inclinées l'une par rapport à l'autre, que l'on se situe plus proche de l'axe optique. La réfraction se fera donc selon une direction différente pour chaque rayon, permettant ainsi de tous les réorienter de sorte qu'ils se rencontrent à nouveau (on dit dans ce cas des rayons qu'ils **convergent** vers un nouveau point). Ce point de convergence, dont la position sera donc conditionnée par la forme de la lentille, constituera alors l'**image** par la lentille du point objet dont ils sont issus.

Dans les faits, cette convergence est la plupart du temps imparfaite. Cependant le modèle que nous allons utiliser permet de supposer qu'elle l'est, et nous admettrons par la suite que les lentilles minces :

- associent à tout point objet  $A$  un point image  $A'$  (hypothèse de **stigmatisme**) ;
- associent à tout objet  $AB$  perpendiculaire à l'axe optique, une image  $A'B'$  également perpendiculaire à l'axe optique (hypothèse d'**aplanétisme**).

## Où l'infini d'une lentille se situe-t-il (à part très loin) ?

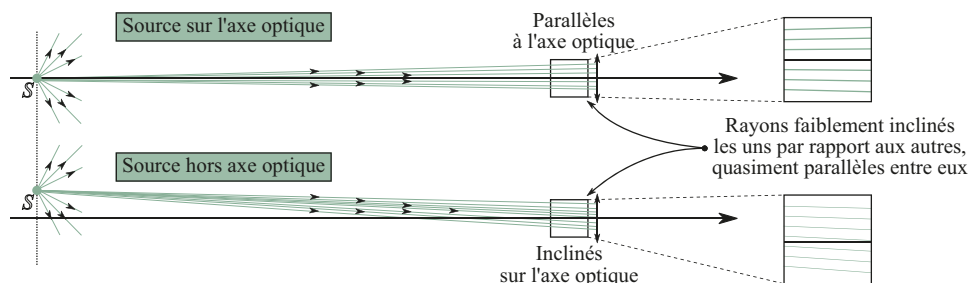
Les propriétés que nous allons voir évoqueront souvent des objets ou des images situés à l'infini de la lentille. Cette condition est évidemment irréalisable en pratique, mais surtout n'a pas besoin de l'être. La physique vise à poser des modèles cherchant à rendre compte au mieux de la réalité, et s'appuie la plupart du temps sur des approximations.

Dans le cas des lentilles minces, on pourra ainsi d'autant mieux considérer qu'une position se situe à l'infini d'une lentille, que sa distance à cette lentille sera plus grande devant la longueur focale de cette dernière.

Concrètement, pour une position  $A$ , on aura :

| Valeur réelle de $\frac{OA}{f'}$                             | $10^1$                                                  | $10^2$                                                 | $10^3$                                                   |
|--------------------------------------------------------------|---------------------------------------------------------|--------------------------------------------------------|----------------------------------------------------------|
| Qualité de l'approximation<br>$\frac{OA}{f'} \approx \infty$ | Écart entre expérience et modèle d'environ <b>10</b> %. | Écart entre expérience et modèle d'environ <b>1</b> %. | Écart entre expérience et modèle d'environ <b>0,1</b> %. |

Par ailleurs, on admettra que des rayons se croisant en un même point (source ou image) sont parallèles entre eux, si et seulement si ce point est situé à l'infini. En effet, le rayon (au sens circulaire) d'une lentille mince est du **même ordre de grandeur que sa longueur focale**. La distance d'un point à l'infini d'une lentille mince étant très supérieure à la longueur focale de cette lentille, elle est également très supérieure à ce rayon. Les seuls rayons lumineux passant par cette lentille et se rencontrant au niveau de ce point de croisement seront donc très peu inclinés les uns par rapport aux autres, et pourront d'autant mieux être considérés comme parallèles entre eux, que leur point de croisement sera plus éloigné.



Ajoutons que ces rayons, parallèles entre eux, seront parallèles à l'axe optique si et seulement si le point de croisement se situe sur l'axe optique ; dans le cas contraire, ils parviendront (ou quitteront) à la lentille inclinés par rapport à l'axe optique.

## — Comment construit-on l'image d'un objet par une lentille mince ?

Nous avons admis plus haut que les lentilles minces associaient à tout faisceau de rayons issus d'un point source commun  $A$ , un point de convergence  $A'$  constituant l'image de ce point source. La lentille réoriente donc tous les rayons afin qu'ils se dirigent vers ce point de convergence. La direction dans laquelle un rayon quelconque va émerger de la lentille est déterminable, mais au prix de calculs complexes.

Cependant 3 rayons parmi la multitude issue du point source, vont présenter des trajectoires particulièrement simples à anticiper. En construisant ces 3 rayons, nous devrions obtenir un point d'intersection correspondant, donc, à l'image recherchée. En admettant que la lentille associe à tout point objet un autre point image, nous pourrions ensuite si nécessaire construire les autres rayons en les faisant passer par ce point image lorsqu'ils émergent de la lentille.

**Remarque :** pour trouver un point d'intersection, 2 rayons seulement suffisent. Cependant les rayons sont souvent très peu inclinés et les zones d'intersection s'étalent sur plusieurs millimètres, entraînant une précision médiocre. La construction des 3 rayons définit non plus seulement une, mais trois intersections qui, donc, doivent coïncider et ce faisant permettent de mieux cadrer le lieu de convergence.

Ces rayons remarquables sont :

| Tout rayon incident dont la direction... | ...émerge de la lentille selon une direction... |
|------------------------------------------|-------------------------------------------------|
| ... passe par le centre optique $O$ ...  | ... identique à sa direction initiale.          |
| ... passe par le foyer objet $F$ ...     | ... parallèle à l'axe optique.                  |
| ... est parallèle à l'axe optique...     | ... passant par le foyer image $F'$ .           |

**Remarque :** ces propriétés sont valables aussi bien pour les lentilles convergentes que pour les lentilles divergentes (en restant attentif/ve au fait que nous évoquons ci-dessus les **directions** des rayons, et pas nécessairement les rayons eux-mêmes). La seule différence dans la méthode de construction d'une image avec ces dernières, provient du fait que les positions des foyers sont inversées. Il peut en résulter une certaine confusion au début, mais après un peu de pratique, leur usage n'est pas plus difficile.

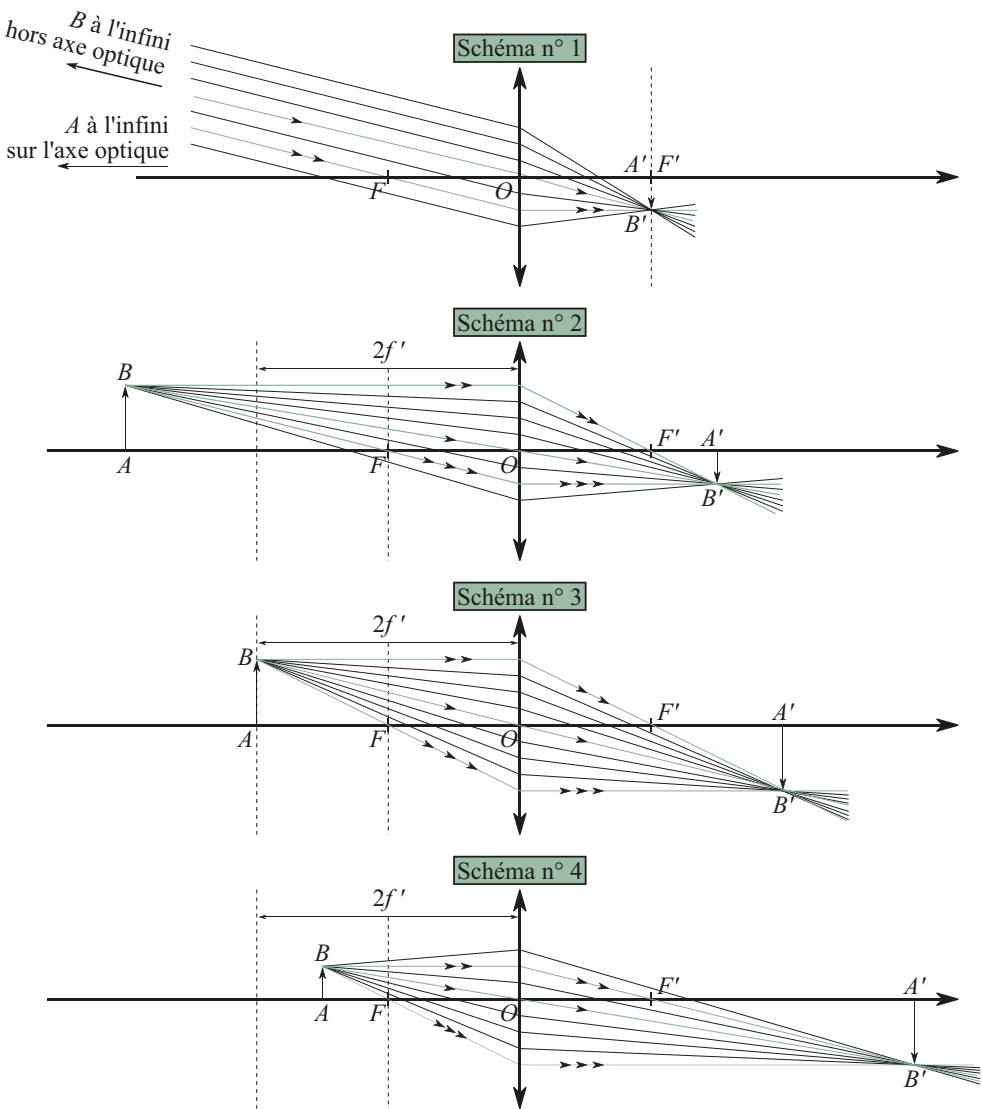
Nous constatons alors que plusieurs cas de figure peuvent se présenter :

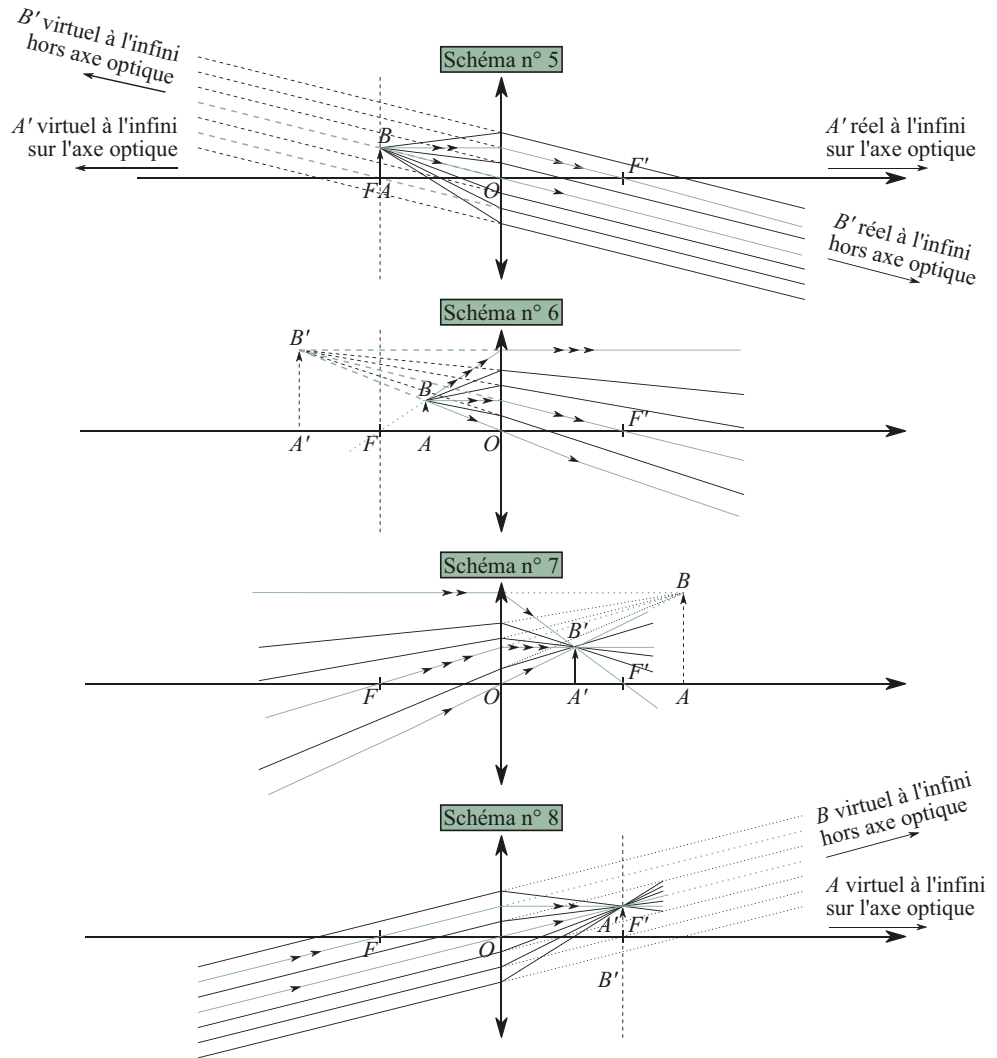
- tant que l'objet se situe **à gauche du plan focal objet**, il existe toujours un point d'intersection commun aux 3 rayons émergents remarquables. L'image conjuguée de l'objet se

situé donc à droite de la lentille (dans ce que l'on appelle l'*espace image* de cette lentille), et il suffit d'intercepter le faisceau lumineux à l'aide d'un écran au niveau du point de convergence, pour visualiser l'image. Une image ainsi projetable est dite **réelle**. Ajoutons qu'elle est toujours **renversée** par rapport à l'objet. Nous pouvons apporter quelques nuances concernant les différentes positions possibles à gauche du plan focal objet :

- si l'objet se situe **à l'infini** (schéma n° 1), l'image se forme dans le plan focal image (notons au passage que sa construction repose, dans ce cas précis, sur 2 rayons remarquables seulement),
- si l'objet se situe à distance finie de la lentille, on distingue 3 régions :

| Schéma n° 2                   | Schéma n° 3                                                         | Schéma n° 4                   |
|-------------------------------|---------------------------------------------------------------------|-------------------------------|
| $\overline{OA} < -2f'$        | $\overline{OA} = -2f'$                                              | $-2f' < \overline{OA} < -f'$  |
| Image plus grande que l'objet | Image de même taille que l'objet (configuration dite de Silbermann) | Image plus petite que l'objet |





- si l'objet se situe **dans le plan focal objet** (schéma n° 5), les rayons issus d'un point objet émergent parallèles entre eux (ici encore, 2 rayons particuliers seulement sont exploitables) : l'image se forme donc à l'infini de la lentille ;
- si l'objet se situe **entre le plan focal objet et la lentille** (schéma n° 6), les rayons émergents divergent, et il semble donc vain d'espérer trouver un point d'intersection. On constate cependant que si l'on trace le prolongement vers la gauche de ces rayons, leurs directions présentent bien un point d'intersection. La lumière ne rebrousse pas chemin après avoir été réfractée, et ce point d'intersection ne peut donc cette fois être intercepté par un écran. Mais si un(e) observateur/trice positionne son œil à droite de la lentille et regarde à travers, il/elle va recevoir un faisceau de rayons lumineux identique à ce que lui aurait envoyé un point objet positionné au niveau de leur point d'intersection. Il existe donc bien une image dans ce cas, mais elle n'est pas matérialisable comme dans les cas précédents : l'image est cette fois dite **virtuelle**, et on le représente en tirets. Notons qu'elle est également droite et que depuis tout point d'observation, elle apparaît sous un angle plus grand que l'objet de départ. Une lentille utilisée dans cette configuration constitue donc simplement une loupe.

**Remarque :** l'image obtenue dans ce cas est plus grande que l'objet, mais également plus éloignée. C'est donc bien l'angle sous lequel sont vus l'un et l'autre qui justifie le fait qu'un objet soit vu plus grand à travers une loupe ; le seul argument de la taille ne suffit pas.

**Remarque :** dans le cas d'un objet situé dans le plan focal objet, l'image est tout à la fois réelle (elle apparaîtra sur un écran positionné à l'infini) et virtuelle (un(e) observateur/trice regardant directement à travers la lentille recevra un faisceau de rayons parallèles entre eux semblant venir d'un objet situé à l'infini, comme dans le cas précédent.

Les schémas n<sup>os</sup> 7 et 8 illustrent une notion que vous n'avez probablement pas vue en terminale, mais qui intervient régulièrement dans les systèmes à plusieurs lentilles qui constitueront votre pain quotidien en CPGE, à savoir les **objets virtuels**.

Vous avez déjà vu des systèmes à plusieurs lentilles (la lunette astronomique par exemple), et connaissez le principe de construction d'une image à travers plusieurs systèmes optiques successifs :

- le premier système forme l'image conjuguée  $A_1B_1$  de l'objet  $AB$  observé ;
- cette image sert d'objet au système suivant, qui en donne une image  $A_2B_2$  ;
- et ainsi de suite jusqu'au dernier système optique qui forme l'image finale.

Il arrive cependant que l'un de ces systèmes donne une image située *après* la lentille suivante. Celle-ci se retrouve donc à devoir former l'image d'un objet situé à sa droite, qui cependant n'a pas d'existence matérielle (d'où le nom d'objet virtuel, et sa représentation en tirets).

Dans ce cas, on procède simplement en représentant l'image qu'aurait formée le système optique précédent en l'absence de la lentille (en pointillés sur le schéma). On repère ensuite, parmi tous les rayons qui convergeaient pour former cette image, nos 3 rayons remarquables.

Il ne reste plus alors qu'à remplacer, au-delà de la lentille, les rayons de départ par leur version réorientée.

Sur le schéma n<sup>o</sup> 8, nous avons représenté le cas d'un objet virtuel situé à l'infini. Les rayons sont donc parallèles entre eux et supposés converger à l'infini vers la droite. Cependant le parallélisme des rayons fait qu'ils convergent également à l'infini vers la gauche. La situation est donc identique au schéma n<sup>o</sup> 1, et nous retrouvons fort heureusement la formation d'une image dans le plan focal image.

**Remarque :** comme vous l'aurez noté, les points objets sont souvent notés par des lettres majuscules, et leurs images conjuguées respectives par les mêmes lettres dotées d'un « prime ». Une **erreur** d'étourderie courante après une période d'inactivité en optique géométrique est de considérer que  $F'$  est l'image conjuguée de  $F$ . Il importe tout d'abord de conserver à l'esprit que  $F$  et  $F'$  ne sont pas des points objets ou des points images, mais **seulement des positions sur l'axe optique**. Il est en revanche vrai que le foyer image est le lieu où se forme l'image conjuguée d'un point objet situé à l'infini sur l'axe optique, et que le foyer objet est la position où doit être placé un point objet pour que son image se situe à l'infini sur l'axe optique. Donc si l'on tient absolument à associer des points conjugués aux foyers, les points en question se situent à l'infini, et non pour chaque foyer au niveau de l'autre foyer.

## — Comment caractérise-t-on la transformation subie entre un objet et son image conjuguée ?

Classiquement, on caractérise le rapport entre un objet et l'image qu'en donne une lentille mince par le grandissement  $\gamma$  subi :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}}$$

$$\frac{\overline{A'B'}}{\overline{AB}} \quad \text{en m}$$

**Remarque :** comme rapport de deux longueurs,  $\gamma$  n'a pas d'unité.

Le grandissement étant une valeur algébrique, il permet d'apprécier deux aspects de l'image par rapport à l'objet :

| Caractéristiques de l'image  | $\gamma < -1$ | $-1 < \gamma < 0$ | $0 < \gamma < 1$ | $1 < \gamma$ |
|------------------------------|---------------|-------------------|------------------|--------------|
| Sens par rapport à l'objet   | Renversée     |                   | Droite           |              |
| Taille par rapport à l'objet | Plus grande   | Plus petite       |                  | Plus grande  |

Cependant, dans le cas d'un objet ou d'une image situé(e) à l'infini, le grandissement devient inutilisable (nul pour un objet à l'infini, infini pour une image à l'infini). Il reste pourtant possible de comparer l'objet et l'image qu'en donne le système optique, mais dans ce cas on ne compare plus les tailles, mais les angles sous lesquels apparaissent respectivement l'objet et l'image (comme dans le cas de la loupe évoqué plus haut).

On a alors recours à une autre grandeur appelée **grossissement**  $G$  et définie comme le rapport de l'angle  $\alpha'$  sous lequel est vue l'image à travers le système, à la plus grande valeur  $\alpha$  possible de l'angle sous lequel est observé l'objet à l'œil nu.

$$G = \frac{\alpha'}{\alpha}$$

$$\alpha', \alpha \text{ en rad ou en } ^\circ.$$

## — Comment peut-on calculer la position de l'image conjuguée d'un objet par une lentille mince ?

Les rayons remarquables utilisés pour construire les image présentent des propriétés géométriques particulières (parallélisme, etc.), qui peuvent être exploitées pour obtenir des relations analytiques entre les positions d'un objet et de son image conjuguée. Ces relations (démontrées essentiellement à partir du théorème de Thalès) sont appelées **relations de conjugaison** ; il en existe plusieurs, nous citerons ici les plus couramment rencontrées en CPGE :

- relation de conjugaison liée au grandissement :

$$\gamma = \frac{\overline{OA'}}{\overline{OA}}$$

$$\overline{OA}, \overline{OA'} \text{ en m.}$$

- relation de conjugaison de Descartes :

$$\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{f'}$$

$$f', \overline{OA'}, \overline{OA} \text{ en m.}$$

**Remarque :** notons au passage que cette relation peut se réécrire :

$$\overline{OA'} = \frac{f' \times \overline{OA}}{f' + \overline{OA}}$$

Ce qui nous permet d'exprimer le grandissement en fonction  $\overline{OA}$  et de  $f'$  :

$$\gamma = \frac{\overline{OA'}}{\overline{OA}} = \frac{f'}{f' + \overline{OA}}$$

- relation de conjugaison de Newton :

$$\overline{F'A'} \times \overline{FA} = -f'^2$$

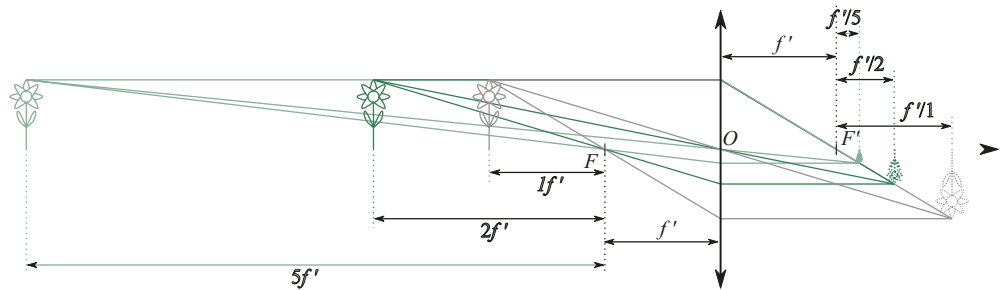
$$\overline{FA}, \overline{F'A'}, f' \text{ en m}$$

Cette dernière relation peut sembler plus compliquée à mettre en œuvre que la précédente (et de fait, elle l'est légèrement) puisque les positions de l'objet et de son image conjuguée ne sont pas repérées par rapport à une origine commune ( $O$  dans la relation de Descartes), mais l'une par rapport au foyer objet  $F$  et l'autre par rapport au foyer image  $F'$ .

En effet, elle peut se réécrire :

$$\left( \frac{\overline{F'A'}}{f'} \right) = - \left( \frac{\overline{FA}}{f'} \right)$$

En d'autres termes, si l'on exprime  $\overline{FA}$  et  $\overline{F'A'}$  non en unités de longueurs usuelles mais en nombres de longueurs focales, la distance de l'image au foyer image varie en raison inverse de celle de l'objet au foyer objet.



L'utilisation de cette relation demande un peu de gymnastique mentale au début, mais une fois maîtrisée elle permet de réduire significativement les quantités de calcul et leur complexité, et d'intuiter beaucoup plus vite les rouages de la plupart des problèmes. Elle deviendra votre principal cheval de bataille en CPGE pour ce qui concerne les lentilles minces.

## 9.3 Les instruments d'optique

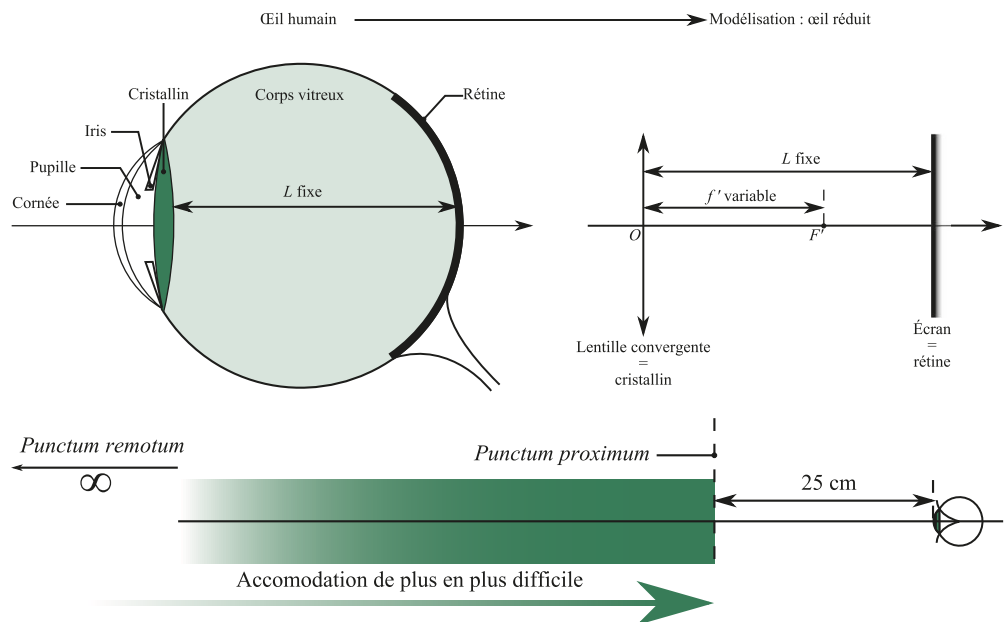
### Quelles sont les caractéristiques de l'œil humain et comment peut-on le modéliser ?

L'**œil humain** est un système optique convergent de **distance focale réglable**. On peut le modéliser par une lentille mince convergente formant l'image des objets sur la rétine, laquelle est pour sa part modélisée par un écran situé à distance fixe de cette lentille.

L'œil voit une image nette si celle-ci se forme sur la rétine. Si la position de l'objet varie, l'œil doit adapter sa longueur focale pour continuer à le voir net, c'est ce que l'on appelle **accommoder**. On appelle alors :

- **punctum remotum** (PR) la position conjuguée de la rétine par la lentille lorsque l'œil est totalement relâché (absence d'accommodation) ;
- **punctum proximum** (PP) la position conjuguée de la rétine par la lentille lorsque l'œil accommode au maximum.

Pour un œil emmétrope (c'est-à-dire sans défaut), le PR se situe à l'infini (vision nette des objets lointains en l'absence d'accommodation). La position du PP varie beaucoup avec l'âge (très proche pour les jeunes, puis de moins en moins à mesure que le sujet vieillit) ; on considère souvent que sa position est située à une distance moyenne d'environ 25 cm de la face de l'œil.



### Quels sont les principaux défauts de l'œil, et comment les corriger ?

Un œil est emmétrope (c'est-à-dire sans défaut) si son PP se situe à l'infini et son PR entre 20 et 30 cm, autrement dit si la convergence de sa face d'entrée est telle, qu'au repos le foyer image se situe sur sa rétine, et que son pouvoir d'accommodation est suffisant pour lui permettre de voir des objets proches de 20 à 30 cm.

Il existe 4 principaux défauts de l'œil :

- **La myopie** : l'œil myope est **trop convergent** par rapport à sa profondeur. Son PR est à distance finie et son PP plus proche que pour un œil emmetrope. On corrige donc ce défaut avec une lentille divergente.
- **L'hypermétropie** : l'œil hypermétrope n'est **pas assez convergent** par rapport à sa profondeur. Il doit déjà utiliser une partie de son pouvoir d'accommodation pour voir net à l'infini ; il lui en reste donc moins pour voir de près, et son PP est plus éloigné que celui d'un œil emmetrope. On corrige donc ce défaut avec une lentille convergente.
- **La presbytie** : il s'agit d'un défaut dû au vieillissement de l'œil, qui consiste en une **diminution du pouvoir d'accommodation**. Ce problème d'accommodation concerne donc uniquement la vision de près. La correction réclame à la base des lentilles convergentes pour suppléer à la faiblesse du pouvoir d'accommodation, cependant en l'absence d'autre défaut de vision, celles-ci ne peuvent être portées que pour la vision de près.
- **L'astigmatisme** : repose sur une **asymétrie de révolution** du cristallin autour de l'axe optique. Il n'assure alors plus la convergence de tous les rayons issus d'un même point objet en un seul et même point image ; la convergence n'est donc qu'approximative, et la vision n'est pas nette. Ce problème se corrige avec des lentilles spéciales dont les dioptries ne sont plus sphériques.

## — Sur quelle mécanique la perception des couleurs repose-t-elle ?

Comme nous l'avons rappelé précédemment, l'image des objets observés par l'œil se forme sur la rétine. À la surface de la rétine se trouvent deux catégories de cellules photosensibles capables de transformer les informations lumineuses leur parvenant en signaux électriques transmis au cerveau. Ces cellules sont :

- les **bâtonnets**, situés à la périphérie de la rétine, insensibles à la couleur mais intervenant aux faibles luminosités (et permettant donc la vision nocturne) ;
- les **cônes**, situés plus près du centre de la rétine dans une zone nommée fovéa, permettent de différencier les couleurs. Il existe trois sortes de cônes (une sensible aux courtes longueurs d'onde et procurant au cerveau la sensation de bleu, une autre aux longueurs d'onde moyennes et procurant la sensation de vert, enfin une troisième sensible aux grandes longueurs d'onde et donnant la sensation de rouge). Ces couleurs directement liées aux cônes sont les **couleurs primaires**. Les autres couleurs (par exemple le jaune, le magenta et le cyan) ne sont pas directement vues en tant que telles mais reconstituées par le cerveau lorsque plusieurs cônes de types différents sont stimulés simultanément sur une même zone de la rétine.

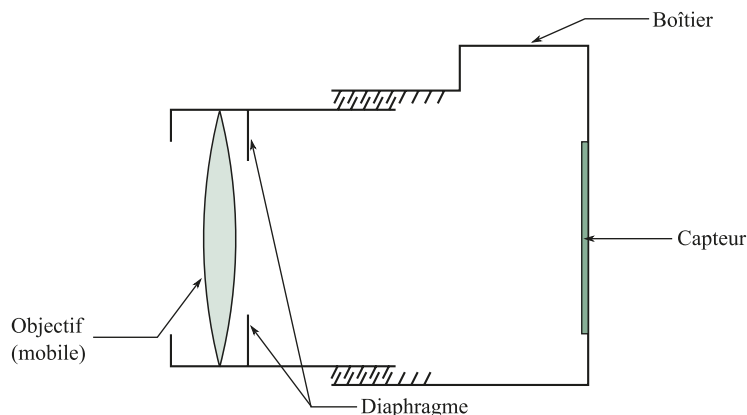
En résumé, si la lumière excite :

- une seule catégorie de cônes, l'œil perçoit une nuance d'une couleur primaire ;
- deux catégories de cônes, l'œil perçoit une nuance d'une couleur secondaire, avec prédominance de la sensation fournie par celle des deux catégories qui est la plus stimulée ;
- trois catégories de cônes, l'œil perçoit une nuance de gris pouvant aller du noir (absence de stimulation) au blanc (saturation de tous les cônes), éventuellement teinté des sensations fournies par les catégories de cônes les plus stimulées.

Le **daltonisme** est une maladie due à un dysfonctionnement d'un ou plusieurs type(s) de cônes, entraînant le plus souvent la confusion entre plusieurs couleurs.

## — Comment un appareil photographique fonctionne-t-il ?

L'appareil photographique peut être modélisé par le schéma suivant :



Un certain parallèle peut être établi avec l'œil :

- le diaphragme est analogue à la pupille ;
- l'objectif est analogue au cristallin ;
- le film ou le capteur est analogue à la rétine.

Les deux systèmes optiques forment des images réelles, sur la rétine pour l'œil, sur le capteur (film photographique ou capteur numérique) pour l'appareil photographique. Il existe néanmoins une grande différence entre les deux au niveau de l'accommodation (appelée *mise au point* dans le cas de l'appareil photographique). En effet, alors que le cristallin se déforme pour adapter la vergence de la face d'entrée de l'œil, la mise au point est quant à elle réalisée par modification de la distance objectif capteur.

**Remarque :** certains objectifs photographiques appelés zooms proposent une focale variable, en offrant la possibilité de modifier la distance entre certaines des lentilles minces qui le composent.

## — Comment construire et calculer les caractéristiques d'une image à travers un système optique composite ?

Les systèmes optiques que nous avons traités jusqu'ici (loupe, œil, appareil photographique) pouvaient être modélisés de manière très simple (une lentille mince unique, éventuellement un écran, point à la ligne). Cependant d'autres systèmes réclament davantage de composants, et vont imposer à la lumière le passage par plusieurs systèmes intermédiaires (lentilles, miroirs, etc.).

Dans ce cas se pose la question de savoir comment traiter ce parcours. Sur le principe, la méthode est très simple et repose, tant pour la construction que pour le calcul, sur un seul et même principe :

- on détermine les caractéristiques de l'image  $A_1B_1$  conjuguée de l'objet  $AB$  de départ, par le premier système optique ;

- cette image sert d'objet pour le deuxième système optique, et l'on détermine les caractéristiques de l'image  $A_2B_2$  conjuguée de l'objet  $A_1B_1$  par ce deuxième système optique ;
- on poursuit de même jusqu'au dernier système optique, qui fournira l'image finale.

Il ne s'agit pas d'une recette magique, mais seulement de logique : qu'un point image soit réel ou virtuel, il entraîne toujours la production d'un faisceau de rayons lumineux en aval du point de convergence où il se situe. Du point de vue phénoménologique, cette situation est identique si l'on remplace ce point de convergence par une source matérielle.

La difficulté de ce type d'exercice repose donc essentiellement sur la mise en œuvre méthodique et rigoureuse des constructions et des calculs :

- les constructions devront se faire en prenant à chaque fois en compte **uniquement le système optique courant**, en faisant abstraction des autres systèmes, traits de construction issus des constructions précédentes, etc. Il est donc essentiel de réaliser des schémas clairs et soignés ;
- de même, les centres optiques, foyers, etc. intervenant dans les différentes relations de conjugaison changeront d'un système optique à l'autre. On aura donc soin d'**indiquer soigneusement** les points, longueurs focales et autres caractéristiques des différents systèmes optiques, afin d'éviter tout risque de confusion.

## — Quel est le principe de fonctionnement de la lunette astronomique ?

Vous avez dû traiter le cas de la lunette astronomique en classe de terminale. Rappelons son cahier des charges :

- il s'agit d'un instrument destiné à observer des objets lointains, vus sous un angle de petite valeur, et généralement peu lumineux ;
- l'image finale est destinée à un œil qui devra l'observer sur de longues durées ; l'image doit donc être formée à l'infini afin d'éviter un effort d'accommodation soutenu.

Notre système doit donc former à l'infini l'image d'un objet également situé à l'infini. Si l'on souhaitait définir des foyers globaux, ils se trouveraient donc tous deux à l'infini. Il n'est donc pas possible de donner une mesure de la longueur focale, ce qui vaut à cet instrument le qualificatif de lunette **afocale**.

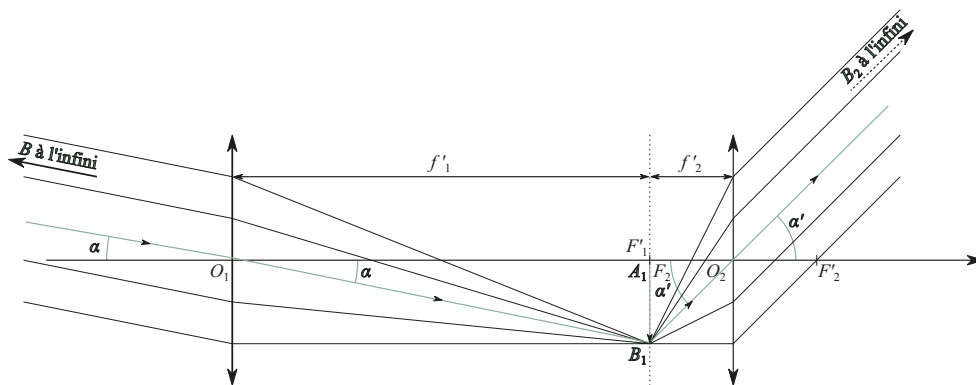
Une manière simple de réaliser cette situation consiste à utiliser deux lentilles, et à faire coïncider le foyer image de la première (où se forme l'image de l'objet observé, puisque celui-ci se situe à l'infini) avec le foyer objet de la seconde. L'image intermédiaire se trouvera ainsi directement dans le plan focal objet de la seconde lentille, qui donc en donnera une image finale positionnée à l'infini.

Concrètement, on appelle :

- **objectif** (focale  $f_1'$ ) la première lentille ; ses dimensions jouent un rôle essentiel à plusieurs niveaux :
  - c'est elle qui collecte la lumière ; or la quantité de lumière qui servira à former l'image finale est proportionnelle à la surface de collecte,
  - ses bords constituent l'obstacle qui va limiter l'entrée de la lumière ; on peut montrer que la monture de cette lentille sera le principal artisan de la diffraction. Les objectifs mesurent en général plusieurs centimètres (voire dizaines de centimètres) de diamètre, et les écarts angulaires dus à la diffraction sont donc faibles. Cependant, s'appliquant à des images dont le diamètre angulaire est lui-même en général très petit, ces effets peuvent rapidement se trouver à hauteur dudit diamètre angulaire, et l'image s'en trouver signi-

ficativement altérée. Les effets diffractifs seront donc d'autant moins sensibles que le diamètre de l'objectif sera plus grand,

- bilan : plus le rayon de la lentille est important, plus l'image sera lumineuse et moins la diffraction sera sensible (et donc plus l'astronome content(e)) ;
- **oculaire** (focale  $f'_2$ ) la seconde lentille, qui aura pour rôle de former une image finale à l'infini.



**Remarque :** pour déterminer la position de l'image conjuguée d'un point objet situé à l'infini d'une lentille (respectivement dans le plan focal objet de celle-ci), un seul rayon suffit : celui passant par le centre optique. En effet, sachant que le point de convergence se situe dans le plan focal image (respectivement à l'infini), il nous suffit d'intercepter ce plan une seule fois pour déterminer la position du point de convergence (respectivement, les rayons émergeant tous parallèles entre eux, il suffit d'un seul d'entre eux pour donner la direction collective).

Une fois ceci en main, notre but est de voir l'image finale sous le plus grand angle possible. En effet, s'agissant d'objet et d'image à l'infini, le grandissement ne nous aiderait pas beaucoup, et seul le grossissement nous permet une approche quantitative des choses.

Or les angles  $\alpha$  et  $\alpha'$  sous lesquels sont respectivement vus l'objet et l'image finale qu'en donne la lunette (nécessaires au calcul du grossissement  $G$ ) interviennent dans deux relations trigonométriques simples :

$$\tan \alpha = \frac{A_1 B_1}{f'_1} \quad \text{et} \quad \tan \alpha' = \frac{A_1 B_1}{f'_2}$$

Les objets observés en astronomie (ainsi que, la plupart du temps, leurs images) sont, nous l'avons dit, généralement vus sous de petits angles. En assimilant dans ce cadre les valeurs en radians des angles avec leurs tangentes, puis en éliminant  $A_1 B_1$  entre les deux égalités ci-dessus, il vient :

$$G = \frac{\alpha'}{\alpha} \approx \frac{f'_1}{f'_2}$$

Nous constatons donc que pour obtenir un grossissement optimal, il est bon de choisir :

- un objectif de focale  $f'_1$  aussi grande que possible ;
- un oculaire de focale  $f'_2$  aussi petite que possible.

## Halte aux idées reçues

Expliquer, pour chacune des affirmations suivantes, pourquoi elle est fausse :

1. Pour disperser une lumière polychromatique, il est nécessaire de faire passer celle-ci à travers un prisme.
2. Un rayon lumineux se propageant parallèlement à l'axe optique d'une lentille mince convergente converge, après franchissement de cette lentille, vers le foyer image de celle-ci.

3. Si un objet ou une image est positionné à plus de 5 m d'une lentille mince, on peut considérer qu'il se trouve à l'infini de celle-ci.

4. Une lentille mince convergente fait toujours converger la lumière qu'elle reçoit.

5. Le grossissement indique à quel point l'image d'un objet est plus grande que l'objet lui-même.

## Du Tac au Tac

(Exercices sans calculatrice)

1. Un rayon lumineux pénètre un prisme de verre, d'indice de réfraction  $n = 1,5$  et d'angle au sommet  $A = 60^\circ$ , avec un angle d'incidence  $i = 4,5^\circ$ .

Déterminer son angle d'incidence sur la deuxième face du prisme.

2. Un rayon lumineux se propage dans une fibre optique dont le cœur a pour indice de réfraction  $n = 1,50$ , et en formant à l'intérieur de celle-ci un angle  $\theta = 30^\circ$  avec son axe longitudinal.

Déterminer quelle longueur de cette fibre aura parcourue la lumière au bout d'une durée  $\Delta t = 1,00$  ms (on rappelle que  $c = 3,00 \cdot 10^8$  m.s<sup>-1</sup>, et que  $\sqrt{3} \approx 1,73$ ).

3. Un objet  $AB$  de 5,0 cm de haut est perpendiculaire à l'axe optique d'une lentille mince convergente  $L$  de centre

optique  $O$  et de vergence  $C = 3,0 \text{ } \delta$ . Il est positionné 50 cm avant cette lentille.

Déterminer la position et la taille de l'image  $A'B'$  de  $AB$  par  $L$ .

4. On considère deux lentilles minces convergentes  $L_1$  et  $L_2$ , dont les centres optiques sont distants de  $\overline{O_1O_2} = 40$  cm. La lentille  $L_1$  donne d'un objet  $AB$  perpendiculaire à l'axe optique une image  $A_1B_1$  positionnée 60 cm après  $L_1$ .

Déterminer la position de l'image finale  $A_2B_2$ , sachant que la lentille  $L_2$  a une vergence  $C_2 = 15 \text{ } \delta$ .

5. Un œil fixe un objet  $AB$  initialement positionné à 2,0 m de son cristallin.

Déterminer quantitativement comment doit être modifiée la vergence de ce cristallin, si  $AB$  est positionné deux fois plus proche de celui-ci.

## Vers la prépa

On positionne un objet lumineux  $AB$  à une distance  $D = 1,00$  m d'un écran. On considère une lentille mince convergente  $L$ , de longueur focale  $f'$  inconnue, et dont la position du centre optique est repérée par son abscisse  $x$ , l'origine étant fixée au niveau du pied  $A$  de l'objet. On déplace  $L$  entre l'objet et l'écran, et l'on constate la formation d'une

image  $A'B'$  nette, pour deux valeurs de  $x$  :  $x_1 = 20,0$  cm et  $x_2 = 80,0$  cm.

Montrer que la seule connaissance de  $D$  et  $d = x_2 - x_1$  permet de déterminer  $f'$ , selon une relation et à une condition que l'on précisera.

## Halte aux idées reçues

**1.** Le passage par un prisme est une condition suffisante, en ce sens qu'elle permet d'obtenir la dispersion de la lumière. Mais elle n'est pas nécessaire, puisqu'il est possible de remplacer le prisme par un autre objet.

Sans même changer radicalement de dispositif, le seul fait de changer de milieu suffit techniquement à disperser la lumière, tous les milieux (à l'exception du vide) étant plus ou moins dispersifs, c'est-à-dire présentant une indice de réfraction variable selon la fréquence de l'onde sinusoïdale qui les traverse.

L'indice de réfraction  $n$  d'un milieu est le rapport de la célérité des ondes lumineuses dans le vide, à celle des ondes lumineuses dans ce milieu :

$$n_X = \frac{c}{c_X} \quad \Leftrightarrow \quad c_X = \frac{c}{n_X}$$

On constate donc que la dépendance de l'indice de réfraction du milieu vis-à-vis de la fréquence de l'onde qui s'y propage est simplement une manifestation de la dépendance de la célérité de cette onde dans ce milieu, vis-à-vis de cette fréquence.

Cette définition éclaire notamment le fait que l'indice de réfraction soit une grandeur sans dimension, et le fait qu'il soit toujours supérieur à 1 (excepté dans le vide, où il vaut exactement 1 par définition).

On utilise cependant couramment le prisme parce que la présence de deux dioptries consécutifs, qui plus est inclinés l'un par rapport à l'autre, permet d'accentuer la dispersion de la lumière, et donc d'observer le phénomène avec plus de confort et de précision.

Il existe par ailleurs d'autres objets permettant de disperser la lumière, notamment les réseaux (grilles de fentes très fines et parallèles entre elles, provoquant des interférences multiples et constructives dans des directions dépendant de la longueur d'onde, et donc de la fréquence, de la radiation considérée).

**2.** L'erreur dans cette affirmation ne porte pas tant sur ce que fait le rayon se propageant au départ parallèlement à l'axe optique (à savoir qu'il est effectivement réorienté par la lentille, en direction de son foyer image), que sur le verbe *converger*.

En effet le concept de convergence, au même titre du reste que celui de divergence, est un concept collectif, qui suppose que plusieurs rayons concourent en un même point. On peut ainsi dire qu'un faisceau de rayons lumineux parallèles entre eux converge, après passage par une lentille, en un point situé dans son plan focal image (point que l'on détermine grâce au rayon

passant par le centre optique, qui n'est pas dévié et intercepte donc le plan focal image au niveau du point en question). On peut même ajouter que si les rayons de ce faisceau sont parallèles à l'axe optique de la lentille, alors ce point n'est autre que le foyer image de la lentille lui-même.

Mais parler de la convergence d'un rayon seul n'a aucun sens, ce concept nécessitant par définition plusieurs rayons.

**3.** L'idée d'un objet se trouvant à l'infini d'un autre est une expression signifiant que cet objet se trouve très éloigné de l'autre. Cette considération suppose donc une mesure de la distance les séparant, et une évaluation de l'importance de cette distance. Or en Physique une grandeur n'est jamais grande ou petite dans l'absolu, mais par comparaison à une autre grandeur de même nature.

Dans le cas de l'optique, le fait qu'un objet soit très éloigné d'une lentille, par exemple, signifie que les rayons lumineux issus d'un même point de cet objet sont pratiquement parallèles entre eux, c'est-à-dire que la valeur maximale de l'angle qu'ils forment les uns avec les autres soit très faible devant  $2\pi$  rad. Dans les faits, on peut montrer que ceci est vérifié dès lors que la distance entre l'objet et la lentille est très supérieure à la longueur focale de la lentille en question.

En général, les optiques utilisées en laboratoire ont des vergences comprises entre 2 et 20  $\delta$ , c'est-à-dire des longueurs focales comprises entre 5 et 50 cm environ. Dans ces conditions, 5 m est une longueur au moins dix fois supérieure à la longueur focale, et la condition est vérifiée en bonne approximation.

Mais si l'on s'intéresse à d'autres cas, comme les télescopes géants (VLT, ELT, GMT...), arborant fièrement des focales de plusieurs mètres voire dizaines de mètres, l'approximation n'est pas vérifiée. À l'autre bout de l'échelle, les optiques de microscope sont souvent de l'ordre de quelques mm de longueur focale, et il est alors inutile d'aller chercher 5 m pour considérer que l'objet ou l'image est à l'infini : 20 ou 30 cm peuvent très bien suffire.

**4.** L'usage a consacré le nom de *lentille convergente*, pour une lentille dont les bords sont moins épais que le centre. Ceci du fait qu'une telle lentille tend à faire converger un ensemble de rayons parallèles entre eux (donc issus d'un point source situé à l'infini de cette lentille), vers un point de son plan focal image.

Si ce point objet se rapproche, la lentille continue à faire converger les rayons issus d'un même point, vers un point situé après le foyer image de la lentille, tant que l'objet se trouve à gauche du plan focal objet de cette lentille.

Si l'objet se trouve au foyer objet de la lentille, cette convergence devient discutable, puisque les rayons émergent alors de la lentille parallèles entre eux, donnant à la fois une image réelle à l'infini (il suffit de positionner un écran à l'infini de la lentille, à droite, pour observer l'image) mais également une image virtuelle à l'infini (un observateur laissant ces rayons lumineux pénétrer son œil aura l'impression qu'ils sont tous issus d'un même point, situé à l'infini sur la gauche, puisqu'ils lui parviendront parallèles entre eux).

Mais si l'objet passe entre le plan focal objet de la lentille et cette dernière, les rayons lumineux divergent. Attention donc : le fait qu'une lentille mince réduise l'éventuelle divergence d'un faisceau de rayons lumineux n'implique pas forcément la convergence de ce dernier.

5. Nous avons vu dans le cours que la grandeur caractérisant le rapport de tailles entre une image et l'objet dont elle est conjuguée est le grandissement.

Le cas du grossissement est plus subtil en ce sens qu'il s'appuie sur les angles sous lesquels sont respectivement vus l'image et l'objet. Or ces angles dépendent certes de leurs tailles, mais également des distances auxquelles ils se trouvent du point d'observation : un petit objet vu de près peut très bien être vu sous un angle plus grand qu'un grand objet vu de loin.

Dans l'éventualité où un objet et son image sont situés à la même distance du point d'observation, alors effectivement les angles sous lesquels sont vus l'objet et son image sont à l'aune de leurs tailles et le grossissement peut, au moins aux petites angles, être assimilé au grandissement. Il s'agit cependant d'un cas particulier et cette affirmation est fautive dans le cas général.

## Du Tac au Tac

(Exercices sans calculatrice.)

1. L'application de la loi de Snell-Descartes pour la réfraction nous permet d'écrire, en prenant pour indice de l'air la valeur 1 et en notant  $r$  l'angle de réfraction dans le verre :

$$\sin(i) = n \sin(r)$$

Sans calculatrice, il est dans ce cas difficile d'évaluer la valeur de  $r$ , nécessaire à la suite de l'exercice. Mais dans le cas présent, l'angle  $i$  est de faible valeur (entendre que sa valeur en radians est faible devant 1). On sait alors que l'on peut approximer le sinus de cet angle par sa valeur en radians. Comme par ailleurs la lumière passe de l'air à un milieu plus réfringent, on sait que le rayon va se rapprocher de la normale au dioptre, donc que  $r < i$ , ce qui permet donc également d'assimiler  $\sin(r)$  à  $r$  en radians. Nous en déduisons :

$$i_{\text{rad}} = nr_{\text{rad}} \quad \Leftrightarrow \quad i = n r \text{ en degrés également}$$

en multipliant de part et d'autre de l'égalité par  $\frac{360}{2\pi} \text{ } ^\circ \cdot \text{rad}^{-1}$ .

L'application numérique donne alors, avec  $i = 4,5^\circ$  et  $n = 1,5$  :  $r = 3,0^\circ$ .

Reste ensuite à déterminer l'angle d'incidence  $i'$  du rayon sur l'autre face du prisme. Il suffit ici de s'appuyer sur le triangle formé par les deux faces du prisme et le rayon se propageant à l'intérieur de celui-ci. Avec  $r = 3,0^\circ$ , nous pouvons déterminer l'angle formé par le rayon avec la première face du prisme, comme complémentaire de  $r$  :  $90,0^\circ - r = 87,0^\circ$ .

Il ne reste plus alors qu'à utiliser le fait que la somme des angles d'un triangle est égale à  $180^\circ$ , et l'on trouve l'angle formé par le rayon lumineux, avec la deuxième face du prisme :  $180,0^\circ - 87,0^\circ - A = 33,0^\circ$ . L'angle d'incidence  $i'$  est le complémentaire de ce dernier, et il vient alors :  $i' = 90,0^\circ - 33,0^\circ = 57,0^\circ$ . Faites un schéma pour vous aider à retrouver ce résultat !

2. Un rayon lumineux se propageant dans une fibre optique avance par réflexions successives sur les interfaces cœur-gaine de cette fibre. Lorsqu'il avance sur une longueur  $L$  de cette fibre, s'il est parallèle à son axe, il parcourt une distance égale à  $L$ . Mais s'il est incliné par rapport à cet axe, il parcourt une distance  $L'$  plus importante, que l'on peut relier à  $L$  par la relation :

$$\cos \theta = \frac{L}{L'} \quad \Leftrightarrow \quad L' = \frac{L}{\cos \theta}$$

où  $\theta$  est l'angle formé par le rayon lumineux avec l'axe de la fibre. Faites un schéma pour vous aider !

Par ailleurs, si ce rayon voyage sur une durée  $\Delta t$  à la célérité  $c_g = \frac{c}{n}$ , il parcourt alors une distance  $L' = c_g \Delta t = \frac{c \Delta t}{n}$ .

En combinant les deux expressions précédentes, nous obtenons :

$$\frac{L}{\cos \theta} = \frac{c \Delta t}{n} \quad \Leftrightarrow \quad L = \frac{c \Delta t}{n} \cos \theta$$

Avec  $c = 3,00 \cdot 10^8 \text{ m} \cdot \text{s}^{-1}$ ,  $\Delta t = 1,0 \cdot 10^{-3} \text{ s}$ ,  $n = 1,50$  et  $\cos \theta = \frac{\sqrt{3}}{2}$ , nous obtenons  $L = 1,00 \cdot 10^5 \sqrt{3} \text{ m}$ , soit 173 km.

3. Pour trouver la position de l'image d'un objet par une lentille, connaissant la position de cet objet par rapport à cette lentille et la longueur focale de celle-ci, il nous suffit d'utiliser la relation de conjugaison de Descartes :

$$\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{f'} \quad \Leftrightarrow \quad \overline{OA'} = \frac{1}{\frac{1}{f'} + \frac{1}{\overline{OA}}}$$

Avec  $\overline{OA} = -0,50 \text{ m}$  d'où  $\frac{1}{\overline{OA}} = -2,0 \delta$  et  $\frac{1}{f'} = C = 3,0 \delta$ , nous trouvons ainsi  $\overline{OA'} = 1,0 \text{ m}$ .

Nous trouvons enfin les caractéristiques de l'image  $A'B'$  en utilisant l'égalité de Thalès présente dans la relation du grandissement :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OA'}}{\overline{OA}} = -2,0 \Rightarrow \overline{A'B'} = -2,0 \overline{AB} = -10 \text{ cm}$$

L'image est donc deux fois plus grande que l'objet, et renversée par rapport à celui-ci.

4. Nous pouvons ici de nouveau utiliser la relation de conjugaison de Descartes, au niveau de la lentille  $L_2$  :

$$\frac{1}{O_2A_2} - \frac{1}{O_2A_1} = \frac{1}{f'_2} \Leftrightarrow \overline{O_2A_2} = \frac{1}{\frac{1}{f'_2} + \frac{1}{O_2A_1}}$$

Le problème étant que  $\overline{O_2A_1}$  n'est pas directement donné. Il nous suffit alors d'utiliser une décomposition de Chasles :

$$\overline{O_2A_1} = \overline{O_2O_1} + \overline{O_1A_1} = +0,20 \text{ m}$$

avec  $\overline{O_2O_1} = -40 \text{ cm}$  et  $\overline{O_1A_1} = 60 \text{ cm}$  d'après l'énoncé. Nous constatons que cette valeur est positive :  $A_1B_1$  est donc un objet virtuel pour la lentille  $L_2$ .

Nous en déduisons  $\frac{1}{\overline{O_2A_1}} = 5,0 \delta$  et de là, avec la relation

de conjugaison de Descartes exprimée précédemment et  $\frac{1}{f'_2} = C_2 = 15 \delta$ ,  $\overline{O_2A_2} = 5,0 \text{ cm}$ .

5. Nous savons que pour un œil, la distance cristallin-rétine ( $\overline{OA'}$ ) est fixe, et que le cristallin adapte sa longueur focale  $f'$  de façon à ce que la rétine soit toujours conjuguée de l'objet observé, par le cristallin. Dans chacune de ces deux situations, la relation de conjugaison de Descartes nous permet donc d'écrire :

$$\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA_1}} = \frac{1}{f'_1} \quad \text{et} \quad \frac{1}{\overline{OA'}} - \frac{1}{\overline{OA_2}} = \frac{1}{f'_2}$$

En éliminant  $\frac{1}{\overline{OA'}}$  entre les deux égalités, puisque l'œil a toujours la même profondeur, nous obtenons ainsi :

$$\frac{1}{\overline{OA_1}} + \frac{1}{f'_1} = \frac{1}{\overline{OA_2}} + \frac{1}{f'_2} \Leftrightarrow \frac{1}{f'_1} - \frac{1}{f'_2} = \frac{1}{\overline{OA_2}} - \frac{1}{\overline{OA_1}}$$

Nous constatons ainsi que si l'objet se rapproche, c'est-à-dire si  $\overline{OA_2} > \overline{OA_1}$  (la valeur absolue diminue, mais s'agissant de valeurs négatives, la valeur algébrique augmente), alors

$\frac{1}{\overline{OA_2}} < \frac{1}{\overline{OA_1}}$ , d'où  $\frac{1}{f'_1} < \frac{1}{f'_2}$  et au final la vergence augmente :

le cristallin devient plus bombé, de sorte qu'il provoque une convergence plus prononcée des rayons.

Dans le cas présent, on nous dit que dans la seconde situation,  $AB$  est positionné deux fois plus proche du cristallin que dans

la première. En d'autres termes :  $\overline{OA_2} = \frac{\overline{OA_1}}{2}$ .

L'égalité précédente devient alors :

$$\frac{1}{f'_1} - \frac{1}{f'_2} = \frac{2}{\overline{OA_1}} - \frac{1}{\overline{OA_1}} \Leftrightarrow \frac{1}{f'_2} = \frac{1}{f'_1} - \frac{1}{\overline{OA_1}}$$

Avec  $\overline{OA_1} = -2,0 \text{ m}$  d'où  $\frac{1}{\overline{OA_1}} = -0,50 \delta$ , nous en déduisons que la vergence du cristallin doit augmenter de  $0,50 \delta$ .

## Vers la prépa

D'après la configuration matérielle utilisée, nous pouvons affirmer que :

$$\underbrace{\overline{AA'}}_D = \underbrace{\overline{AO}}_x + \overline{OA'} \Leftrightarrow \overline{OA'} = D - x$$

En injectant l'expression précédente dans la relation de conjugaison de Descartes, et avec  $\overline{OA} = -x$ , nous obtenons :

$$\frac{1}{D-x} + \frac{1}{x} = \frac{1}{f'} \Leftrightarrow x^2 - Dx + Df' = 0$$

Nous constatons donc que  $x = \overline{AO}$  (la longueur décrivant la position de la lentille par rapport à l'objet) est solution d'une équation du deuxième degré dont les coefficients dépendent uniquement de  $D$  et  $f'$ .

Nous savons qu'une telle équation possède deux racines réelles distinctes entre elles, si et seulement si son discriminant est strictement positif, soit en l'occurrence :

$$D^2 - 4Df' > 0 \Leftrightarrow D > 4f'$$

La condition évoquée par l'énoncé est donc que l'écran doit être éloigné de l'objet, d'une distance supérieure à 4 fois la longueur focale de la lentille, c'est-à-dire plus éloigné que dans une configuration de Silbermann.

Les solutions de l'équation précédente ont pour expressions respectives :

$$x_+ = \frac{D + \sqrt{D^2 - 4Df'}}{2} \quad \text{et} \quad x_- = \frac{D - \sqrt{D^2 - 4Df'}}{2}$$

En notant alors  $d = x_2 - x_1 = x_+ - x_- = 60,0 \text{ cm}$ , nous obtenons :

$$d = \sqrt{D^2 - 4Df'}$$

De l'équation précédente, nous déduisons :

$$d^2 = D^2 - 4Df' \Leftrightarrow f' = \frac{D^2 - d^2}{4D} = 16,0 \text{ cm}$$



### 10.1 Phénomènes, grandeurs et lois de base de l'électrocinétique

#### — Qu'est-ce que le courant électrique et comment peut-il circuler ?

On appelle **courant électrique** un mouvement ordonné de particules électriquement chargées. Ces particules véhiculent de l'énergie et il est donc possible, moyennant des systèmes adaptés, d'utiliser cette énergie à des fins diverses.

Pour pouvoir être le siège d'un courant électrique, un matériau doit donc disposer de particules électriquement chargées et libres de se déplacer, appelées **porteurs de charge libres**. Concrètement, on en trouve principalement dans deux types de matériaux :

- les métaux, où les porteurs de charge libres sont des électrons (statistiquement, quelques électrons par atomes sont libres de passer d'un atome à un autre) ;
- les solutions électrolytiques, où les porteurs de charge libres sont des ions solvatés (cf. conductimétrie).

Un matériau capable de conduire le courant électrique est dit **conducteur** ; un matériau ne le permettant pas est dit **isolant**.

On représente alors le sens de circulation du courant électrique dans un conducteur par une flèche apposée sur celui-ci :

Sens du courant  
→

#### — Qu'est-ce qu'un dipôle électrique ?

On appelle **dipôle électrique** tout objet possédant deux bornes (on dit également *pôles*) permettant de le connecter à d'autres dipôles électriques (ces connexions sont généralement réalisées au moyen de fils métalliques). L'association entre plusieurs dipôles électriques constitue alors un **circuit électrique**, dont le comportement dépendra de la nature des différents dipôles dont il est constitué, ainsi que de la façon dont ils sont agencés les uns par rapport aux autres.

Il existe une très grande variété de dipôles, mais il est possible dans un premier temps de les classer en deux catégories :

- les dipôles **générateurs** (également dits *actifs*), qui **imposent** la circulation d'un courant électrique ;
- les dipôles **récepteurs** (également dits *passifs*), qui **subissent** la circulation du courant électrique imposé par un éventuel générateur.

## Comment décrire l'architecture d'un circuit électrique ?

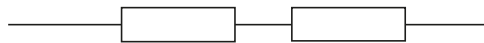
La très grande diversité des associations possibles entre dipôles nécessite un vocabulaire spécifique dédié à la description d'un circuit. Les éléments de ce vocabulaire sont motivés par les propriétés générales des circuits, qui seront détaillées plus loin. En particulier, on appelle :

- **nœud** d'un circuit, tout point de ce circuit où concourent au moins 3 fils ;
- **branche** d'un circuit, toute portion de ce circuit comprise entre 2 nœuds consécutifs (pas d'autre nœud entre ceux situés aux extrémités) ;
- **maille** d'un circuit, toute succession fermée de branches ne passant pas deux fois par la même branche.

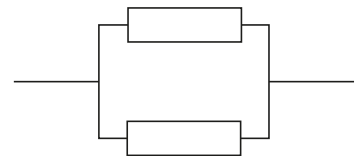
Il existe de multiples manières d'associer des dipôles entre eux. Nous en retiendrons deux ; deux dipôles sont ainsi dits associés :

- **en série** lorsqu'ils possèdent une unique connexion commune, et que cette connexion n'est pas un nœud du circuit ;
- **en parallèle** (on dit également en dérivation) lorsque leurs bornes sont connectées deux à deux.

Dipôles associés en série



Dipôles associés en parallèle



## Qu'est-ce qu'une tension électrique et comment la mesure-t-on ?

Toute particule électriquement chargée génère, en toute position  $M$  de l'espace, une grandeur appelée potentiel électrique, notée  $V(M)$  (USI : le volt, symbole V).

Dans le cadre des circuits électriques que nous traitons ici, nous considérerons que le potentiel électrique généré par une charge  $q$  :

- est d'autant plus élevé que  $q$  est plus grande ;
- ne varie jamais le long d'un fil, mais diffère en général entre les deux bornes d'un dipôle.

On appelle alors **tension électrique** aux bornes d'un dipôle, la différence de potentiel (souvent abrégée *ddp*) entre les bornes de ce dipôle. L'USI de tension électrique est évidemment la même que celle de potentiel électrique, soit le volt.

Si nous considérons alors un dipôle de bornes respectives  $A$  et  $B$ , nous constatons qu'en l'absence de précision il est possible de considérer deux tensions à ses bornes :

$$U_{AB} = V_A - V_B \quad \text{et} \quad U_{BA} = V_B - V_A$$

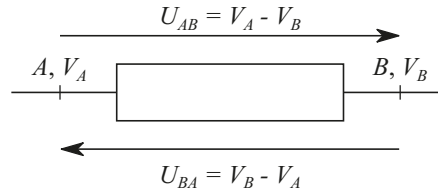
Nous notons que :

$$U_{BA} = -U_{AB}$$

Il est donc important de préciser avec laquelle de ces deux tensions nous choisissons de travailler, une confusion risquant d'introduire ou au contraire d'oublier un signe « - » dans les équations régissant le comportement du circuit, et donc de conduire à des résultats faux.

On dispose alors de deux options pour préciser rapidement la tension qui nous intéresse :

- préciser en indice de cette tension l'ordre dans lequel sont différenciés les potentiels ;
- représenter l'ordre dans lequel ils sont différenciés sur le schéma du circuit, à l'aide d'une flèche suivant la convention suivante :  $U_{\text{flèche}} = V_{\text{pointe}} - V_{\text{origine}}$

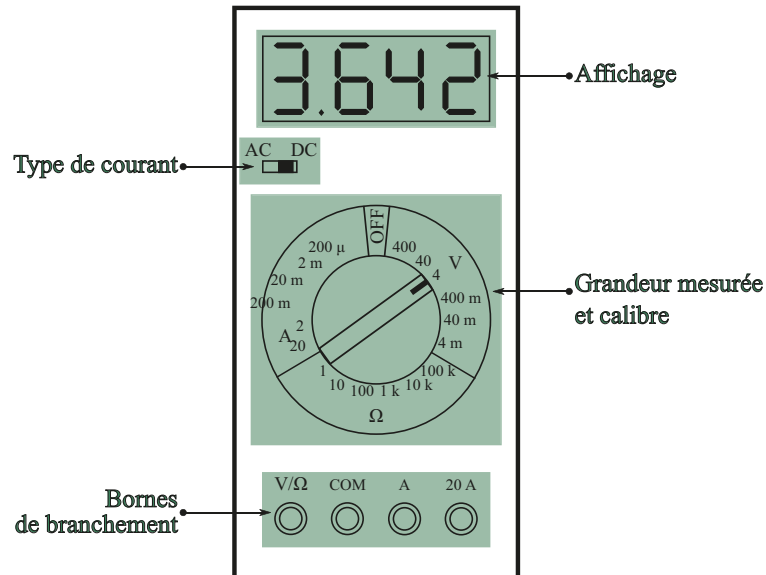


Nous verrons par la suite comment est effectué le choix de l'une ou l'autre de ces tensions, en fonction notamment de la nature du dipôle considéré (générateur ou récepteur).

**Remarque :** la flèche représente uniquement un ordre de différenciation, et en aucun cas un vecteur ; la tension électrique est bien une grandeur scalaire, non une grandeur vectorielle.

La mesure de la tension électrique se fait à l'aide d'un voltmètre. Cependant de nos jours, cet instrument est généralement combiné avec d'autres instruments au sein d'un **multimètre numérique**. Celui-ci comporte 4 zones essentielles, dont la forme et l'emplacement varient d'un modèle à l'autre et que vous devez toujours repérer avant d'effectuer votre mesure :

- un écran, où s'affichera la valeur mesurée ;
- une zone de choix entre courant continu (CC en français, DC en anglais) et courant alternatif (CA en français, AC en anglais) ;
- un cadran permettant de sélectionner la grandeur à mesurer ainsi que le calibre de mesure ;
- des prises de connexion.



**Remarque :** attention, les modalités de calcul de la valeur affichée dépendent du type de courant. Par exemple la moyenne temporelle (celle affichée en mode continu) d'une tension purement sinusoïdale est nulle, même si celle-ci oscille avec une grande amplitude. De même la grandeur efficace (celle affichée en mode alternatif) est nulle pour une tension continue, même si celle-ci présente

une valeur élevée. Il ne s'agit donc pas d'une précision anecdotique, mais qui conditionne la pertinence des résultats obtenus, voire la sécurité.

Dans le cas où l'on souhaite mesurer la tension électrique aux bornes d'un dipôle, on va offrir au voltmètre deux points de mesure (un à chaque borne), entre lesquels il mesurera la différence de potentiel qu'il affichera ensuite à l'écran.

La procédure à suivre pour effectuer une mesure est donc :

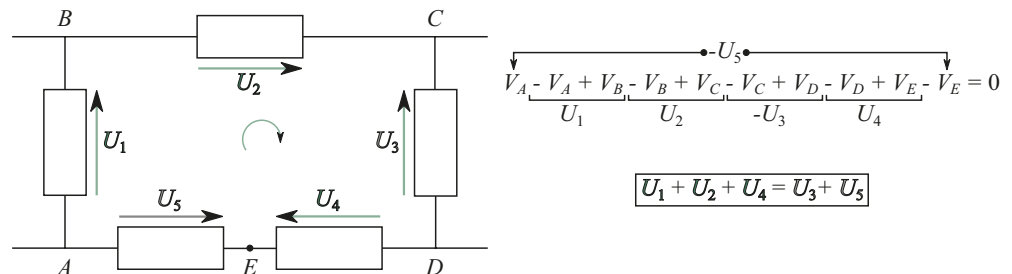
1. Sélectionner le type de courant (alternatif ou continu).
2. Sélectionner la grandeur à mesurer.
3. Placer le bouton sur le **plus grand calibre** disponible pour la grandeur mesurée (protection du matériel).
4. Sélectionner les bornes de branchement ; pour une tension, il s'agira des bornes **V** et **COM**.
5. Brancher le voltmètre **aux bornes** du dipôle aux bornes duquel on souhaite mesurer la tension.
6. Abaisser le calibre, jusqu'à la plus petite des valeurs proposées, qui reste supérieure à la valeur mesurée (optimisation de la précision). La tension affichée à l'écran correspond alors à la différence de potentiel entre les bornes V et COM :  $U_{\text{aff}} = V_V - V_{\text{COM}}$ .

**Remarque :** en électrocinétique, on tâche en général de privilégier la mesure de valeurs positives (il y a déjà bien assez d'occasion de se tromper de signe sans en rajouter). Donc si l'on sait dans quel sens la tension est positive (autrement dit de quel côté du dipôle le potentiel est le plus élevé), on a tout intérêt à brancher la borne V du voltmètre au niveau du potentiel le plus élevé.

Précisons que du point de vue du circuit, le voltmètre se comporte comme un interrupteur ouvert : il possède une résistance interne très élevée (de l'ordre de 10 MΩ) par rapport aux résistances usuellement rencontrées (au plus de 100 kΩ en général). Il peut donc être branché aux bornes d'un dipôle sans que ses points de branchements puissent être considérés comme des nœuds (puisque'ils n'occasionneront pratiquement aucune fuite de courant). On peut ainsi considérer que sa présence n'altère pas le comportement du circuit.

## — Quelle relation existe-t-il entre les tensions au sein d'une même maille ?

Une maille consistant en une succession fermée de branches, si on la parcourt on se retrouve nécessairement au même point, donc au même potentiel. Ce retour obligatoire à la valeur de départ a pour conséquence une loi, dite **loi des mailles**, et qui s'énonce ainsi : dans toute maille fermée d'un circuit, la somme des tensions dont les flèches représentatives sont orientées dans un sens, est égale à la somme des tensions dont les flèches sont orientées dans le sens opposé.



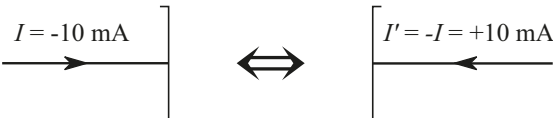
# — Qu'est-ce qu'une intensité électrique et comment la mesure-t-on ?

En courant continu, on appelle **intensité**  $I$  d'un courant électrique dans une branche de circuit, le rapport de la charge électrique  $Q$  transitant à travers cette branche en une durée  $\Delta t$ , à la durée en question :

$$I = \frac{Q}{\Delta t}$$

À la base son USI est le rapport des USI de charge électrique et de durée, soit donc du C.s<sup>-1</sup>. Toutefois on lui préfère en général l'ampère (symbole : A).

Notons que si les porteurs de charge libres sont des électrons, la charge électrique véhiculée est négative, et par suite l'intensité de ce courant d'électrons l'est également. Or on peut montrer que lorsqu'un courant d'intensité  $I$  circule dans un certain sens dans une branche de circuit, tout se passe comme si un courant d'intensité opposée circulait dans l'autre sens. En effet, une charge négative qui entre (respectivement qui sort) dans une branche, équivaut à une charge positive qui sort (respectivement qui entre).



Comme pour les tensions, nous voici donc libres de travailler avec l'un ou l'autre de ces deux courants. La réalité physique inviterait plutôt à considérer le courant d'électrons d'intensité négative, qui circule des zones de faible vers celles de fort potentiel (« du – vers le + », pour faire court). Ceci dit, encore une fois, on préfère privilégier les valeurs positives, et par convention on travaille toujours avec le **courant d'intensité positive**, qui circule, donc, des zones de fort vers celles de faible potentiel (ou encore « du + vers le – »).

La mesure de l'intensité du courant électrique se fait également au moyen d'un multimètre, cependant les modalités diffèrent sur plusieurs points :

| Multimètre utilisé en tant que... | Voltmètre                                              | Ampèremètre                                     |
|-----------------------------------|--------------------------------------------------------|-------------------------------------------------|
| Sélection du type courant         | Selon le type d'alimentation (continue ou alternative) |                                                 |
| Cadran                            | V                                                      | A                                               |
| Calibre initial                   | Le plus grand possible                                 |                                                 |
| Bornes de branchement             | V et COM                                               | A et COM                                        |
| Branchement                       | En parallèle aux bornes du dipôle                      | En série dans la branche du dipôle              |
| Calibre final                     | Le plus petit qui reste supérieur à la valeur mesurée  |                                                 |
| Valeur affichée                   | $U_{\text{aff}} = V_V - V_{\text{COM}}$                | $I_{\text{aff}} = I_{A \rightarrow \text{COM}}$ |
| Assimilable dans le circuit à...  | Un interrupteur ouvert                                 | Un fil                                          |

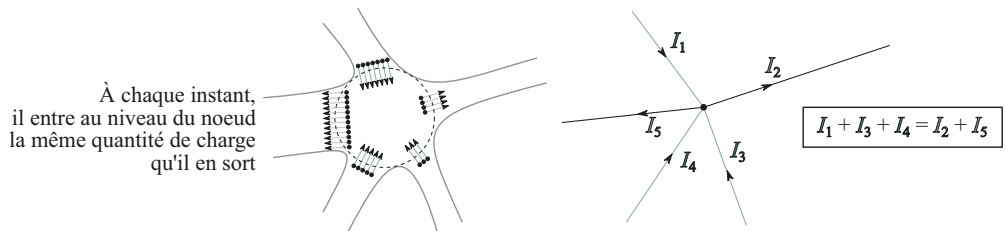
De la même façon qu'un voltmètre se comporte comme un interrupteur ouvert pour ne pas modifier l'intensité du courant qui circule à travers le dipôle, la présence de l'ampèremètre ne doit pas modifier la tension aux bornes du dipôle. Lorsque le multimètre est utilisé en tant qu'ampèremètre, il présente donc cette fois une résistance pratiquement nulle, qui le rend assimilable à un fil (doté d'un écran affichant l'intensité du courant qui le traverse de sa borne A vers sa borne COM, mais un fil quand même).

## — Quelle relation existe-t-il entre les intensités des courants concourant en un même nœud ?

La charge électrique étant par nature une grandeur conservative, on comprend que :

- tant que dans une portion de circuit le courant électrique ne possède qu'une seule voie d'entrée et une seule voie de sortie électrique, la charge transitant à travers cette portion est la même en chacun de ses points en une même durée  $\Delta t$ . En d'autres termes, l'intensité du courant électrique est la même en tout point d'une même branche ;
- à l'inverse, s'il existe un nœud dans une portion de circuit, la somme des charges se dirigeant vers ce nœud en une durée  $\Delta t$  sera nécessairement égale à la somme de celles s'éloignant de ce nœud pendant la même durée.

Ce dernier résultat a pour conséquence la **loi des nœuds**, qui s'énonce ainsi : en tout nœud d'un circuit, la somme des intensités des courants électriques orientés vers ce nœud est égale à la somme des intensités des courants électriques orientés à l'opposé de ce nœud.



## — Qu'est-ce qu'un régime transitoire ?

En courant continu comme en courant alternatif, un circuit tend toujours vers un régime de fonctionnement où il reproduira constamment le même comportement :

- tensions et intensités constantes dans le cas d'un fonctionnement en courant continu ;
- tensions et intensités périodiques à la fréquence imposée par le générateur dans le cas d'un fonctionnement en courant alternatif.

Lorsque l'une ou l'autre de ces deux situations est établie, on dit du circuit qu'il fonctionne en **régime permanent** et, encore une fois, il finit toujours par y aboutir. Cependant, si les conditions de fonctionnement du circuit changent (typiquement une variation de la tension imposée par le générateur, un interrupteur changeant d'état, etc.), les caractéristiques de ce régime permanent vont évoluer pour se fixer à d'autres valeurs. Et le temps que cette transition s'opère, le circuit va adopter un comportement particulier. On dit dans ce cas qu'il fonctionne en **régime transitoire**.

L'étude de ces régimes transitoires, que vous avez abordée en classe de terminale sur l'exemple du circuit RC, reviendra en force en CPGE sur ce circuit et sur d'autres, et il importe donc que vous vous y habituez.

## — Les définitions et lois vues précédemment sont-elles encore valables en régime transitoire ?

Les deux lois vues ci-dessus, fondées sur la conservativité de la charge électrique pour l'une (loi des nœuds) et sur celle du potentiel électrique pour l'autre (loi des mailles), doivent être vérifiées à chaque instant, et peuvent donc jusqu'à nouvel ordre être considérées comme systématiquement valables.

La définition de la tension électrique à partir du potentiel fonctionne également.

En revanche la définition de l'intensité du courant électrique, elle, fait explicitement intervenir la variable temps. La question de savoir ce qu'elle devient dans le cas transitoire se pose donc sérieusement. Il est en fait nécessaire de faire ce que nous faisons toujours en physique lorsqu'une grandeur évolue en fonction d'une certaine variable : décomposer la variation globale en une série de variations élémentaires, et passer de l'étude d'une valeur moyenne à celle d'une valeur instantanée.

L'intensité instantanée véhiculée par un courant électrique sur une durée  $dt$  se définit alors comme :

$$i(t) = \frac{\delta q}{dt}$$

avec  $\delta q$ , charge véhiculée par ce courant dans la branche de circuit considérée, pendant la durée  $dt$ .

De manière générale, lorsque l'on étudie des grandeurs évoluant au cours du temps (aussi bien les variations d'un courant alternatif que celles résultant d'un régime transitoire), on explicite éventuellement cette dépendance, mais surtout on remplace les notations majuscules par des notations minuscules :

| Valeurs constantes | Valeurs évoluant au cours du temps |
|--------------------|------------------------------------|
| $U, I$             | $u(t), i(t)$                       |

## — Quel est le principe général de traitement d'un problème d'électrocinétique ?

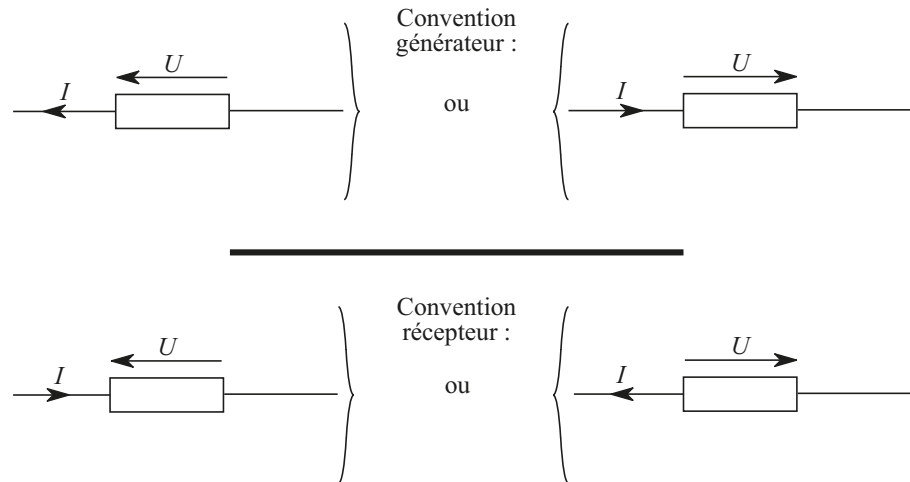
La loi des nœuds et la loi des mailles, respectivement appelées première et deuxième **lois de Kirchhoff**, constituent les deux piliers qui permettent d'établir les équations décrivant le comportement d'un circuit électrique. S'appuyant respectivement sur la conservativité de la charge électrique et celle du potentiel électrique, elles expriment comment sont liées entre elles les intensités des courants d'une part, et les tensions d'autre part. En leur adjoignant des relations liant les tensions aux intensités (ce que nous allons voir sous peu), on obtient un système d'équations dont les inconnues sont justement ces tensions et ces intensités, et dont la résolution permettra de connaître le détail de chacune de celles-ci.

# 10.2 Les dipôles et leurs caractéristiques

## — Quels sont les différents types de dipôles ?

Nous avons déjà évoqué au début de ce chapitre l'existence de dipôles générateurs et de dipôles récepteurs, cependant avant d'aller plus loin il est nécessaire de préciser quelques points.

En particulier, les générateurs imposant une différence de potentiel, ils vont forcer la circulation d'un courant d'intensité positive dans la partie du circuit qui leur est extérieure, de leur borne de fort potentiel vers celle de faible potentiel. Il est alors intéressant de noter que les flèches représentatives de la tension positive aux bornes d'un générateur, et du courant d'intensité positive le traversant, sont forcément orientées dans le même sens.



Les récepteurs, pour leur part, subissent le passage du courant électrique et développent sous cet effet une différence de potentiel entre leurs bornes : celle par laquelle arrivent les électrons (donc le courant d'intensité négative) et par laquelle part le courant d'intensité positive développera donc un potentiel faible. Réciproquement, la borne par laquelle partent les électrons (et par laquelle arrive, donc, le courant d'intensité positive) développera un potentiel élevé. Il s'ensuit que les flèches représentatives de la tension positive aux bornes d'un dipôle récepteur et du courant d'intensité positive le traversant, sont orientées à contre-sens l'une de l'autre.

Comme dit précédemment, on préfère travailler autant que possible sur des valeurs positives, ce qui fait que l'on orientera systématiquement chaque dipôle selon la convention correspondant à sa nature.

**Remarque :** on peut juger cette convention inutile, et choisir de travailler sur des dipôles orientés selon la convention contraire à leur nature. Il importe cependant de bien comprendre les implications d'un tel choix, puisqu'outre prendre le risque d'indisposer votre correcteur/trice, les relations caractéristiques que nous allons voir devront être modifiées pour rester cohérentes malgré le changement de signe d'une des grandeurs qu'elles engagent, et pas de l'autre. Ainsi par exemple la loi d'Ohm pour un conducteur ohmique orienté en convention générateur ne s'écrit-elle plus  $u_R = R \times i_R$ , mais bien  $u_R = -R \times i_R$ .

## — Qu'est-ce que la caractéristique d'un dipôle ?

Ainsi que nous l'avons dit plus haut, les lois de Kirchhoff permettent d'établir une série d'équation liant entre elles les tensions d'une part, et les intensités des courants d'autre part. Cependant pour pouvoir résoudre ce système d'équations, il est nécessaire d'y intégrer les relations entre tensions aux bornes des dipôles, et intensités des courants qui les traversent respectivement.

On appelle ainsi **caractéristique** d'un dipôle électrique, la relation liant la tension aux bornes d'un dipôle à l'intensité du courant qui le traverse. Lorsque cette relation peut s'exprimer sous une forme analytique simple, on peut y associer une représentation graphique également qualifiée de caractéristique du dipôle.

Pour certains dipôles (en tête desquels le condensateur), la caractéristique ne met pas en relation la tension et l'intensité, mais la dérivée de l'une ou de l'autre, avec l'autre ou l'une. Dans ce cas la mise en équation du circuit aboutira à une équation différentielle, mais la démarche reste la même : résolution de l'équation (fût-elle différentielle), dont les solutions seront les fonctions du temps recherchées.

## — Qu'est-ce qu'un générateur idéal de tension et quelle est sa caractéristique ?

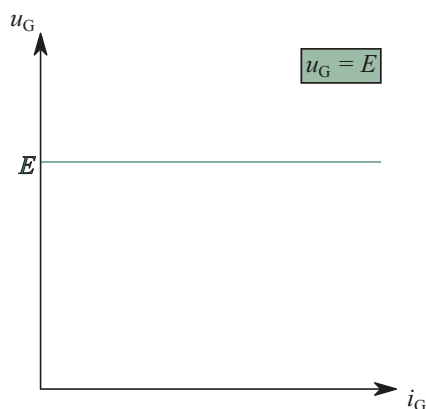
On appelle **générateur de tension** un générateur dont le paramètre fixe est la tension qu'il impose à ses bornes (par opposition au générateur de courant que vous découvrirez plus tard et qui, lui, fixe l'intensité du courant qu'il débite). Les générateurs de tension sont ceux que nous côtoyons le plus souvent au quotidien : piles, batteries, accumulateurs, etc. affichent et maintiennent une valeur de tension électrique (1,5 V, 4,5 V, 9,0 V, etc.).

Un générateur **idéal** est alors un modèle de générateur capable de maintenir sa tension quelle que soit l'intensité du courant qu'il doit débiter.

Sa caractéristique est des plus simples : la tension à ses bornes étant constante, elle est indépendante de l'intensité du courant qui le traverse, et se résume donc, pour un générateur dont la tension à vide est  $E$ , à :

$$u_G = E$$

On peut donc lui associer la représentation graphique suivante :



## — Qu'est-ce qu'un conducteur ohmique et quelle est sa caractéristique ?

Il existe une catégorie de dipôles récepteurs présentant la particularité, lorsqu'ils sont soumis à une tension électrique, de laisser passer un courant électrique dont l'intensité est directement proportionnelle à la tension en question. On appelle alors conductance  $G$  le coefficient de proportionnalité tel que :

$$i_R = G \times u_R$$

L'USI de conductance serait donc l'ampère (d'intensité débitée) par volt (de tension imposée), soit  $A \cdot V^{-1}$ , auquel on préfère généralement le siemens (symbole S, majuscule, attention à ne pas confondre avec le « s » minuscule de la seconde).

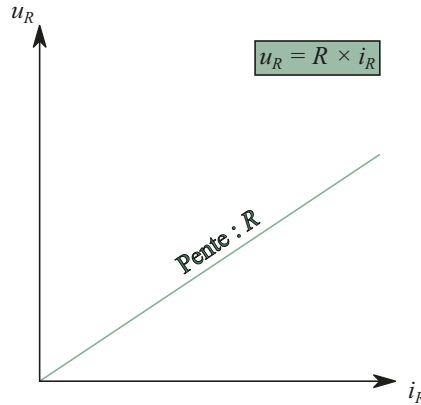
Cette conductance est donc d'autant plus grande que l'intensité du courant circulant l'est aussi pour une tension donnée. On lui préfère souvent la résistance  $R$ , définie par la relation inverse connue sous le nom de **loi d'Ohm** :

$$i_R = G \times u_R \Leftrightarrow u_R = \frac{1}{G} \times i_R \Leftrightarrow u_R = R \times i_R$$

avec  $R = \frac{1}{G}$  qui varie en raison inverse de la conductance : un conducteur ohmique qui conduit bien le courant électrique (conductance élevée) résiste peu à son passage (résistance faible) et réciproquement.

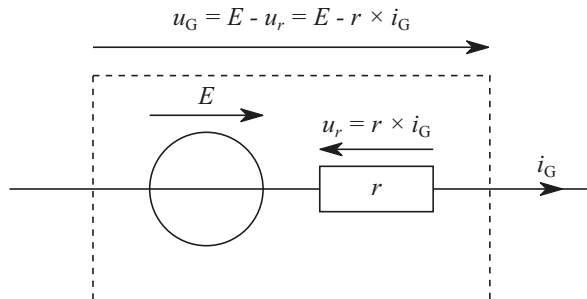
L'USI de résistance serait donc le volt (de tension aux bornes du conducteur ohmique) par ampère (d'intensité du courant le traversant), soit le  $\text{V} \cdot \text{A}^{-1}$ . On pourrait encore proposer le  $\text{S}^{-1}$ , mais dans les faits comme vous le savez c'est l'ohm (symbole  $\Omega$ ) qui est le plus couramment utilisé.

La caractéristique d'un conducteur ohmique est donc la loi d'Ohm, à laquelle nous pouvons associer la représentation graphique suivante :



## — Qu'est-ce qu'un générateur réel de tension et quelle est sa caractéristique ?

Dans les faits, la tension aux bornes d'un générateur autonome tend à baisser lorsque l'intensité du courant qu'il débite augmente. On constate expérimentalement que cette variation est souvent bien décrite par une fonction affine de coefficient directeur négatif. Or cette relation est celle que l'on obtiendrait si le générateur était constitué d'un générateur idéal de tension, en série avec un conducteur ohmique. On modélise donc couramment les générateurs de tension réels de la manière suivante :



On constate ainsi que la caractéristique du générateur idéal est effectivement une relation affine, de pente égale à  $-r$  et d'ordonnée à l'origine  $E$  :

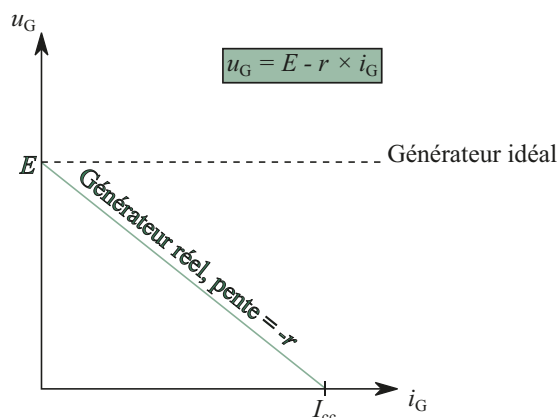
$$u_G = -r \times i_G + E$$

Il est alors intéressant de noter qu'un tel générateur ne peut débiter une intensité supérieure à une valeur limite, appelée **intensité de court-circuit**  $I_{cc}$ . Celle-ci correspond à la situation où le générateur est court-circuité, c'est-à-dire où ses deux bornes sont directement reliées l'une à l'autre par un fil conducteur, entraînant une tension nulle aux bornes du générateur réel :

$$u_G = 0 \Leftrightarrow -r \times i_G + E \Leftrightarrow i_G = I_{cc} = \frac{E}{r}$$

Pour des valeurs usuelles (typiquement  $E = 10 \text{ V}$  et  $r = 10 \Omega$ ), l'intensité ne pourra donc pas excéder 1A.

On peut alors associer à cette relation caractéristique la représentation graphique suivante :



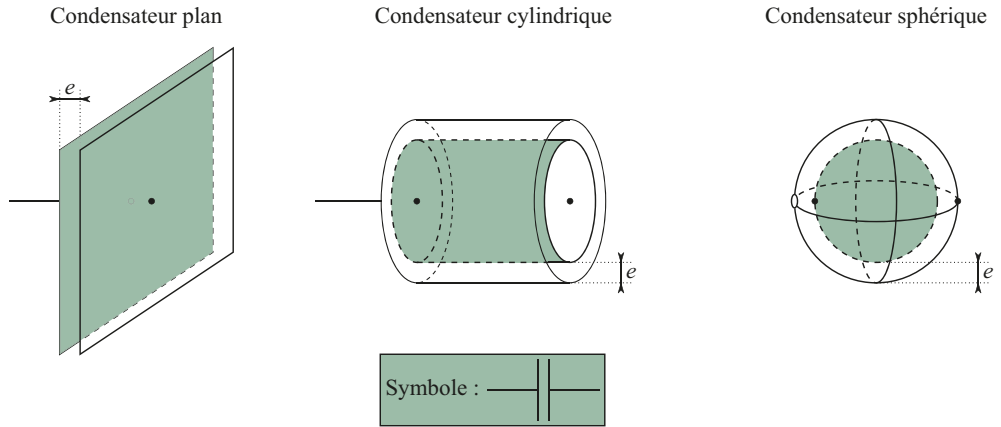
Nous constatons donc une différence significative entre les comportements d'un générateur réel et d'un générateur idéal. Cependant, tant que l'intensité réclamée par les conditions d'utilisation d'un générateur réel reste très inférieure à  $I_{cc}$ , la valeur de la tension aux bornes du générateur réel restera voisine de celle qu'imposerait le seul générateur idéal :

$$i_G \ll I_{cc} = \frac{E}{r} \Leftrightarrow E \gg r \times i_G \Leftrightarrow u_G = E - r \times i_G \simeq E$$

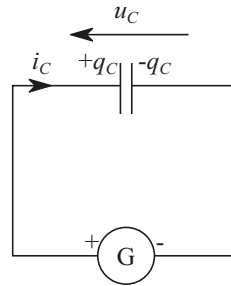
Dans ces conditions, il est donc licite d'approximer le générateur réel par un générateur idéal.

## — Qu'est-ce qu'un condensateur et quelle est sa caractéristique ?

Un condensateur est un dipôle constitué de deux surfaces conductrices en regard l'une de l'autre, séparées d'une distance  $e$  faible devant leurs dimensions. La forme d'un condensateur peut varier ; typiquement on rencontre des condensateurs dont les surfaces conductrices sont des plans parallèles entre eux, des cylindres coaxiaux ou des sphères concentriques. L'espace entre ces plaques est occupé par un isolant.



Ainsi vu, un condensateur semble s'apparenter à un interrupteur ouvert, ce qui offre des perspectives d'application assez limitées. On peut en réalité montrer que si les surfaces en regard sont soumises à une différence de potentiel, alors elles vont accumuler des charges électriques égales en valeur absolue, et de signes opposés, et que ces charges seront proportionnelles à la différence de potentiel imposée.



On appelle alors **capacité**  $C$  d'un condensateur, la constante de proportionnalité telle que :

$$q_C = C \times u_C$$

où  $q_C$ , donc, est la charge portée par l'armature du côté de laquelle arrive le courant.

L'USI de  $C$  serait donc le  $C.V^{-1}$ , auquel on préfère le farad (symbole F).

La relation vue plus haut ne fait pas explicitement apparaître l'intensité d'un éventuel courant dans la branche contenant le condensateur. On voit cependant rapidement que si un courant circule dans cette branche, les charges portées par les armatures vont varier.

Concrètement, considérons **l'armature du condensateur par laquelle arrive le courant**. En une durée  $dt$ , le courant d'intensité  $i_C$  véhicule une charge :

$$\delta q = i_C \times dt$$

Puisque nous nous intéressons à l'armature par laquelle arrive le courant, cette charge (algébrique) va venir s'ajouter à celle déjà présente sur cette armature. Pendant la durée  $dt$ , la charge portée par cette armature passe donc de  $q_C(t)$  à :

$$q_C(t + dt) = q_C(t) + \delta q \Leftrightarrow dq_C = i_C \times dt \Leftrightarrow i_C = \frac{dq_C}{dt}$$

En injectant alors dans cette égalité l'expression  $q_C = C \times u_C$ , nous obtenons la relation caractéristique du condensateur :

$$i_C = \frac{d(C \times u_C)}{dt} \Leftrightarrow i_C = C \times \frac{du_C}{dt}$$

si la capacité est constante, autrement dit que la géométrie du condensateur ne varie pas au cours du temps.

La relation caractéristique ne met donc pas aux prises directement l'intensité du courant dans la branche du condensateur avec la tension aux bornes de ce dernier, mais avec **la dérivée de cette tension**. Lui associer une représentation graphique n'est plus aussi simple, mais la mise en équation, elle, est toujours possible.

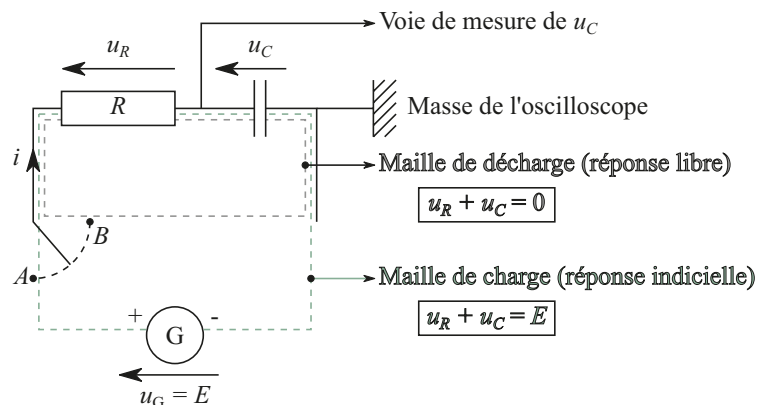
## 10.3 Exemple d'étude d'un régime transitoire : cas du circuit RC

### — Qu'appelle-t-on circuit RC ?

On appelle circuit RC l'association (dans le cas présent, en série) d'un conducteur ohmique (dont nous noterons ci-après la résistance  $R$ ) avec un condensateur (capacité notée  $C$ ). Dans le cadre qui nous occupe, ce circuit subira deux types de changements :

- **réponse à un échelon de tension** (également dite *réponse indicielle*) : partant d'un condensateur initialement déchargé ( $q_{C(t=0)} = 0$ ), et d'un courant d'intensité nulle ( $i_{(t=0)} = 0$ ), donc d'une tension globale initialement nulle aux bornes de l'ensemble RC, on impose soudainement une tension non nulle. Le condensateur va donc se charger, mais le courant à l'origine de cette charge devra passer par le conducteur ohmique ;
- **réponse libre** : c'est la situation inverse, où partant d'un condensateur chargé, on boucle l'association condensateur/conducteur ohmique sur un fil. Le condensateur va donc cette fois se décharger à travers le conducteur ohmique.

Pour pouvoir étudier alternativement ces deux réponses, on utilise en général le montage suivant :



Les orientations ont été choisies sur la base de la tension positive aux bornes du générateur, qui orienté en convention générateur nous donne le sens du courant. En prenant alors les dipôles récepteurs en convention récepteur, nous obtenons le sens des flèches de tension aux bornes de ceux-ci.

Quelles sont les équations qui régissent le fonctionnement d'un circuit RC ?

Le circuit ne comportant aucun nœud (la maille de travail varie selon la position de l'interrupteur, mais il n'en existe qu'une seule à chaque fois), la loi des nœuds ne nous est d'aucun secours. Tout au mieux nous permet-elle d'affirmer une évidence : le courant circulant à travers le conducteur ohmique aura la même intensité que celui intervenant dans la (dé)charge du condensateur :

i\_R = i\_C = i

La loi des mailles, elle, nous donne les deux équations figurant sur le schéma précédent :

| Réponse indicielle | Réponse libre   |
|--------------------|-----------------|
| $u_R + u_C = E$    | $u_R + u_C = 0$ |

En détaillant les expressions caractéristiques du conducteur ohmique et du condensateur, il vient alors :

{ u\_R = R x i\_R ; i\_C = C x du\_C/dt } => u\_R = R x C x du\_C/dt

puisque i\_R = i\_C = i. En substituant cette expression dans les équations fournies par la loi des mailles, il vient respectivement :

RC du\_C/dt + u\_C = E et RC du\_C/dt + u\_C = 0

On note que le produit R x C est homogène à une durée. La chose se voit d'ores et déjà en observant l'équation, et peut facilement être corroborée par l'analyse dimensionnelle :

[RC] = [R] x [C] = [U]/[I] x [Q]/[U] = [Q]/[I] = [Q] x ([Q]/[T])^-1 = [T]

On pose alors pour alléger l'écriture tau = RC et, en divisant à gauche et à droite par tau dans les équations précédentes, il vient :

| Réponse indicielle              | Réponse libre               |
|---------------------------------|-----------------------------|
| $du_C/dt + 1/tau x u_C = E/tau$ | $du_C/dt + 1/tau x u_C = 0$ |

Nous trouvons donc des équations différentielles linéaires à coefficients constants d'ordre 1, avec second membre constant dans le premier cas, et sans second membre dans le second.

## — Quelle est la dynamique de fonctionnement d'un circuit RC en réponse indicielle ?

Nous savons que les solutions de la première équation sont de la forme :

$$u_C(t) = Ke^{-\frac{t}{\tau}} + E$$

où  $K$  est une constante d'intégration fixée par une condition particulière.

Dans le cas présent, on peut invoquer la condition initiale portant sur la tension  $u_C$  aux bornes du condensateur :  $u_{C(t=0)} = 0$  (condensateur initialement déchargé). Il convient cependant d'être prudent : le changement de structure imposé au circuit (bascule de l'interrupteur de la position  $A$  à la position  $B$ ) est susceptible d'entraîner des variations brutales des différentes grandeurs dont il est le siège. Or la condition initiale évoquée ci-dessus est établie **avant la bascule** de l'interrupteur, et il serait cavalier d'affirmer qu'elle possède forcément la même valeur juste après.

Nous verrons cependant dans le chapitre portant sur les aspects énergétiques de l'électrocinétique que **la tension aux bornes d'un condensateur est toujours une fonction continue du temps**.

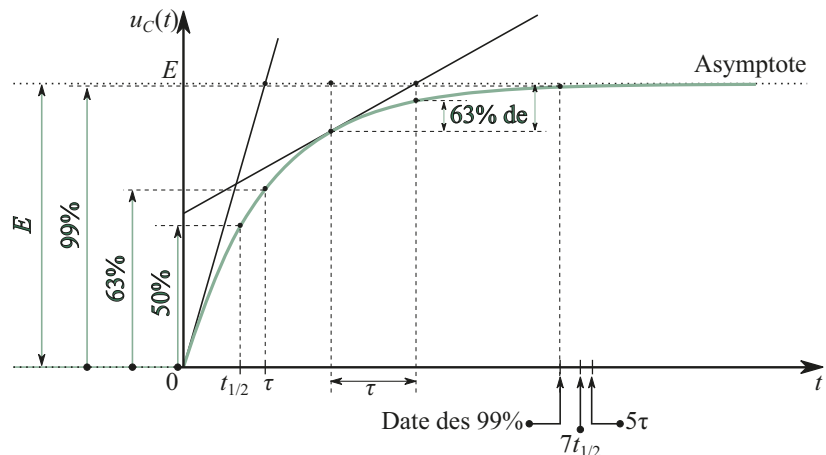
Dans ce cas, nous pouvons donc écrire :

$$u_{C(t=0^+)} = u_{C(t=0^-)} \Leftrightarrow \underbrace{Ke^{-\frac{0}{\tau}} + E}_{\text{Expression valable après fermeture}} = \underbrace{0}_{\text{Valeur avant fermeture}} \Rightarrow K = -E$$

Connaissant la valeur de la constante  $K$ , nous en déduisons l'expression complète de  $u_C(t)$  :

$$u_C(t) = -Ee^{-\frac{t}{\tau}} + E \Leftrightarrow u_C(t) = E \left( 1 - e^{-\frac{t}{\tau}} \right)$$

L'allure générale est alors la suivante :



Nous y retrouvons les propriétés usuelles des exponentielles :

- asymptote horizontale pour  $t \rightarrow \infty$  ;
- durée de demi-charge indépendante de la valeur initiale de la tension, ayant pour effet que toute attente d'une même durée entraînera la couverture d'une même fraction (63%) de l'intervalle restant à couvrir jusqu'à l'asymptote ;
- lecture de  $\tau$  comme écart temporel entre une date quelconque, et l'intersection de la tangente à la courbe à la date en question, avec l'asymptote ;
- durée de demi-charge liée au paramètre  $\tau$  par la relation :  $t_{1/2} = \tau \times \ln(2)$  ;
- pourcentage de couverture de l'intervalle entre valeur initiale et valeur asymptotique :
  - indépendant de la valeur initiale,
  - 50 % au bout d'une durée  $t_{1/2}$  (par définition),
  - 63 % au bout d'une durée  $\tau$ ,
  - > 99 % au bout d'une durée  $5\tau$  (reste 0,7 %) ou  $7t_{1/2}$  (reste 0,8 %).

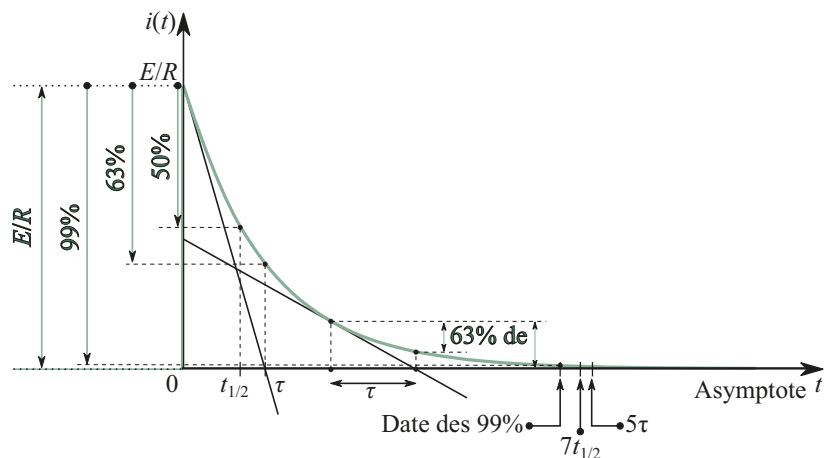
Notons que s'agissant d'une exponentielle, la valeur asymptotique n'est à strictement parler jamais atteinte. Cependant dès  $7\tau$  elle est atteinte à plus de 99 %. Or pour des valeurs usuelles de  $R$  et  $C$  (typiquement,  $R = 1 \text{ k}\Omega$  et  $C = 1 \text{ }\mu\text{F}$ ),  $\tau$  est en général de l'ordre de la milliseconde. Nous pouvons donc considérer qu'au bout de quelques secondes,  $u_C = E$  (et par suite que l'intensité, dérivant de  $u_C$  désormais constante, sera nulle). Nous observons ici, comme annoncé, un transitoire de durée très courte.

Par ailleurs, disposant désormais d'une forme explicite de  $u_C(t)$ , nous pouvons déduire l'expression de l'intensité dans le circuit de deux manières :

- en utilisant la loi des mailles :  $R \times i + u_C = E \Leftrightarrow i = \frac{E - u_C}{R} = \frac{E}{R} e^{-\frac{t}{\tau}}$
- en utilisant l'expression caractéristique du condensateur :

$$i = C \times \frac{du_C}{dt} = C \times E \underbrace{\frac{d}{dt} \left( 1 - e^{-\frac{t}{\tau}} \right)}_{+\frac{1}{\tau} \times e^{-\frac{t}{\tau}}} \Leftrightarrow i = \frac{C}{\tau} \times E \times e^{-\frac{t}{\tau}} = \frac{E}{R} e^{-\frac{t}{\tau}}$$

puisque  $\tau = RC$ . Nous trouvons bien la même expression dans les deux cas.



**Remarque :** l'intensité dans une branche de circuit contenant un condensateur, contrairement à la tension aux bornes de ce dernier, subit une importante discontinuité à la date de bascule, passant brutalement de 0 à  $\frac{E}{R}$ . Attention donc : toutes les grandeurs ne sont pas nécessairement continues au cours du temps, et nous ne sommes pas libres de faire valoir n'importe quelle condition initiale.

## — Quelle est la dynamique de fonctionnement d'un circuit RC en réponse libre ?

Nous savons que les solutions de la seconde équation vue plus haut sont de la forme :

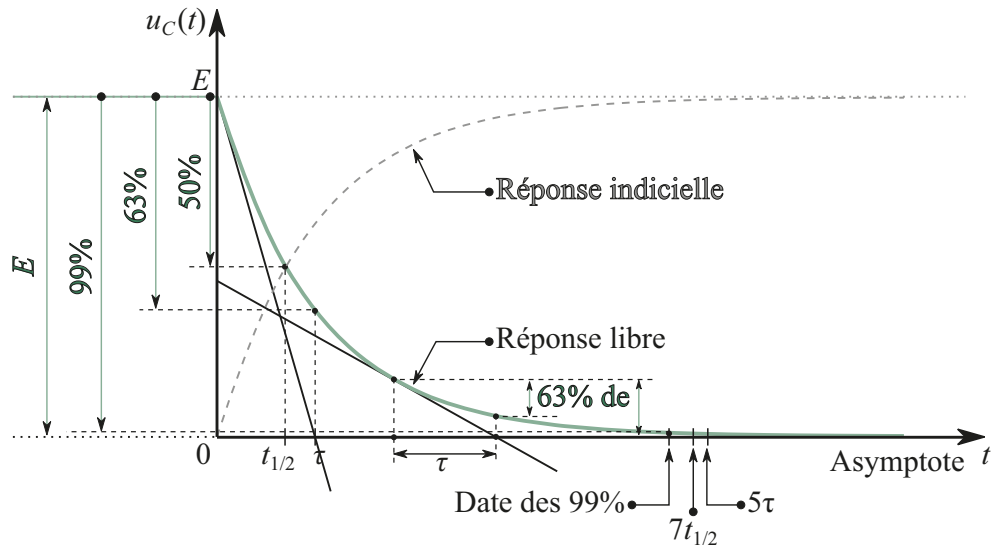
$$u_C(t) = K e^{-\frac{t}{\tau}}$$

où  $K$  est une constante d'intégration fixée par une condition particulière. Nous avons vu qu'à la suite de la phase de charge, la tension aux bornes du condensateur était égale à  $E$  (et accessoirement que l'intensité était nulle). La condition initiale portant sur la tension  $u_C$  est donc cette fois  $u_{C(t=0)} = E$ . En utilisant à nouveau la continuité de la tension aux bornes d'un condensateur, il vient cette fois :

$$u_{C(t=0^+)} = u_{C(t=0^-)} \Leftrightarrow \underbrace{K e^{-\frac{0}{\tau}}}_{\substack{\text{Expression} \\ \text{valable} \\ \text{après} \\ \text{fermeture}}} = \underbrace{E}_{\substack{\text{Valeur} \\ \text{avant} \\ \text{fermeture}}} \Leftrightarrow K = E$$

Connaissant la valeur de la constante  $K$ , nous en déduisons l'expression complète de  $u_C(t)$  :

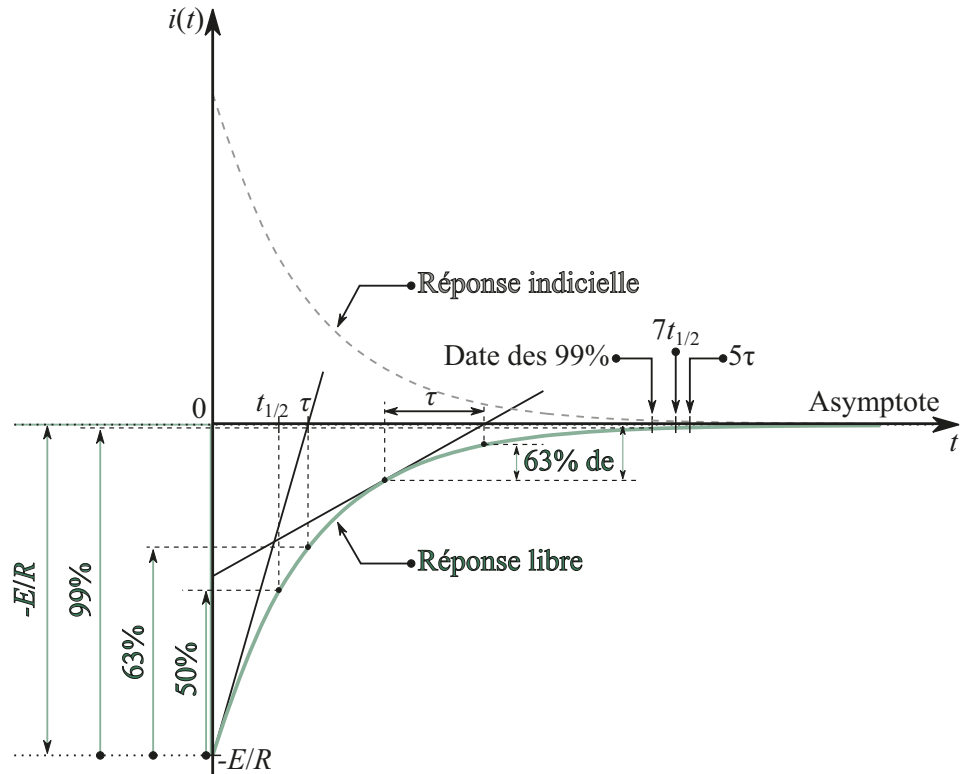
$$u_C(t) = E e^{-\frac{t}{\tau}}$$



L'expression de l'intensité dans le circuit se trouve comme dans le cas précédent, et donne cette fois :

$$i = -\frac{E}{R} e\left(-\frac{t}{\tau}\right)$$

L'allure de la courbe représentative associée donne :



Nous observons à nouveau une discontinuité à l'origine, mais il y a plus remarquable : mathématiquement, l'intensité du courant considéré est nécessairement négative, point qui contrevient à l'habitude évoquée plus haut de travailler avec des grandeurs positives.

Il faut cependant comprendre que dans cette situation le statut du condensateur est équivoque. En effet, en l'absence de générateur, il prend le départ avec une tension à ses bornes, contre laquelle le conducteur ohmique ne peut aller puisque lui-même ne développe une différence de potentiel à ses bornes que s'il est traversé par un courant. Le condensateur se comporte donc effectivement comme un générateur. Cependant comme nous l'avons orienté en convention récepteur, ce comportement entraîne forcément une différence de signe entre la tension et l'intensité du courant, associées aux flèches respectivement choisies pour les représenter.

## Halte aux idées reçues

Expliquer, pour chacune des affirmations suivantes, pourquoi elle est fausse :

1. Dans un circuit électrique, la flèche apposée sur un fil représente le sens de l'intensité électrique.
2. Il n'existe que deux façons d'associer deux dipôles entre eux : en série et en dérivation.
3. Dans un circuit comprenant un générateur électrique, le courant circule toujours du pôle positif du générateur, vers le pôle négatif.

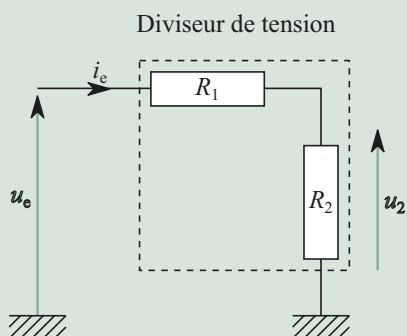
4. L'intensité instantanée est la dérivée de la charge électrique par rapport au temps.

5. Pour étudier le régime transitoire d'un circuit RC à l'aide d'un multimètre, il est nécessaire de placer celui-ci en mode alternatif.

## Du Tac au Tac

(Exercices sans calculatrice)

1. Montrer que l'association en série ou en dérivation de plusieurs conducteurs ohmiques équivaut à un conducteur ohmique unique dont on précisera dans chaque cas la résistance. Reprendre la question avec



Exprimer :

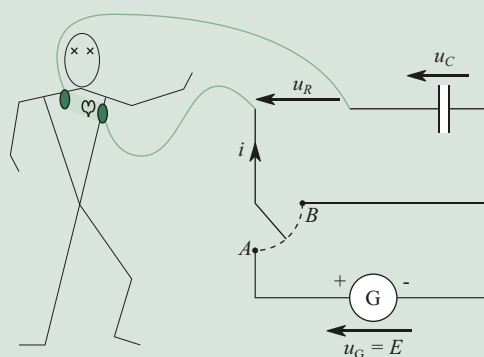
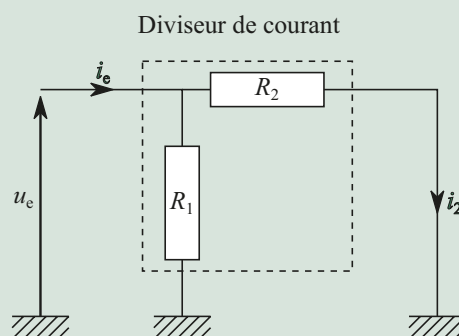
- la tension de sortie  $u_2$  en fonction de la tension d'entrée  $u_e$  dans le premier ;
- l'intensité de sortie  $i_2$  en fonction de l'intensité d'entrée  $i_e$  dans le second ;

et justifier les noms de ces circuits.

3. Les défibrillateurs utilisés pour secourir les personnes en arrêt cardiaque fonctionnent sur la base d'un condensateur. Celui-ci se charge grâce à une batterie, puis se décharge entre les deux palettes placées de part et d'autre du cœur de la victime. Le choc électrique ainsi administré permet de faire cesser une éventuelle fibrillation et donne au cœur une chance de reprendre un rythme normal. Le choc est efficace tant que la tension entre les palettes reste au moins égale à 50% de sa valeur initiale, et doit être administré en une durée maximale de l'ordre de 14 ms.

des condensateurs, et déterminer cette fois la capacité équivalente.

2. En électrocinétique, on utilise couramment les deux montages suivants, respectivement qualifiés de **diviseur de tension** et **diviseur de courant** :

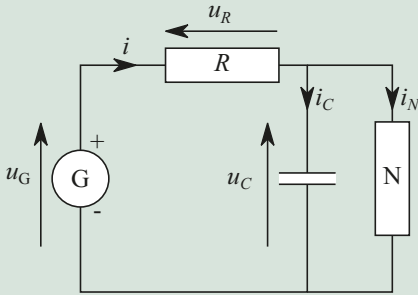


Déterminer la valeur de la capacité du condensateur à utiliser, ainsi qu'un ordre de grandeur de l'intensité du courant qui va traverser la victime durant l'administration du choc.

Données :

- tension du condensateur avant décharge :  $u_{C(t=0)} = E = 1,5 \text{ kV}$  ;
- résistance électrique opposée au passage du courant par la portion de torse située entre les palettes :  $R = 50 \text{ } \Omega$  ;
- approximation :  $\ln 2 \approx 0,70$ .

4. Une lampe au néon peut être représentée par le circuit suivant :



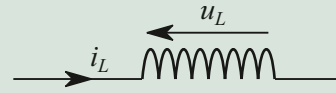
Le générateur délivre une tension constante  $u_G = E = 200 \text{ V}$ , le conducteur ohmique présente une résistance  $R = 10 \text{ M}\Omega$ , le condensateur une capacité  $C = 150 \text{ nF}$  et le composant N, qui assure la partie lumineuse, se comporte comme suit :

- tant que la lampe est éteinte, il se comporte comme un interrupteur ouvert ;
- elle s'allume seulement si la tension à ses bornes est supérieure ou égale à la valeur d'allumage  $u_{N,\text{all}} = 90 \text{ V}$  ;

- à partir du moment où elle est allumée, il se comporte comme un conducteur ohmique de résistance  $R_N = 10 \text{ k}\Omega$  ;
- elle s'éteint seulement si la tension à ses bornes est inférieure ou égale à la valeur d'extinction  $u_{N,\text{ext}} = 70 \text{ V}$ .

Justifier que celle-ci émette périodiquement des flashes lumineux, et déterminer l'expression  $u_C(t)$  de la tension aux bornes du condensateur dans les différentes phases de fonctionnement de la lampe.

5. Une bobine de fil enroulé sur un cylindre constitue un dipôle récepteur parfois qualifié de *self*. On la représente par le symbole suivant :



La relation entre la tension  $u_L$  aux bornes d'un tel dipôle et l'intensité  $i_L$  du courant qui la traverse est (en convention récepteur) de la forme :

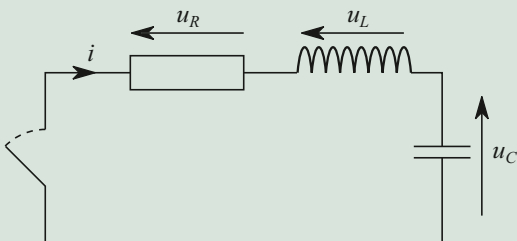
$$u_L = L \times \frac{di_L}{dt}$$

où  $L$  est une constante positive appelée **inductance** de la bobine (USI : le henry, symbole : H).

Sachant que l'intensité  $i_L$  du courant électrique circulant à travers une bobine est forcément une fonction continue du temps, déterminer la réponse indicielle d'un circuit RL, constitué par l'association en série d'une telle bobine avec un conducteur ohmique de résistance  $R$ .

## Vers la prépa

On associe en série un condensateur de capacité  $C = 1,00 \text{ } \mu\text{F}$ , un conducteur ohmique de résistance  $R = 100 \text{ } \Omega$  et une bobine d'inductance  $L = 1,0 \text{ H}$  (cf. Du Tac au Tac n° 5 concernant ce dernier dipôle).



Le condensateur, initialement chargé, présente une tension  $u_{C(t=0)} = E = 5,0 \text{ V}$ .

1. Montrer que la fonction suivante est solution de l'équation décrivant la réponse libre de ce circuit :

$$u_C(t) = K e^{-\frac{t}{\tau}} \cos(\omega t + \varphi)$$

à condition que les paramètres  $\tau$  et  $\omega$  vérifient des expressions que l'on précisera en fonction des différents paramètres du circuit.

2. Déterminer les valeurs des constantes d'intégration  $K$  et  $\varphi$ .

3. Décrire l'allure de la courbe représentative de  $u_C(t)$  en précisant les valeurs numériques de ses différentes caractéristiques.

## Halte aux idées reçues

1. Cet énoncé commet une erreur de forme, qui cependant trahit souvent une mécompréhension de fond : l'intensité N'EST PAS une grandeur orientée. Tout au mieux est-elle signée, et ce signe dépend-il d'un phénomène physique orienté, à savoir la circulation du courant électrique auquel est associée cette intensité. De manière générale, vous devez faire très attention à ne pas confondre, en physique :

- les phénomènes étudiés : il s'agit d'événements au cours desquels un changement a été observé dans les caractéristiques physiques du système d'étude ;
- les objets engagés dans ces phénomènes : éléments matériels constituant le système d'étude, et éventuellement affectés par le phénomène étudié ;
- les grandeurs physiques mesurées sur ces objets : nombre de fois (valeur) où une certaine caractéristique d'un objet reproduit un étalon de référence de cette caractéristique (unité).

Dans le cas qui nous occupe :

- le phénomène étudié est la circulation de particules électriquement chargées, autrement dit un courant électrique ;
- les objets affectés par la circulation de ce courant électrique sont les différentes pièces constitutives du circuit : dipôles, fils, ampèremètre. La flèche figurant sur le circuit représente le sens du courant dont ces objets sont le siège, et que l'on choisit d'étudier ;
- la grandeur caractérisant le phénomène étudiée est l'intensité.

2. Nous n'avons abordé dans le cours que deux façons d'associer deux dipôles, mais celles-ci sont en fait très restrictives :

- l'association en dérivation impose que les bornes des deux dipôles soient liées deux à deux ;
- l'association en série impose que les deux dipôles soient liés par une seule borne, **et que celle-ci ne constitue pas un nœud du circuit.**

Ainsi si deux dipôles possèdent une unique borne commune ils sont effectivement associés directement, mais si un troisième dipôle intervient également sur cette borne, les deux premiers ne seront plus considérés comme étant associés en série.

3. Comme nous l'avons dit, dans toute branche d'un circuit nous pouvons choisir de travailler avec un courant électrique

orienté dans un certain sens, ou avec le courant d'intensité opposée et orienté dans le sens contraire. Formellement, les deux reviennent exactement au même. Il est juste nécessaire de se montrer vigilant(e) par rapport au fait que beaucoup des résultats que vous tenez pour acquis supposent implicitement que vous avez travaillé avec un **courant d'intensité positive**.

Or l'intensité du courant délivré par un générateur est positive, uniquement si ce courant circule du pôle positif du générateur, vers son pôle négatif.

La question, qui peut sembler superflue en courant continu, revêtira une importance particulière lorsque vous apprendrez à travailler en courant alternatif. En effet dans ce cas, les bornes de potentiels respectivement fort et faible alternent sans cesse, et avec elles le sens du courant électrique d'intensité positive. Plutôt que de prendre en compte ces interversions à répétitions, on préfère fixer un pôle positif et un pôle négatif. La convention générateur permet alors de fixer le sens du courant électrique, et les interversions de bornes se résument finalement à des changements de signe pour la tension aux bornes du générateur, et avec elle l'intensité du courant qu'il débite.

4. Il est vrai que nous avons défini l'intensité instantanée dans une branche de circuit, comme le rapport d'une charge électrique élémentaire  $\delta q$  transitant à travers cette branche en une durée élémentaire  $dt$ , à la durée en question. La forme ainsi obtenue :

$$i = \frac{\delta q}{dt}$$

fait effectivement penser, formellement, à une dérivée. Vous ne devez cependant pas perdre de vue que la dérivée d'une fonction par rapport à une variable est la limite du taux d'accroissement de cette fonction par rapport à cette variable, lorsque l'intervalle sur lequel est calculé ce taux tend vers 0 :

$$\left( \frac{df}{dt} \right)_{t=t_0} = \lim_{t \rightarrow t_0} \left[ \frac{f(t) - f(t_0)}{t - t_0} \right]$$

Le concept de dérivée n'a donc de sens que si nous sommes en mesure d'identifier une fonction dont la valeur varie au cours du temps.

Par exemple dans le cas où l'intensité considérée est celle d'un courant circulant dans une branche contenant un condensateur, nous avons vu que toute charge transitant dans cette branche venait s'accumuler sur une armature du condensateur. Nous avons de ce fait pu assimiler la charge  $\delta q$  transitant dans la branche à la variation  $dq_C$  de la charge  $q_C$  portée par

l'armature du condensateur par laquelle arrive le courant, et ce faisant obtenir la relation :

$$i_C = \frac{dq_C}{dt}$$

Donc dans le cas particulier où la branche dans laquelle circule le courant étudié contient une zone où les charges peuvent s'accumuler (autrement dit un condensateur), nous pouvons effectivement décrire l'intensité du courant comme la dérivée temporelle de la charge portée par cette zone. Mais dans le cas contraire, nous ne serions même pas en mesure d'identifier clairement où se situe la charge dont nous sommes en train de parler.

Il est vrai que le formalisme utilisé en physique dans les affaires de grandeurs élémentaires pousse à la faute. En effet :

- pour une grandeur  $G$  échangée, on note une **portion élémentaire échangée**  $\delta G$ , avec un «  $\delta$  » minuscule, donc. Si l'on considère un échange macroscopique, on notera simplement :

$$G \begin{matrix} \text{éch} \\ \text{entre} \\ t_1 \text{ et } t_2 \end{matrix} = \int_{t_1}^{t_2} \delta G$$

- si en revanche cette grandeur n'est pas une grandeur échangée entre deux systèmes, mais stockée au sein d'un système, on note alors la **variation élémentaire**  $dG$ , avec un «  $d$  » minuscule et non un  $\delta$ . Et si l'on veut évoquer une variation macroscopique, on notera alors :

$$(\Delta G) \begin{matrix} \text{entre} \\ t_1 \text{ et } t_2 \end{matrix} = G(t_2) - G(t_1) = \int_{t_1}^{t_2} dG$$

C'est donc cette fois à la variation macroscopique qu'est affectée la lettre « delta », cette fois prise dans sa version majuscule.

En résumé :

| Échelle       | Grandeur échangée | Grandeur stockée |
|---------------|-------------------|------------------|
| Élémentaire   | $\delta G$        | $dG$             |
| Macroscopique | $G_{\text{éch}}$  | $\Delta G$       |

**Remarque :** ce point peut sembler anecdotique. Il recouvre cependant une importance essentielle, à la fois du point de vue physique (bien connaître la signification des grandeurs utilisées, et toujours être capable de les interpréter) mais également du point de vue mathématique. La question du caractère intégrable ou

non d'une fonction pose en effet de vraies questions de fond, en particulier concernant la conservativité de certaines grandeurs. Ces considérations sont au centre des propriétés de certaines forces en mécanique, et surtout des flux d'énergie en thermodynamique, qui pose la question de la réversibilité des phénomènes.

5. Cet énoncé assimile les régimes transitoires et le fonctionnement en courant alternatif. La confusion est compréhensible, les deux introduisant l'idée de grandeurs électrocinétiques variant au cours du temps.

Les régimes transitoires, nous l'avons dit, consistent en l'évolution de ces grandeurs sous l'effet d'un changement structurel du circuit. Les points de fonctionnement des différents dipôles vont en effet devoir évoluer pour s'adapter aux nouvelles contraintes du circuit, et c'est de cette évolution dont nous parlerons.

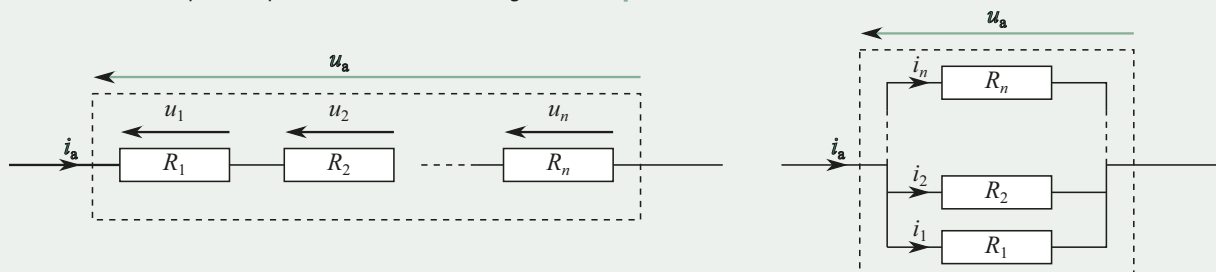
Cependant ces régimes transitoires (du moins ceux que nous abordons pour le moment) finissent tous par s'amortir en exponentielle (les solutions générales de l'équation sans second membre), et disparaître au profit des régimes permanents (les solutions particulières de l'équation avec second membre). Or ces solutions permanentes ne concernent pas uniquement les fonctionnements en courant continu. Les tensions et intensités relatives aux dipôles d'un circuit alimenté en courant alternatif, bien qu'oscillant au cours du temps, effectueront ces oscillations avec des amplitudes et des déphasages précis par rapport au générateur, mais dépendant du circuit. Ce sont les évolutions de ces paramètres, que décrivent les régimes transitoires.

Il importe par ailleurs de conserver à l'esprit que le choix du mode continu ou alternatif sur un multimètre conditionne le mode de calcul de la valeur affichée : valeur moyenne dans le cas continu, valeur efficace (c'est-à-dire moyenne quadratique des écarts à la moyenne, mais laissons-en un peu pour vos années en CPGE) dans le cas alternatif.

## Du Tac au Tac

1. L'énoncé nous demande de montrer qu'une association de conducteurs ohmiques équivaut à un conducteur ohmique unique. La première chose à faire est de bien cadrer ce qu'est un conducteur ohmique. Le cours nous dit qu'il s'agit d'un dipôle répondant à la loi d'Ohm, c'est-à-dire dont la caractéristique est une relation linéaire entre tension à ses bornes et intensité du courant qui le traverse.

Considérons donc les dipôles résultant des deux associations évoquées, et tâchons d'établir le lien entre la tension à leur borne et l'intensité du courant qui les traverse :



Dans le cas de l'association en série, l'association ne présente aucun nœud et la loi des nœuds ne nous apporte donc pas d'information, si ce n'est que tous les conducteurs ohmiques sont traversés par le même courant, d'intensité  $i_a$ .

La loi des mailles, elle, nous donne l'égalité :

$$u_a = u_1 + u_2 + \dots + u_n$$

En exprimant alors la loi d'Ohm pour chacun des conducteurs ohmiques, il vient :

$$u_a = R_1 \times i_a + R_2 \times i_a + \dots + R_n \times i_a \Leftrightarrow$$

$$u_a = (R_1 + R_2 + \dots + R_n) \times i_a$$

Nous trouvons donc bien une relation de proportionnalité entre la tension aux bornes du dipôle constitué par l'association en série de ces conducteurs ohmiques, et l'intensité du courant qui le traverse. Tout se passe donc comme si cette association était en fait un conducteur ohmique unique (elle suit la loi d'Ohm), dont la résistance équivalente serait :

$$R_{\text{eq},s} = \sum_{k=1}^n R_k$$

Dans le cas de l'association en parallèle, la tension est la même aux bornes de tous les dipôles et c'est cette fois la loi des mailles qui ne nous est d'aucun secours.

La loi des nœuds, en revanche, nous donne :

$$i_a = i_1 + i_2 + \dots + i_n$$

En exprimant alors la loi d'Ohm pour chacun des conducteurs ohmiques, il vient :

$$i_a = G_1 \times u_a + G_2 \times u_a + \dots + G_n \times u_a \Leftrightarrow$$

$$i_a = (G_1 + G_2 + \dots + G_n) \times u_a$$

Nous retrouvons donc à nouveau la loi d'Ohm : tout se passe comme si cette association était un conducteur ohmique unique, dont la résistance équivalente répondrait cette fois à l'égalité :

$$G_{\text{eq},p} = \sum_{k=1}^n G_k \Leftrightarrow \frac{1}{R_{\text{eq},p}} = \sum_{k=1}^n \frac{1}{R_k}$$

**Remarque :** on montre facilement que pour une association en parallèle de deux conducteurs ohmiques seulement, la résistance équivalente se résume à :

$$R_{\text{eq},p} = \frac{R_1 \times R_2}{R_1 + R_2}$$

Cette expression suffit en général, quitte, pour plus de deux conducteurs ohmiques, à procéder par associations successives (on calcule d'abord la résistance équivalente aux deux premiers, puis l'on associe celle-ci avec la troisième et ainsi de suite). Il importe surtout d'avoir conscience que cette égalité croît rapidement en complexité avec le nombre de conducteurs ohmiques engagés.

La démarche est exactement la même pour des condensateurs. Il faut seulement prendre garde au fait que cette fois la capa-

cité multiplie un terme de tension pour donner l'intensité, de la même façon que la conductance dans le cas des conducteurs ohmiques. Nous trouverons donc :

|                                                           | Association<br>en série                                  | Association<br>en parallèle                              |
|-----------------------------------------------------------|----------------------------------------------------------|----------------------------------------------------------|
| <b>Conducteurs<br/>ohmiques<br/>(version résistance)</b>  | $R_{\text{eq},s} = \sum_{k=1}^n R_k$                     | $\frac{1}{R_{\text{eq},p}} = \sum_{k=1}^n \frac{1}{R_k}$ |
| <b>Conducteurs<br/>ohmiques<br/>(version conductance)</b> | $\frac{1}{G_{\text{eq},s}} = \sum_{k=1}^n \frac{1}{G_k}$ | $G_{\text{eq},p} = \sum_{k=1}^n G_k$                     |
| <b>Condensateurs<br/>(capacités)</b>                      | $\frac{1}{C_{\text{eq},s}} = \sum_{k=1}^n \frac{1}{C_k}$ | $C_{\text{eq},p} = \sum_{k=1}^n C_k$                     |

**2.** Il s'agit ici simplement de mettre en œuvre les lois de base de l'électrocinétique de manière à isoler une grandeur particulière en fonction des autres paramètres du problème.

Le circuit **diviseur de tension**, ne présente aucun nœud. La loi des mailles, quant à elle, nous donne l'égalité :

$$u_e = u_1 + u_2$$

en notant  $u_1$  la tension aux bornes du premier conducteur ohmique en convention récepteur.

Nous pouvons compléter cette égalité en faisant intervenir la loi d'Ohm :

$$u_1 = R_1 \times i_e \quad \text{et} \quad u_2 = R_2 \times i_e$$

La loi des mailles nous donne alors :

$$u_e = R_1 \times i_e + R_2 \times i_e \Leftrightarrow i_e = \frac{u_e}{R_1 + R_2}$$

En réinjectant cette expression sur la loi d'Ohm pour le second conducteur ohmique, nous obtenons alors :

$$u_2 = R_2 \times \frac{u_e}{R_1 + R_2} \Leftrightarrow u_2 = \frac{u_e}{1 + \frac{R_1}{R_2}}$$

La tension en sortie du circuit (entendre : aux bornes du second conducteur ohmique) est donc directement proportionnelle à la tension d'entrée, *via* un coefficient qui va en diminuer la valeur (division par un nombre nécessairement  $> 1$ ), d'où son titre de diviseur de tension.

Dans le circuit **diviseur de courant**, la loi des nœuds nous donne cette fois :

$$i_e = i_1 + i_2$$

en notant  $i_1$  l'intensité du courant circulant dans le premier conducteur ohmique en convention récepteur.

Nous pouvons compléter cette égalité en faisant intervenir la loi d'Ohm :

$$i_1 = \frac{u_e}{R_1} \quad \text{et} \quad i_2 = \frac{u_e}{R_2}$$

La loi des nœuds nous donne alors :

$$i_e = \frac{u_e}{R_1} + \frac{u_e}{R_2} \Leftrightarrow u_e = i_e \times \frac{R_1 \times R_2}{R_1 + R_2}$$

En réinjectant cette expression sur la loi d'Ohm pour le second conducteur ohmique, nous obtenons alors :

$$i_2 = i_e \times \frac{R_1 \times R_2}{R_1 + R_2} \times \frac{1}{R_2} \Leftrightarrow u_2 = \frac{i_e}{1 + \frac{R_2}{R_1}}$$

L'intensité en sortie du circuit (entendre : dans la branche contenant le second conducteur ohmique) est donc directement proportionnelle à l'intensité en entrée, *via* un coefficient qui va en diminuer la valeur (division par un nombre nécessairement  $> 1$ ), d'où son titre de diviseur de courant.

**3.** La donnée clé pour déterminer la capacité du condensateur est naturellement la durée de demi-décharge, puisque l'énoncé précise que cette tension ne doit pas baisser de plus de 50% en une durée de 14ms. Nous pouvons donc écrire :

$$\frac{t_1}{2} = 14 \text{ ms}$$

Or nous savons que lors de la décharge d'un condensateur, la constante de temps caractéristique de l'exponentielle et la durée nécessaire à sa diminution de moitié sont telles que :

$$\tau = RC \text{ et } \frac{t_1}{2} = \tau \times \ln 2 \Leftrightarrow \frac{t_1}{2} = RC \times \ln 2$$

Il nous suffit alors d'isoler la capacité :

$$C = \frac{1}{R} \times \frac{\frac{t_1}{2}}{\ln 2} = \frac{1}{50 \Omega} \times \frac{14 \text{ ms}}{\ln 2} = 0,40 \text{ mF}$$

Sachant que la tension aux bornes du condensateur a pour expression :

$$u_C(t) = E \times e^{-\frac{t}{\tau}}$$

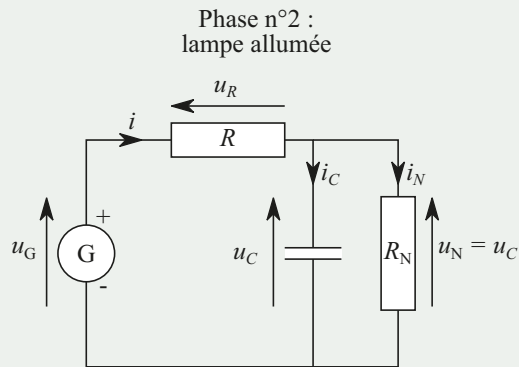
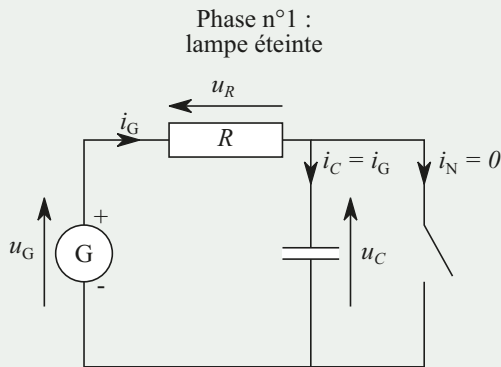
nous trouvons l'intensité associée grâce à la relation :

$$i = C \times \frac{du_C}{dt} = -\frac{1}{\tau} \times C \times E \times e^{-\frac{t}{\tau}} \rightarrow i = -\frac{E}{R} e^{-\frac{t}{\tau}}$$

L'intensité du courant est donc, en valeur absolue, de l'ordre de :

$$\frac{E}{R} = \frac{1,5 \cdot 10^3 \text{ V}}{50 \Omega} = 30 \text{ A}$$

**4.** Le dipôle proposé dans cet exercice est relativement sophistiqué ; procédons donc à son analyse méthodique. Changeant d'état selon la valeur de la tension à ses bornes, il équivaut donc selon les situations à l'un ou l'autre des circuits ci-dessous :



Tâchons donc déjà de comprendre pourquoi la lampe émettrait des flashes :

- au début, la lampe est éteinte et le circuit équivaut donc au premier schéma ci-dessus, dont nous connaissons bien le comportement : le condensateur va se charger selon une exponentielle à argument négatif renversée, de constante de temps associée :  $\tau_1 = RC = 1,5 \text{ s}$  en tendant asymptotiquement vers la valeur  $u_{C,\infty} = E$  ;
- or le néon se trouvant associé en parallèle avec le condensateur, la tension à ses bornes est égale à celle aux bornes du condensateur et croît donc de manière identique.  $u_{N,\text{all}}$  étant inférieure à  $E$ , la tension aux bornes du néon va finir par atteindre la valeur  $u_{N,\text{all}}$ , provoquant l'allumage du néon ;

- le néon étant allumé, le circuit équivalent est désormais celui de droite. Le condensateur va se comporter comme un second générateur et va donc se décharger à travers le conducteur ohmique de résistance  $R_N$  ;
- ce faisant, la tension aux bornes du condensateur (et donc celle aux bornes du composant N) va baisser. Si l'on suppose qu'elle finit par tomber au niveau de la tension d'extinction  $u_{N,\text{ext}}$  (ce que nous invite clairement à supposer l'énoncé, autrement la présence de flashes serait inexplicable), la lampe va finir par s'éteindre et le circuit redeviendra celui de gauche ;
- nous retrouvons le circuit RC de départ, la tension aux bornes du condensateur recommence à monter (mais au départ de  $u_{N,\text{ext}}$ , cette fois-ci), et ainsi de suite.

La première phase (lorsque la lampe est éteinte et que le condensateur se charge) correspond à une réponse indicielle (nouvelle tension supérieure à la précédente), donc la solution est de la forme :

$$u_{C,1}(t) = E \left( 1 - e^{-\frac{t}{\tau_1}} \right)$$

Dans la deuxième phase, en revanche, l'intensité du courant n'est plus la même à travers le conducteur ohmique de résistance  $R$  et à travers la branche contenant le condensateur. Il nous faut donc établir les équations auxquelles obéit le régime transitoire dans ce nouveau contexte.

Procédons comme d'habitude : lois de Kirchhoff et caractéristiques des différents dipôles, en conservant à l'esprit que nous recherchons une équation différentielle portant sur  $u_C$ .

- La loi des nœuds nous donne :  $i_G = i_C + i_N$ .
- La loi des mailles, outre l'égalité des tensions aux bornes du condensateur et du conducteur ohmique de résistance  $R_N$ , nous donne :  $u_G = u_R + u_C$ .
- En appliquant les lois caractéristiques des différents dipôles récepteurs et en remplaçant la tension aux bornes du générateur par sa valeur, nous trouvons :

$$E = Ri_G + u_C \Leftrightarrow E = R \times (i_C + i_N) + u_C \Leftrightarrow$$

$$E = R \times \left( C \frac{du_C}{dt} + \frac{u_N}{R_N} \right) + u_C$$

Il ne nous reste plus qu'à faire valoir le fait que  $u_N = u_C$ , et nous trouvons enfin :

$$E = R \times \left( C \frac{du_C}{dt} + \frac{u_C}{R_N} \right) + u_C \Leftrightarrow \frac{du_C}{dt} + \frac{1}{C} \left( \frac{1}{R} + \frac{1}{R_N} \right) u_C = \frac{E}{RC}$$

Nous retrouvons donc une équation différentielle linéaire du premier ordre à coefficients constants avec second membre constant, dont nous identifions sans peine :

- la constante de temps caractéristique  $\tau_2$  :

$$\frac{1}{\tau_2} = \frac{1}{C} \left( \frac{1}{R} + \frac{1}{R_N} \right) = \frac{1}{RC} \left( 1 + \frac{R}{R_N} \right) \Leftrightarrow \tau_2 = \frac{\tau_1}{1 + \frac{R}{R_N}}$$

- l'asymptote horizontale :

$$\frac{1}{C} \left( \frac{1}{R} + \frac{1}{R_N} \right) u_{C,2,\infty} = \frac{E}{RC} \Leftrightarrow u_{C,2,\infty} = \frac{E}{1 + \frac{R}{R_N}} = 0,20 \text{ V}$$

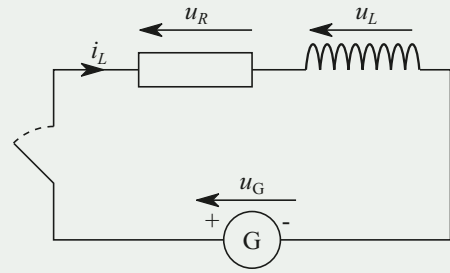
Nous constatons que cette valeur est effectivement inférieure à  $u_{N,\text{ext}}$ , donc la lampe va bien s'éteindre.

L'expression finale de cette seconde phase est donc :

$$u_{C,2}(t) = u_{C,2,\infty} e^{-\frac{t}{\tau_2}}$$

5. Commençons déjà par représenter le circuit RL évoqué dans l'énoncé : s'agissant d'une réponse indicielle, nous nous

interrogeons sur ce qu'il advient de l'intensité  $i_L$  lorsque nous fermons l'interrupteur dans le circuit ci-dessous :



Le circuit ne comportant aucun nœud, la loi des nœuds ne nous apporte rien si ce n'est le fait que l'intensité du courant est la même à travers tous les dipôles du circuit.

La loi des mailles, elle, nous donne (une fois l'interrupteur fermé) :

$$u_G = u_R + u_L$$

En utilisant les relations caractéristiques du conducteur ohmique et de la bobine, celle-ci devient, avec  $u_G = E$  et  $i_R = i_L$  :

$$u_G = R \times i_L + L \times \frac{di_L}{dt} \Leftrightarrow \frac{di_L}{dt} + \frac{R}{L} i_L = \frac{E}{L}$$

Nous retrouvons une équation différentielle linéaire d'ordre 1 à coefficients constants, avec second membre constant, dont nous savons que les solutions sont de la forme :

$$i_L(t) = K e^{-\frac{t}{\tau}} + \frac{E}{R}$$

avec  $\tau = \frac{L}{R}$ .

Nous déterminons la valeur de la constante  $K$  à l'aide de la condition initiale fournie, et de l'hypothèse de continuité donnée par l'énoncé concernant l'intensité du courant circulant à travers une bobine :

$$i_{L(t=0^+)} = i_{L(t=0^-)} \Leftrightarrow K \underbrace{e^{-\frac{0}{\tau}}}_1 + \frac{E}{R} = 0 \Leftrightarrow K = -\frac{E}{R}$$

En injectant ce résultat dans l'expression de  $i_L(t)$  trouvée plus haut, nous obtenons :

$$i_L(t) = -\frac{E}{R} e^{-\frac{t}{\tau}} + \frac{E}{R} \Leftrightarrow i_L(t) = \frac{E}{R} \left( 1 - e^{-\frac{t}{\tau}} \right)$$

Le formalisme rencontré ici est donc très proche de celui utilisé dans le cas du dipôle RC, mais les phénomènes mis en jeu diffèrent. L'expression consacrée n'est plus ici charge (respectivement décharge) comme dans le cas du condensateur, mais simplement **établissement (respectivement rupture) du courant** à travers la bobine.

Par ailleurs la grandeur continue n'est plus ici une tension mais une intensité, tandis que la grandeur éventuellement discontinue n'est plus l'intensité mais la tension aux bornes de la bobine. Ceci a notamment pour effet d'entraîner de possibles

surtensions et la création d'étincelles (dites *étincelles de rupture*) si l'on interrompt brutalement la circulation du courant électrique (cas d'une réponse libre).

## Vers la prépa

Une nouvelle fois, le circuit ne présente pas de nœud ; la loi des nœuds est donc sans objet, si ce n'est nous affirmer que l'intensité du courant est la même à travers chacun des dipôles présents :

$$i_R = i_L = i_C = i$$

Exprimons à présent la loi des mailles ; il vient :

$$u_L + u_R + u_C = 0 \Leftrightarrow LC \frac{d^2 u_C}{dt^2} + RC \frac{du_C}{dt} + u_C = 0 \Leftrightarrow$$

$$\frac{d^2 u_C}{dt^2} + \frac{R}{L} \frac{du_C}{dt} + \frac{1}{LC} u_C = 0$$

avec  $i_L = i_R = i_C = C \frac{du_C}{dt}$  puisque le circuit ne comporte aucun nœud.

Pour vérifier que la fonction proposée est solution de l'équation différentielle ci-dessus, nous allons avoir besoin de ses dérivées successives. Celles-ci promettent d'être relativement copieuses, aussi allons-nous les calculer une par une au préalable :

- fonction de base :  $u_C(t) = Ke^{-\frac{t}{\tau}} \cos(\omega t + \varphi)$
- dérivée première, obtenue comme dérivée d'un produit de fonctions :

$$\frac{du_C}{dt} = \left[ -\frac{1}{\tau} Ke^{-\frac{t}{\tau}} \right] \cos(\omega t + \varphi) + Ke^{-\frac{t}{\tau}} \left[ -\omega \sin(\omega t + \varphi) \right]$$

$$\text{soit encore : } \frac{du_C}{dt} = -Ke^{-\frac{t}{\tau}} \left[ \frac{1}{\tau} \cos(\omega t + \varphi) + \omega \sin(\omega t + \varphi) \right]$$

- dérivée seconde :

$$\frac{d^2 u_C}{dt^2} = \left[ +\frac{1}{\tau} Ke^{-\frac{t}{\tau}} \right] \left[ \frac{1}{\tau} \cos(\omega t + \varphi) + \omega \sin(\omega t + \varphi) \right]$$

$$- Ke^{-\frac{t}{\tau}} \left[ -\frac{\omega}{\tau} \sin(\omega t + \varphi) + \omega^2 \cos(\omega t + \varphi) \right]$$

et finalement :

$$\frac{d^2 u_C}{dt^2} = Ke^{-\frac{t}{\tau}} \left[ \left( \frac{1}{\tau^2} - \omega^2 \right) \cos(\omega t + \varphi) + 2 \frac{\omega}{\tau} \sin(\omega t + \varphi) \right]$$

Substituons alors ces différentes expressions dans l'équation différentielle ; le facteur  $Ke^{-\frac{t}{\tau}}$  est commun aux trois termes et en factorisant par celui-ci, nous obtenons :

$$Ke^{-\frac{t}{\tau}} \left\{ \begin{aligned} & \left[ \left( \frac{1}{\tau^2} - \omega^2 \right) \cos(\omega t + \varphi) + 2 \frac{\omega}{\tau} \sin(\omega t + \varphi) \right] \\ & - \frac{R}{L} \left[ \frac{1}{\tau} \cos(\omega t + \varphi) + \omega \sin(\omega t + \varphi) \right] \\ & + \frac{1}{LC} \cos(\omega t + \varphi) \end{aligned} \right\} = 0$$

Sachant que l'exponentielle est toujours strictement positive et en supposant que la constante  $K$  est non nulle (autrement le problème n'ira pas bien loin), nous simplifions ; par ailleurs en regroupant les termes relatifs respectivement au sinus et au cosinus, nous trouvons :

$$\left[ \left( \frac{1}{\tau^2} - \omega^2 \right) - \frac{R}{L} \times \frac{1}{\tau} + \frac{1}{LC} \right] \cos(\omega t + \varphi)$$

$$+ \left[ \omega \left( \frac{2}{\tau} - \frac{R}{L} \right) \right] \sin(\omega t + \varphi) = 0$$

Or un cosinus et un sinus contenant le même argument ne sont jamais nuls simultanément (et même si tel était le cas, une montre arrêtée est à l'heure deux fois par jour, or l'équation différentielle doit être satisfaite à tout instant). L'égalité ci-dessus ne peut donc être vérifiée en tout instant que si les coefficients portant respectivement sur le sinus et le cosinus sont nuls, autrement dit à la double condition suivante :

$$\left( \frac{1}{\tau^2} - \omega^2 \right) - \frac{R}{L} \times \frac{1}{\tau} + \frac{1}{LC} = 0 \quad \text{et} \quad \omega \left( \frac{2}{\tau} - \frac{R}{L} \right) = 0$$

Ici encore, si  $\omega$  était nulle, le problème n'irait pas loin ; la seconde condition se résume donc à :

$$\tau = \frac{2L}{R}$$

Nous pouvons alors injecter cette expression dans la première condition, qui devient :

$$\left[ \left( \frac{R}{2L} \right)^2 - \omega^2 \right] - \frac{R}{L} \times \frac{R}{2L} + \frac{1}{LC} = 0 \Leftrightarrow \omega^2 = \frac{1}{LC} - \left( \frac{R}{2L} \right)^2$$

**Remarque :** notons que  $\omega^2$  étant nécessairement positif, cette égalité est possible uniquement si :

$$\frac{1}{LC} - \left( \frac{R}{2L} \right)^2 \geq 0 \Leftrightarrow R \leq 2\sqrt{\frac{L}{C}} = 2,00 \text{ k}\Omega$$

Cette valeur est appelée **résistance critique** du circuit : si la résistance est trop élevée, les effets dissipatifs de la résistance empêchent la naissance des oscillations.

Nous pouvons alors calculer les valeurs des deux paramètres demandés :

$$\tau = \frac{2L}{R} = 20 \text{ ms} \quad \text{et} \quad \omega = \sqrt{\frac{1}{LC} - \left( \frac{R}{2L} \right)^2} = 1,0 \cdot 10^3 \text{ rad.s}^{-1}$$

Nous pouvons alors associer à la pulsation une fréquence et une période, cette dernière nous indiquant sur quelle durée s'étale une oscillation de la fonction sinusoïdale :

$$f = \frac{\omega}{2\pi} = 1,6 \cdot 10^2 \text{ Hz} \Leftrightarrow T = \frac{1}{f} = 6,3 \text{ ms}$$

Il nous reste enfin à déterminer les valeurs des constantes  $K$  et  $\varphi$ . Comme d'habitude, celles-ci sont fixées par des conditions particulières, en l'occurrence les conditions initiales, que nous pouvons faire valoir sur les grandeurs continues en fonction du temps :

- tension aux bornes du condensateur :

$$u_{C(t=0^+)} = u_{C(t=0^-)} \Leftrightarrow K \cos \varphi = E$$

- intensité du courant à travers la bobine :

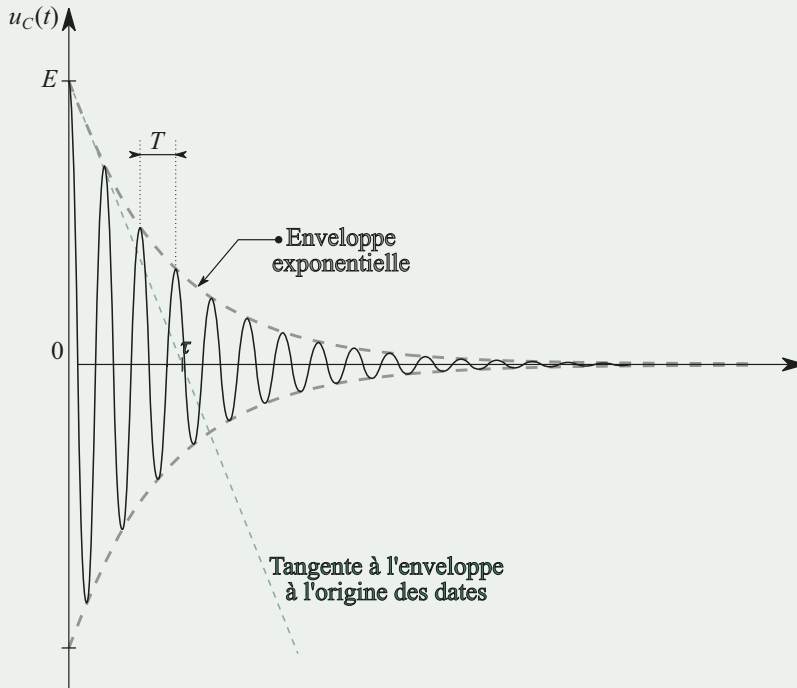
$$i_{L(t=0^+)} = i_{L(t=0^-)} \Leftrightarrow \frac{\cos \varphi}{\tau} + \omega \sin \varphi = 0 \Leftrightarrow$$

$$\varphi = \text{Arctan}\left(-\frac{1}{\omega\tau}\right) = -5,0 \cdot 10^{-2} \text{ rad}$$

Disposant désormais de  $\varphi$ , nous complétons la résolution en déterminant la valeur de la constante  $K$  :

$$K = \frac{E}{\cos \varphi} = 5,0 \text{ V}$$

La courbe représentative présente finalement l'allure suivante :





## 11.1 Énergie et puissance

### — Qu'est-ce que l'énergie ?

On appelle **énergie** d'un système, une grandeur quantifiant la capacité de ce système à transformer son environnement ou lui-même. Son USI est le joule (symbole : J), équivalent à du  $\text{kg.m}^2.\text{s}^{-2}$ .

Cette définition très générale permet de rendre compte du fait que l'énergie se manifeste à travers de multiples phénomènes, dont elle est le dénominateur commun. En effet, s'il est une pierre angulaire de la physique, c'est le fait que **l'énergie est une grandeur conservative** : elle ne peut pas plus être créée *ex nihilo* qu'elle ne peut disparaître.

### — Où l'énergie est-elle localisée ?

Formellement, on considère régulièrement l'énergie en tant que grandeur **stockée** par un système physique, ou **échangée** entre deux systèmes. On peut faire l'analogie entre l'argent stocké sur un compte en banque, et l'argent échangé entre deux personnes : l'une constitue un capital susceptible de varier, l'autre une quantité qui transite entre le capital d'un(e) épargnant(e) et celui d'un(e) autre.

L'énergie stockée pouvant varier au cours du temps, on note :

- $dE$  une variation élémentaire de l'énergie, c'est-à-dire se produisant en une durée élémentaire  $dt$ . Physiquement, ceci signifie que cette variation est calculée sur une durée très courte par rapport à l'échelle de temps caractéristique sur laquelle le phénomène se déroule ;
- $\Delta E = E_f - E_i$  une variation macroscopique de l'énergie, c'est-à-dire se produisant en une durée  $\Delta t = t_f - t_i$  sur laquelle le système évolue significativement (les indices « i » et « f » désignent respectivement les états initial et final dans lesquels est considéré le système d'énergie  $E$ ). On peut écrire :

$$\Delta E = E_f - E_i = \int_{t_i}^{t_f} dE$$

Les principales formes de stockage que vous rencontrerez seront :

|                                            | Domaine mécanique                                                                                                                                                                                     | Domaine thermodynamique                                                                                                                                                                   |
|--------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Énergie principale...                      | Énergie mécanique $E_m$                                                                                                                                                                               | Énergie interne $U$                                                                                                                                                                       |
| Se détaillant sous les formes associées... | <ul style="list-style-type: none"> <li>• au mouvement du système <b>énergie cinétique <math>E_c</math></b></li> <li>• à la position du système <b>énergie potentielle <math>E_p</math></b></li> </ul> | <ul style="list-style-type: none"> <li>• à la température</li> <li>• aux changements d'état</li> <li>• aux transformations chimiques</li> <li>• aux transformations nucléaires</li> </ul> |

L'énergie échangée, elle, peut seulement *s'identifier* à la variation d'une grandeur : elle ne constitue pas un capital mais une monnaie d'échange dont le montant peut s'identifier à la variation d'un capital. Dans ce cas, on note :

- $\delta E_{\text{éch}}$  une portion élémentaire d'énergie échangée, c'est-à-dire se produisant en une durée élémentaire  $dt$ ;
- $E_{\text{éch}}$  une énergie échangée à l'échelle macroscopique, reliée à la portion élémentaire  $\delta E_{\text{éch}}$

par : 
$$E_{\text{éch},t_i \rightarrow t_f} = \int_{t_i}^{t_f} \delta E_{\text{éch}}.$$

Vous rencontrerez deux formes d'énergie échangée :

|                                               | Travail $W$                                                                                                                                                                                               | Transfert thermique $Q$ |
|-----------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------|
| Correspond à un transfert entre les formes... | <ul style="list-style-type: none"><li>cas d'une force conservative :<br/><math>E_c \Leftrightarrow E_p</math></li><li>cas d'une force non conservative :<br/><math>E_c \Leftrightarrow U</math></li></ul> | $U \leftrightarrow U$   |

L'énergie échangée n'est donc pas localisée sur un système, mais en transit entre deux systèmes. Le formalisme utilisé en matière énergétique considère régulièrement les deux aspects, et les confondre mène inévitablement à des non-sens.

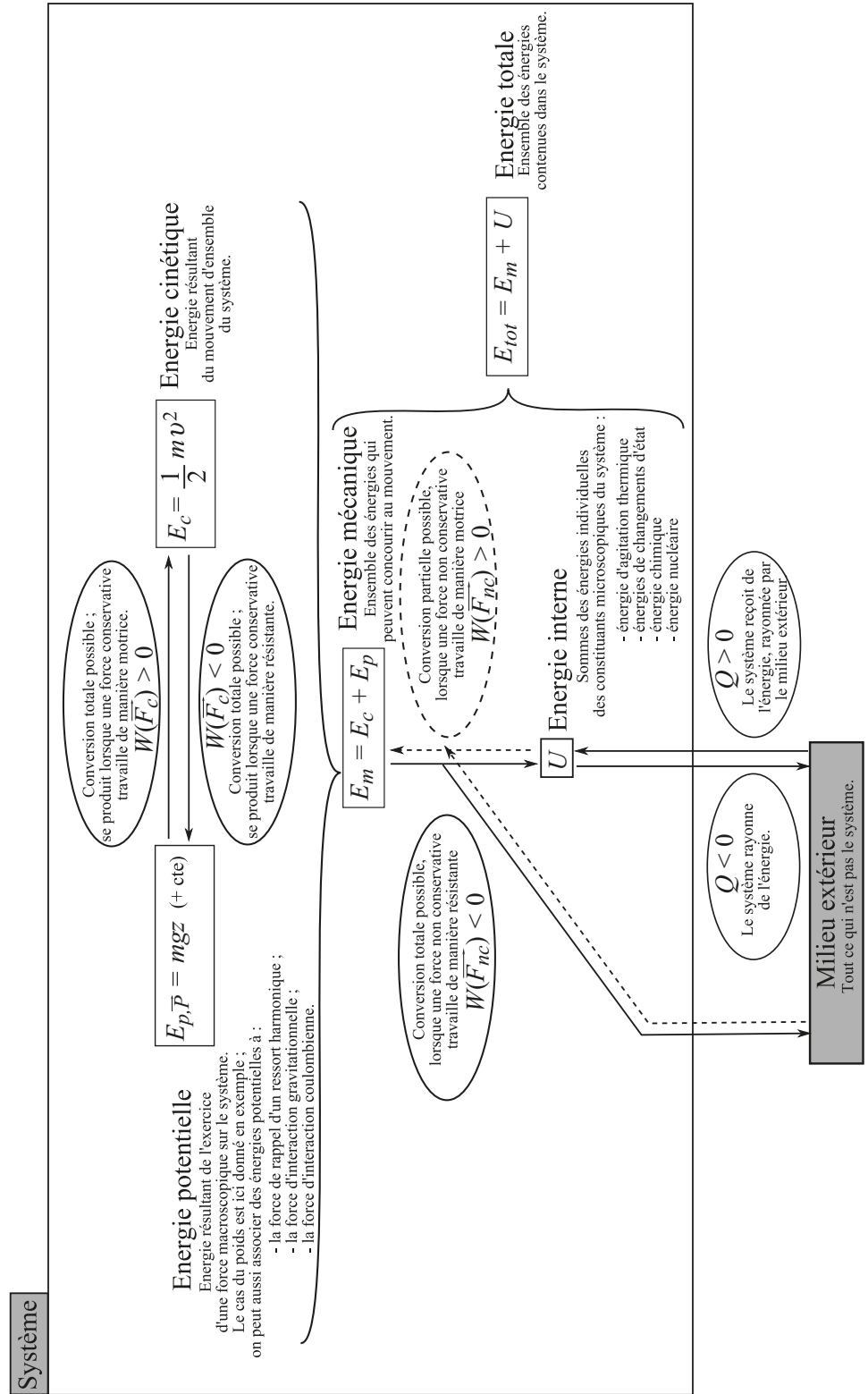
## Comment l'énergie se manifeste-t-elle ?

Il importe, tant pour la culture scientifique (qui reste votre première arme face à un résultat numérique suspect) que pour votre sécurité, d'avoir conscience que les quantités d'énergie engagées par les différents phénomènes qui nous entourent couvrent une gamme très variée d'ordres de grandeur (les valeurs ci-dessous sont données en joule d'énergie par kilogramme de matière subissant le phénomène indiqué) :

| Chute libre de 1 m dans le champ de pesanteur terrestre | Variation de température de 1 °C | Changement d'état        |                          |                             |
|---------------------------------------------------------|----------------------------------|--------------------------|--------------------------|-----------------------------|
|                                                         |                                  | physique                 | chimique                 | nucléaire                   |
| $10^1 \text{ J.kg}^{-1}$                                | $10^3 \text{ J.kg}^{-1}$         | $10^5 \text{ J.kg}^{-1}$ | $10^7 \text{ J.kg}^{-1}$ | $10^{12} \text{ J.kg}^{-1}$ |

Réciproquement,  $10^9 \text{ J}$  d'énergie ne se manifestent pas de manière aussi spectaculaire selon la forme sous laquelle ils sont engagés :

| Sous forme mécanique                        | Sous forme thermique                                        | Changement d'état               |                                                      |                                            |
|---------------------------------------------|-------------------------------------------------------------|---------------------------------|------------------------------------------------------|--------------------------------------------|
|                                             |                                                             | physique                        | chimique                                             | nucléaire                                  |
| Tour Eiffel tombant en chute libre sur 10 m | Mène 3000 L d'eau de la température ambiante à l'ébullition | Fonte de 10 bonshommes de neige | Combustion de 50 L de combustible fossile (un plein) | Combustion de 1 g de combustible nucléaire |



## — Qu'est-ce que la puissance ?

Outre l'énergie sous toutes ses formes, vous serez régulièrement amené(e) à travailler avec la puissance. Celle-ci est définie comme le rapport de la variation d'énergie dont un système est le siège en une certaine durée, à la durée en question ; en version instantanée, ceci donne :

$$P = \frac{dE}{dt}$$

L'USI associée serait donc le joule par seconde ( $\text{J.s}^{-1}$ , correspondant à des  $\text{kg.m}^2.\text{s}^{-3}$ ), auquel on préfère le watt (symbole : W).

La puissance peut également être considérée comme échangée, auquel cas on remplacera  $dE$  par le travail élémentaire  $\delta W$  ou le transfert thermique élémentaire  $\delta Q$  correspondant.

La puissance est à l'énergie ce que la vitesse est à la position : un taux d'accroissement par rapport au temps qui permet de quantifier le rythme de ses variations. Par exemple l'énergie déployée en quelques secondes par une explosion thermonucléaire est du même ordre que celle consommée sous forme électrique par une ville telle que Paris sur une année entière : les effets d'un transfert d'énergie dépendent énormément de la durée sur laquelle celui-ci s'étale.

**Remarque :** en anglais, le terme « *power* » désigne à la fois les mots « pouvoir » et « puissance ». Il est intéressant de noter qu'en français ce dernier terme n'est lui-même ni plus ni moins qu'un dérivé du subjonctif du verbe « pouvoir » (tout comme la croissance est le fait de croître, la réjouissance le fait de se réjouir, etc.). On trouve ici une illustration sémantique de cette idée que la capacité de transformation d'un système, si elle est directement liée à l'énergie qu'il est capable de mettre en œuvre, doit également être considérée au regard de la durée sur laquelle s'étale cette mise en œuvre.

## — 11.2 Aspects énergétiques des phénomènes mécaniques

### — Existe-t-il une grandeur conservée au cours d'un mouvement mécanique ?

On constate que dans certaines circonstances, un système en mouvement se retrouve en plusieurs positions de sa trajectoire, animé d'une même vitesse. Par exemple lors d'un tir parabolique en l'absence de frottements, la valeur de la vitesse avec laquelle le projectile passe par une certaine altitude à la descente est la même que celle dont il était doté à la montée à l'altitude en question.

La loi de vitesse se démontre facilement à l'aide de la 2<sup>e</sup> loi de Newton, mais pose la question d'une grandeur qui serait possiblement conservée tout au long du mouvement (la quantité de mouvement, par exemple, ne l'est pas, et son taux de variation est précisément égal à la résultante des forces extérieures exercées sur le système).

## Comment peut-on mettre en évidence l'énergie cinétique ?

On peut chercher cette grandeur en se fondant sur la 2<sup>e</sup> loi de Newton ; en supposant l'étude menée par rapport à un référentiel galiléen, nous avons :

$$m \frac{d\vec{v}}{dt} = \sum \vec{F}_{\text{ext}} \Leftrightarrow m d\vec{v} = \sum \vec{F}_{\text{ext}} \times dt$$

L'intégration de cette équation à ce stade nous donnerait la loi de vitesse, mais suppose de connaître l'expression des forces extérieures en fonction du temps. Or la physique nous donne les expressions des forces par rapport à des variables d'espace. Une intégration du terme de forces par rapport aux variables d'espace aurait donc davantage de chances d'aboutir.

Nous savons que pour passer d'une durée à une longueur, l'intermédiaire est la vitesse. Plus précisément, le déplacement élémentaire  $d\vec{l} = dx \cdot \vec{e}_x + dy \cdot \vec{e}_y + dz \cdot \vec{e}_z$  réalisé en une durée  $dt$  par un système ponctuel animé d'une vitesse  $\vec{v}$  a pour expression :

$$d\vec{l} = \vec{v} \times dt$$

**Remarque :** l'expression du déplacement élémentaire fournie plus haut vous est donnée en coordonnées cartésiennes. Vous verrez en CPGE que ce déplacement élémentaire peut également s'exprimer dans les autres systèmes de coordonnées (cylindro-polaires, sphériques), plus adaptées à certaines forces. Les forces gravitationnelle et électrostatique, par exemple, sont à symétrie sphériques et se prêtent beaucoup mieux à un traitement en coordonnées sphériques qu'en coordonnées cartésiennes.

En effectuant de part et d'autre de l'équation donnée par le PFD le produit scalaire par la vitesse, on trouve ainsi :

$$m\vec{v} \cdot d\vec{v} = \sum \vec{F}_{\text{ext}} \cdot \vec{v} \cdot dt \Leftrightarrow m \underbrace{\vec{v} \cdot d\vec{v}}_{d\left(\frac{v^2}{2}\right)} = \sum \vec{F}_{\text{ext}} \cdot d\vec{l}$$

En intégrant de part et d'autre entre les positions initiale et finale du système, il vient alors :

$$\left[ \frac{1}{2}mv^2 \right]_i^f = \sum \int_{\mathcal{C}_{i \rightarrow f}} \vec{F}_{\text{ext}} \cdot d\vec{l}$$

où  $\mathcal{C}_{i \rightarrow f}$  désigne la courbe suivie par le système entre ses positions initiale et finale.

Nous constatons donc que l'énergie cinétique varie au cours du mouvement, à l'aune des intégrales le long du chemin suivi des produits scalaires des forces s'exerçant sur le système (on appelle ces intégrales les *circulations* des forces sur le chemin suivi).

On appelle alors :

- **énergie cinétique**  $E_c$  d'un système de masse  $m$  animé d'une vitesse de valeur  $v$ , la grandeur :

$$E_c = \frac{1}{2}mv^2$$

- **travail**  $W_{\mathcal{C}_{i \rightarrow f}}(\vec{F})$  d'une force  $\vec{F}$  le long d'un chemin  $\mathcal{C}_{i \rightarrow f}$ , l'intégrale :

$$W_{\mathcal{C}_{i \rightarrow f}}(\vec{F}) = \int_{\mathcal{C}_{i \rightarrow f}} \vec{F}_{\text{ext}} \cdot d\vec{l}$$

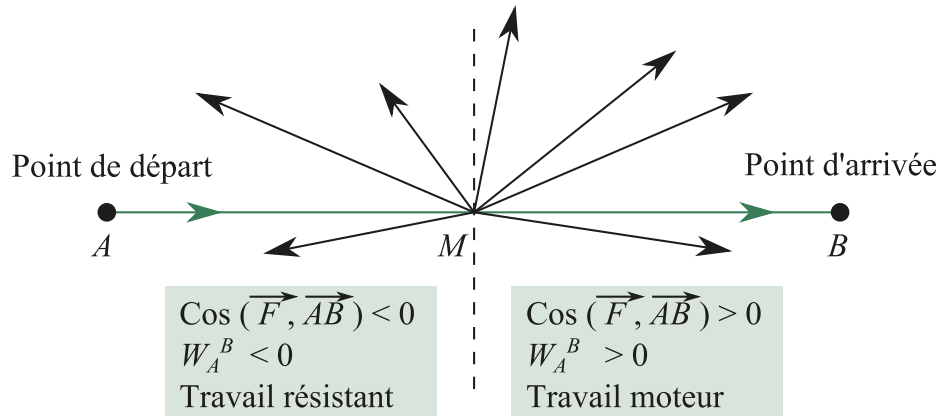
Dans le cas où la force  $\vec{F}$  est constante sur tout le chemin  $\mathcal{C}_A^B$ , on peut sortir  $\vec{F}$  de l'intégrale. L'intégrale de  $d\vec{OM}$  entre  $M = A$  et  $M = B$  donne alors  $\vec{OB} - \vec{OA} = \vec{AB}$ , d'où :

|                                                                                       |                       |
|---------------------------------------------------------------------------------------|-----------------------|
| $W_{\mathcal{C}_A^B}(\vec{F} = \text{cte}) = W_A^B(\vec{F}) = \vec{F} \cdot \vec{AB}$ | $W_A^B(\vec{F})$ en J |
|                                                                                       | $F$ en N              |
|                                                                                       | $AB$ en m             |

**Remarque :** dans le cas d'une force constante, le travail ne dépend que des positions  $A$  et  $B$  entre lesquelles se déplace le système sans aucune considération du chemin suivi pour se rendre de l'un à l'autre. Cette propriété concerne également certaines forces non constantes. De telles forces sont dites **conservatives**, et seront abordées dans le cadre du théorème de l'énergie mécanique.

On note alors que :

- Si la coordonnée de  $\vec{F}$  selon la direction du déplacement du système favorise ce déplacement, le travail est positif : il est alors dit **moteur**.
- Si  $\vec{F}$  est normale au déplacement, le travail est nul.
- Si la composante de  $\vec{F}$  selon la direction du déplacement du système contrevient à ce déplacement, le travail est négatif : il est alors dit **résistant**.



## — Qu'exprime le théorème de l'énergie cinétique ?

L'égalité vue plus haut peut alors se réécrire :

$$\Delta E_c = \sum W_{\mathcal{C}_{i \rightarrow f}}(\vec{F})$$

Nous trouvons ici l'expression du **théorème de l'énergie cinétique**, selon lequel la variation de l'énergie cinétique d'un système ponctuel entre deux positions de sa trajectoire, évaluée par rapport à un référentiel galiléen, est égale à la somme des travaux des forces s'exerçant sur ce système le long de cette trajectoire, entre ces deux mêmes positions.

**Remarque :** vous verrez en CPGE qu'en réalité si les forces internes au système travaillent, elles doivent également être prises en compte. Par ailleurs, notez bien que l'énergie cinétique se résume à  $1/2mv^2$  uniquement pour un système ponctuel : dans le cas contraire, des termes supplémentaires

liés à la rotation du système sur lui-même doivent également être pris en compte. Enfin, cette expression ne devra pas être abusivement étendue au cas d'un système relativiste : si la vitesse n'est plus négligeable devant la célérité de la lumière dans le vide, l'expression de  $E_c$  doit également être revue.

## — Qu'est-ce qu'une énergie potentielle ?

Il est intéressant de noter que les travaux des forces s'exprimant comme des intégrales, on peut se demander si ces intégrales correspondent ou non à la variation d'une primitive. En d'autres termes, si l'expression d'une force dépend uniquement des coordonnées de position du système sur lequel elle s'exerce, et qu'il existe une fonction dont cette force dérive, alors son travail ne dépendrait plus que des positions initiale et finale du système, à l'exclusion notamment du chemin suivi entre ces deux positions.

Lorsqu'une force vérifie cette condition, elle est dite **conservative** (nous verrons plus loin le sens de cette expression). La primitive qui lui est associée est homogène au produit d'une force par une longueur, autrement dit il s'agit d'une énergie.

On appelle alors **énergie potentielle** associée à une force  $\vec{F}$  conservative, la grandeur  $E_{p,\vec{F}}$  dont la variation entre deux positions du système s'identifie à l'opposé du travail de cette force entre ces deux positions :

$$dE_{p,\vec{F}} = -\delta W(\vec{F}) = -\vec{F} \cdot d\vec{l}$$

## — Quelle relation lie une force à l'énergie potentielle qui lui est associée ?

La variation d'énergie potentielle d'un système entre deux de ses positions  $A$  et  $B$  est donc égale à l'opposé du travail de cette force entre ces deux positions :

$$E_{p,B} - E_{p,A} = -W_A^B(\vec{F})$$

$$E_{p,A}, E_{p,B}, W_A^B(\vec{F}) \text{ en J}$$

En particulier, pour le travail élémentaire d'une force  $\vec{F} = F(x)\vec{e}_x$  ne dépendant que d'une seule coordonnée  $x$  et dirigée suivant l'axe de cette coordonnée, nous avons  $dE_p = -\delta W = -F(x)dx$ , d'où finalement :

$$F(x) = -\frac{dE_p}{dx}$$

|       |      |
|-------|------|
| $E_p$ | en J |
| $F$   | en N |
| $x$   | en m |

On constate ainsi qu'une force est conservative si et seulement si elle s'identifie à la dérivée d'une fonction par rapport aux coordonnées d'espace. Cette fonction est alors l'énergie potentielle associée à cette force. On dit d'une force conservative qu'elle **dérive d'une énergie potentielle**.

## Quelles sont les énergies potentielles associées aux forces les plus courantes ?

On trouve donc l'énergie potentielle  $E_p(x)$  d'une force  $F(x)$  à partir de la primitive :

|                           |       |      |
|---------------------------|-------|------|
| $E_p(x) = - \int F \, dx$ | $E_p$ | en J |
|                           | $F$   | en N |
|                           | $x$   | en m |

**Remarque :** lors du calcul de l'énergie potentielle associée à une force conservative, apparaît en réalité toujours une constante d'intégration. Ainsi trouverez-vous parfois un sybillin « +cte » dans l'expression d'une énergie potentielle. L'énergie potentielle n'est donc en réalité définie qu'à une constante additive près. Celle-ci ne doit cependant pas vous perturber : ne considérant à chaque fois qu'une **variation** d'énergie entre deux positions du système, la constante se supprimera d'elle-même lorsque nous effectuerons la différence d'énergie potentielle. Cette constante n'ayant aucun impact sur la suite du problème, nous sommes libres de la fixer à une valeur qui nous agréer. Typiquement dans le cas du poids, on la prend égale à 0, avec pour effet d'offrir une énergie potentielle de pesanteur nulle à l'altitude  $z = 0$ .

**L'énergie potentielle gravitationnelle :** deux corps de masses respectives  $m_A$  et  $m_B$ , ponctuels ou dont les répartitions de masses sont à symétrie sphérique, dont les centres de masses respectifs  $A$  et  $B$  sont séparés d'une distance  $r$ , possèdent une énergie potentielle d'interaction :

|                                                      |               |        |
|------------------------------------------------------|---------------|--------|
| $E_{p,g}(r) = -\mathcal{G} \frac{m_A \times m_B}{r}$ | $E_{p,g}$     | en J   |
|                                                      | $\mathcal{G}$ | en USI |
|                                                      | $m_A, m_B$    | en kg  |
|                                                      | $r$           | en m   |

**L'énergie potentielle de pesanteur :** un corps de masse  $m$ , situé à une altitude  $z$  au-dessus de la surface d'un astre générant un champ de pesanteur uniforme  $\vec{g}_A$ , possède une énergie potentielle d'interaction avec cet astre :

|                        |           |                |
|------------------------|-----------|----------------|
| $E_{p,p}(z) = m g_A z$ | $E_{p,p}$ | en J           |
|                        | $m$       | en kg          |
|                        | $g_A$     | en $N.kg^{-1}$ |

**Remarque :** lorsque  $z \ll R_A$ , rayon de l'astre, on peut montrer que cette énergie potentielle se déduit de la précédente, tout comme l'expression du poids d'un corps sur un astre se déduit de celle de l'attraction gravitationnelle exercée par cet astre sur ce corps, dans les mêmes conditions.

**L'énergie potentielle élastique d'un corps attaché à un ressort :** tout système attaché à un ressort de constante de raideur  $k$ , de longueur à vide  $l_0$  et étiré sur une longueur  $l$  possède une énergie potentielle de rappel :

|                                        |            |               |
|----------------------------------------|------------|---------------|
| $E_{p,él} = \frac{1}{2} k (l - l_0)^2$ | $E_{p,él}$ | en J          |
|                                        | $k$        | en $N.m^{-1}$ |
|                                        | $l, l_0$   | en m          |

## — Qu'est-ce que l'énergie mécanique, et dans quelles circonstances est-elle conservée ?

En séparant les forces conservatives d'une part, et les forces non conservatives d'autre part, le théorème de l'énergie cinétique devient :

$$\Delta E_c = \sum W_{C_i \rightarrow f}(\vec{F}_c) + \sum W_{C_i \rightarrow f}(\vec{F}_{nc})$$

Identifions alors les travaux des forces conservatives à l'opposé de la variation des énergies potentielles associées :

$$\Delta E_c = -\sum \Delta E_{p, \vec{F}_c} + \sum W_{C_i \rightarrow f}(\vec{F}_{nc}) \Leftrightarrow \Delta E_c + \sum \Delta E_{p, \vec{F}_c} = \sum W_{C_i \rightarrow f}(\vec{F}_{nc})$$

Ainsi, les forces conservatives (définies comme les forces dont le travail est indépendant du chemin suivi) sont des forces dont le travail permet un stockage de l'énergie cinétique sous une forme récupérable : l'énergie potentielle. L'énergie d'un système associée à son mouvement d'ensemble peut donc en fait :

- soit être exprimée sous forme de mouvement (énergie cinétique) ;
- soit être stockée sous une forme dormante mais récupérable (énergie potentielle), liée à la position du système.

On appelle alors **énergie mécanique**  $E_m$  du système d'étude, la somme de son énergie cinétique et des énergies potentielles associées aux différentes forces conservatives qui s'exercent sur lui :

$$E_m = E_c + \sum E_{p, \vec{F}_c}$$

Le théorème de l'énergie cinétique adopte alors une autre forme, appelée **théorème de l'énergie mécanique**, et selon laquelle la variation d'énergie mécanique d'un système entre deux positions, évaluée par rapport à un référentiel galiléen, est égale à la somme des travaux des forces non conservatives s'exerçant sur ce système, entre ces deux positions :

$$\Delta E_m = \sum W_{C_i \rightarrow f}(\vec{F}_{nc})$$

Nous constatons ainsi que si un système est soumis **uniquement à des forces conservatives**, son **énergie mécanique est conservée** au cours du mouvement (d'où le terme de forces *conservatives*, autrement dit forces qui permettent la conservation, en l'occurrence de l'énergie mécanique). On dit dans ce cas que l'énergie mécanique est une **intégrale première du mouvement**, expression signifiant que cette grandeur est conservée au cours du mouvement.

Donc pour répondre à la question posée en premier titre de cette section : dans le cas où un système est soumis uniquement à des forces conservatives, oui, il existe une grandeur conservée, appelée énergie mécanique. Dans le cas contraire, la variation de cette énergie mécanique entre deux positions du système s'identifie au travail des forces non conservatives s'exerçant sur lui entre ces deux positions. Ce travail dépend du chemin suivi et doit donc prendre en compte non seulement les positions de départ et d'arrivée, mais également le chemin par lequel le système se rend de l'une à l'autre.

## — Parmi les forces d'usage courant, lesquelles ne sont pas conservatives ?

Les principales forces non conservatives que vous rencontrerez sont les forces de **frottements**. Il s'agit de forces dissipatives, dont le travail systématiquement négatif (elle s'opposent par nature au mouvement, autrement elles feraient très mal leur travail de frottement) entraîne une diminution non moins systématique de l'énergie mécanique du système : si ça frotte, le système ira forcément moins vite et/ou moins loin.

Vous rencontrerez diverses formes de forces de frottements : constantes, colinéaires à la vitesse du système, etc. L'essentiel est de bien comprendre que dans ce cas, vous devez considérer la circulation de la force sur **l'intégralité du chemin suivi** : si le système passe d'une position *A* à une position *B* puis revient à sa position de départ, on doit considérer la circulation sur l'ensemble de l'aller-retour, et il aura perdu de l'énergie mécanique sur chacune des deux phases. La force étant non conservative, il ne récupèrera pas au retour ce qu'il avait perdu à l'aller.

## — 11.3 Aspects énergétiques des phénomènes électriques

### — Qu'est-ce que l'énergie électrique ?

Dans un circuit électrique, des porteurs de charge libres se déplacent sous l'effet d'une différence de potentiel électrique  $U = \Delta V$ . Nous savons que la force électrique  $\vec{F}_e$  subie par une particule porteuse d'une charge électrique  $q$  en présence d'un champ électrique  $\vec{E}$  s'exprime :

$$\vec{F}_e = q \times \vec{E}$$

Le travail élémentaire de cette force sur un déplacement  $d\vec{l}$  s'en déduit comme :

$$\delta W_e = q \times \vec{E} \cdot d\vec{l}$$

Or on peut montrer que le champ électrique lui-même dérive du potentiel électrique, et l'on a ainsi :

$$\delta W_e = -q dV$$

La portion de travail peut ainsi s'identifier à la variation d'une énergie potentielle dont l'expression serait :

$$dE_{p,e} = -\delta W_e = +q dV \Rightarrow \Delta E_{p,e} = q \times \Delta V \Leftrightarrow \Delta E_{p,e} = q \times U$$

avec  $U = \Delta V$  la différence de potentiel électrique sous l'effet de laquelle circule la particule de charge  $q$ .

Ainsi les particules électriquement chargées acquièrent-elles, sous l'effet d'un champ électrique, une énergie couramment qualifiée d'**énergie électrique**. Cette énergie peut être restituée sous diverses formes au moyen de dispositifs dédiés. Le fonctionnement réciproque est du reste tout-à-fait envisageable :

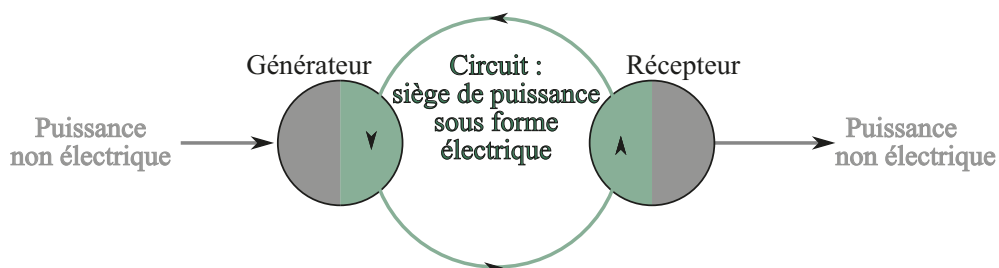
| $E_{p,e}$ restituée sous forme... | Mécanique                                        | Lumineuse                                | Thermique             | Chimique                       | Nucléaire                               |
|-----------------------------------|--------------------------------------------------|------------------------------------------|-----------------------|--------------------------------|-----------------------------------------|
| Dispositif associé                | Moteur                                           | Lampe                                    | Radiateur             | Électrolyseur                  | Accélérateur de particules              |
| Dispositif réciproque             | Alternateur (éolienne, centrale hydroélectrique) | Cellule photovoltaïque (panneau solaire) | Centrale géothermique | Centrale à combustible fossile | Centrale à fission, ou fusion contrôlée |

**Remarque :** dans les faits, les centrales géothermiques et à combustible (que ce dernier soit chimique ou nucléaire) sont toutes des centrales thermiques : une source de chaleur est utilisée pour provoquer l'évaporation de l'eau, la vapeur ainsi générée met en mouvement une turbine reliée à un alternateur, dont la mise en mouvement génère du courant électrique. On peut donc détailler les centrales à combustible comme une chambre de combustion convertissant de l'énergie (chimique ou nucléaire) en énergie thermique, suivie d'une turbine convertissant l'énergie thermique en énergie mécanique, et terminée par un alternateur convertissant l'énergie mécanique en énergie électrique.

## — Comment évaluer l'énergie électrique échangée par un dipôle ?

On peut considérer le circuit électrique comme un marché énergétique au sein duquel :

- certains dipôles fournissent de l'énergie au circuit sous forme électrique, à partir d'une source d'énergie extérieure non électrique : les **générateurs** ;
- d'autres convertissent l'énergie électrique portée par les porteurs de charge en déplacement, en une forme non électrique, utilisable par un(e) utilisateur/trice extérieur(e) : les **récepteurs**.



On peut alors montrer que la **puissance fournie au circuit** sous forme électrique par un dipôle orienté en **convention générateur** s'exprime en fonction de la tension  $u_d$  entre les bornes de ce dipôle, et de l'intensité  $i_d$  du courant électrique qui le traverse, comme :

$$P_{e,d} = u_d \times i_d$$

$P$  en W  
 $U$  en V  
 $I$  en A

**Remarque :** parlant depuis le départ d'énergie électrique, on peut se demander pourquoi nous passons ici à une rhétorique de puissance. Il se trouve que l'expression de la puissance électrique

échangée par un dipôle électrique avec le milieu extérieur répond à la forme ci-dessus, particulièrement simple et très générale. Si nous souhaitons passer à l'énergie il devient nécessaire d'intégrer cette relation, qui donnera des résultats très différents d'un dipôle à l'autre. Nous donnons donc ici une unique relation clef, que l'on déclinera (ou plus exactement que l'on intégrera) au gré des expressions de  $u_d$  et de  $i_d$ , lesquelles seront du reste liées entre elles par la caractéristique du dipôle envisagé.

Si le dipôle est orienté en **convention récepteur**, cette même expression désigne cette fois la puissance électrique **consommée** au niveau de ce dipôle, c'est-à-dire une **puissance retirée au circuit**.

|                                                         | Dipôle générateur                                         | Dipôle récepteur                                          |
|---------------------------------------------------------|-----------------------------------------------------------|-----------------------------------------------------------|
| Convention générateur<br>(puissance fournie au circuit) | $u_d, i_d$ de même signe<br>→ $P_{e,G}$ fournie $> 0$     | $u_d, i_d$ de signes opposés<br>→ $P_{e,R}$ fournie $< 0$ |
| Convention récepteur<br>(puissance retirée du circuit)  | $u_d, i_d$ de signes opposés<br>→ $P_{e,G}$ retirée $< 0$ | $u_d, i_d$ de même signe<br>→ $P_{e,R}$ retirée $> 0$     |

La convention d'orientation d'un dipôle va donc plus loin que le seul souci de travailler avec des grandeurs de même signe : elle permet de rendre compte du caractère consommateur ou fournisseur d'énergie d'un dipôle à l'égard du système que constitue le circuit. En effet, en comptant les énergies du point de vue du circuit, on comptabilise positivement l'énergie reçue de la part de ce dipôle si celui-ci est générateur, et négativement si elle est « reçue » de la part d'un récepteur (une réception négative ayant la même saveur que de se voir « offrir » une facture). Si à l'inverse on choisit d'effectuer un bilan énergétique du point de vue du milieu extérieur au circuit, le récepteur est une source d'énergie (qui prend au circuit pour donner au milieu extérieur), et l'on compte positivement ce qu'il retire au circuit. Le générateur en convention récepteur, lui, « fournit » au milieu extérieur une puissance négative, autrement dit il ne constitue pas une source mais une fuite d'énergie du point de vue extérieur.

## — Quel usage un dipôle électrique peut-il faire de l'énergie véhiculée par un courant électrique ?

La nature de la puissance non électrique que peut utiliser un générateur pour fournir de la puissance électrique, ou de celle sous laquelle un récepteur va restituer l'énergie électrique qu'il a reçue, dépend de sa structure (ou plus exactement, sa structure est conçue en sorte d'utiliser ou restituer un certain type d'énergie). Par un exemple un conducteur ohmique va par nature convertir de l'énergie électrique en énergie thermique (il agit comme un milieu soumettant les porteurs de charges à une force de frottements, menant à la conversion de l'énergie mécanique de ces porteurs en énergie d'agitation thermique du dipôle qu'ils traversent).

On peut donc imaginer, moyennant une certaine dose de créativité, des dipôles réalisant toutes les conversions énergétiques possibles et imaginables (et l'on ne s'en prive pas : c'est même l'une des bases du métier d'ingénieur(e)).

**Remarque :** les conversions sont souvent multiples. Par exemple un moteur électrique convertit certes une énergie électrique en énergie mécanique, mais il présentera inévitablement des pertes sous forme thermique.

Cependant du point de vue strictement électrocinétique, nous nous limiterons dans un premier temps à quelques cas simples :

- **le conducteur ohmique** : nous savons que dans ce cas, la tension et l'intensité sont liées par la loi d'Ohm :

$$u_R = R \times i_R \Rightarrow P_{e,R} = R \times i_R^2 = \frac{u_R^2}{R}$$

Donc si nous connaissons l'expression  $i_R(t)$  de l'intensité du courant à travers un conducteur ohmique (comme dans le cas d'un régime transitoire de circuit RC, par exemple), nous connaissons également la puissance électrique instantanée que dissipe ce conducteur ohmique par unité de temps :

$$P_{e,R} = \frac{\delta Q_R}{dt} \Leftrightarrow \delta Q_R = P_{e,R} \times dt = R \times i_R^2(t) \times dt$$

où  $\delta Q_R$  représente l'énergie électrique dissipée sous forme thermique par le conducteur ohmique en une durée  $dt$ .

Nous pouvons ainsi en déduire l'énergie thermique dissipée en une certaine durée comme :

$$Q_R = \int_{t_1}^{t_2} R i_R^2(t) dt = \int_{t_1}^{t_2} \frac{u_R^2(t)}{R} dt$$

Cet effet dissipatif constitue ce que l'on appelle l'**effet Joule** ;

- **le condensateur** : dans ce cas, la relation caractéristique nous donne :

$$i_C = C \frac{du_C}{dt} \Rightarrow P_{e,C} = C \times u_C \times \frac{du_C}{dt} = \frac{d}{dt} \left[ \frac{1}{2} C u_C^2(t) \right] = \frac{d}{dt} \left[ \frac{1}{2} \frac{q_C^2(t)}{C} \right]$$

La puissance apparaît donc ici directement comme la dérivée d'une fonction par rapport au temps. Il s'ensuit que l'énergie peut cette fois être considérée non comme échangée, mais comme stockée : tout comme certaines forces permettent d'emmagasiner l'énergie mécanique sous forme potentielle, le condensateur est un dipôle qui permet d'**emmagasiner de l'énergie électrique**, qu'il sera par la suite en mesure de restituer. C'est en effet ce qui apparaît lorsque nous enchaînons réponse indiciaire et réponse libre de la tension aux bornes du condensateur : là où le conducteur ohmique se contente de dissiper l'énergie électrique (on n'a jamais vu un tel dipôle se refroidir pour faire circuler un courant électrique), le condensateur la stocke.

Nous pouvons donc écrire qu'un condensateur aux bornes duquel la tension a pour valeur  $u_C$  est porteur d'une forme d'énergie potentielle électrique dont l'expression serait :

$$E_{e,C} = \int P_{e,C} \times dt = \frac{1}{2} C u_C^2(t)$$

en fixant la constante d'intégration à 0, nous affectons une valeur nulle à l'énergie d'un condensateur, lorsque la tension à ses bornes (ou la charge portée par ses armatures) est nulle, ce qui semble un choix plutôt parlant (les perspectives d'extraction d'énergie à partir d'un condensateur déchargé sont assez limitées) ;

- citons enfin, puisque nous l'avons abordée au chapitre d'électrocinétique, **la bobine**. Vous n'avez *a priori* pas vu ce dipôle au lycée, mais vous le verrez rapidement en CPGE, associé avec condensateur et/ou conducteur ohmique (cf. exercice « Vers la Prépa » du chapitre d'électrocinétique), et il constitue un intéressant pendant au condensateur. En effet ce dernier emmagasine de l'énergie sous forme électrique, la bobine quant à elle est également capable de l'emmagasiner mais vous verrez que ce sera cette fois sous forme magnétique.

Sa relation caractéristique est, rappelons-le :

$$i_L = L \times \frac{di_L}{dt}$$

où le coefficient de proportionnalité  $L$  est appelé coefficient d'auto-inductance (souvent abrégé *inductance* tout court), et a pour USI le henry (symbole : H).

Nous aurons donc dans ce cas :

$$P_{e,L} = L \times \frac{di_L}{dt} \times i_L = L \times i_L \times \frac{di_L}{dt} = \frac{d}{dt} \left[ \frac{1}{2} Li_L^2(t) \right]$$

On trouve une nouvelle fois une puissance s'exprimant comme la dérivée d'une fonction : l'énergie n'est donc pas dissipée, mais peut bien être emmagasinée par une bobine. Pour un courant d'intensité  $i_L(t)$ , cette énergie a alors pour expression :

$$E_{e,L} = \int P_{e,L} \times dt = \frac{1}{2} Li_L^2(t)$$

## — Un dipôle électrique se résume-t-il toujours à un composant simple ?

Vous travaillerez souvent sur des modèles de dipôles très idéalisés. En réalité, un dipôle ne répond jamais parfaitement à ces modèles. Un condensateur chargé abandonné à lui-même dans un circuit ouvert, par exemple, va inévitablement se décharger à mesure que le temps passe. On peut rendre compte de cette fuite de courant en associant au condensateur un conducteur ohmique en parallèle, ce qui amène à travailler sur un dipôle de type RC (mais en parallèle ; celui que nous avons traité en électrocinétique reposait sur une association en série). Vous irez donc progressivement vers des modèles plus complexes, qui vous permettront de nuancer la représentation de la réalité.

## — Comment le rendement d'un dipôle électrique est-il défini ?

La question peut de prime abord sembler absurde : l'énergie étant par nature conservative, un dipôle va nécessairement puiser autant d'énergie d'un côté qu'il va en libérer de l'autre. En réalité le rendement est moins une question de physique fondamentale que de physique appliquée, et fait intervenir la notion de **puissance utile**, c'est-à-dire servant un certain objectif.

On appelle ainsi **rendement**  $\eta$  d'un dipôle électrique, le rapport de la puissance utile  $P_{\text{utile}}$  qu'il génère, rapportée à la puissance totale  $P_{\text{conso}}$  qu'il consomme :

$$\eta = \frac{P_{\text{utile}}}{P_{\text{conso}}}$$

Il est donc capital sur des exercices faisant intervenir le rendement, de bien cadrer :

- quelles sont les puissances consommées et les puissances fournies par le dipôle ;
- parmi les puissances fournies, la(les)quelle(s) nous intéressent.

**Remarque :** la notion de rendement énergétique va bien au-delà du seul cadre électrocinétique, et peut être utilisée pour tout système convertissant une forme d'énergie en une autre. Un moteur à explosion, par exemple, convertit l'énergie chimique du combustible qu'il brûle, en énergie mécanique. Les notions d'énergies (ou de puissances) utile et consommée sont donc parfaitement transposables.

## 11.4 Aspects énergétiques des phénomènes thermodynamiques

### — Qu'est-ce que la thermodynamique ?

On appelle thermodynamique, littéralement, l'étude des transferts thermiques qui peuvent affecter un système donné. Cependant la limite n'est pas simple à établir, tant les idées qui s'y trouvent traitées peuvent également porter sur le travail des forces de pression, l'énergie des liaisons de covalence assurant la cohésion des molécules, voire celle des interactions nucléaires fortes assurant celle des noyaux...

Dans le cas présent, nous nous limiterons à une introduction et considérerons qu'elle concerne les échanges énergétiques entre systèmes macroscopiquement au repos.

L'étude d'objets qui se contentent de ne pas bouger peut sembler sans intérêt, mais ce serait ignorer des phénomènes aussi variés que les changements de température, les changements d'état physique, les transformations chimiques, etc. qui représentent autant de situation où un objet peut, même en restant parfaitement immobile du point de vue macroscopique, être le siège d'échanges énergétiques conséquents.

Souvenons-nous en effet des ordres de grandeur vus en début de chapitre : le fait qu'une même dose d'énergie se manifeste de manière beaucoup plus spectaculaire sous forme mécanique que sous forme thermodynamique, s'il attire moins l'attention, ne doit pas vous faire perdre de vue qu'il témoigne en même temps de ce que les énergies engagées par unité de masse sont beaucoup plus importantes en thermodynamique qu'en mécanique.

Il y a donc certainement beaucoup de choses à faire de l'énergie contenue dans la matière, moins du fait de son mouvement d'ensemble, que du fait de l'énergie individuelle dont est dotée chaque molécule, voire chaque atome, voire (soyons fous) de chaque noyau constituant un échantillon de matière.

### — Qu'est-ce que l'énergie interne d'un système ?

L'énergie interne (usuellement notée  $U$ ) d'un système est ce qu'il reste d'énergie à ce système une fois exclue son énergie mécanique macroscopique. En d'autres termes, si l'on note  $E_{\text{tot}}$  l'énergie totale dont est doté un système, on peut écrire :

$$E_{\text{tot}} = E_{\text{m}} + U \quad \Leftrightarrow \quad \Delta E_{\text{m}} + \Delta U = \Delta E_{\text{tot}}$$

Or en postulant que l'énergie est une grandeur conservée (point dont la contestation nécessiterait la remise à plat de toute la physique contemporaine, dont les modèles donnent tout de même à ce jour d'assez bon résultats), on peut affirmer que l'énergie totale d'un système peut varier uniquement si ce système a procédé à des échanges avec le milieu extérieur. On obtient ainsi :

$$\Delta E_{\text{tot}} = \sum E_{\text{éch}} \quad \Leftrightarrow \quad \Delta E_{\text{m}} + \Delta U = \sum E_{\text{éch}}$$

Dans l'hypothèse d'un corps macroscopiquement au repos, l'énergie mécanique ne varie pas et nous pouvons écrire :

$$\Delta E_m = 0 \Leftrightarrow \Delta U = \sum E_{\text{éch}}$$

Cette dernière égalité traduit donc la conservation de l'énergie pour un corps au repos. Vous en verrez une version plus détaillée en CPGE, mais devez pour le moment retenir qu'elle formalise un postulat de conservation de l'énergie connu sous le nom de **premier principe de la thermodynamique**.

**Remarque :** qui dit premier sous-entend qu'il existe au moins un deuxième. Le premier principe postule la conservation de l'énergie mais ne dit rien quant au sens dans lequel se produisent ses transferts. Par exemple imaginer, comme nous l'évoquions plus haut, un conducteur ohmique se refroidissant pour générer du courant électrique, ne viole pas en soi la conservativité de l'énergie : cela envisage seulement un transfert d'énergie dans un sens contraire à tout ce que nous indique l'expérience. La question du sens du transfert pose également celle du sens de la flèche du temps. La grandeur fixant ce sens est appelée entropie, et c'est sur elle que porte ce deuxième principe ; vous découvrirez l'une et l'autre en CPGE.

## — Comment l'énergie interne d'un système varie-t-elle lorsque sa température change ?

La température d'un corps est une grandeur macroscopique quantifiant le niveau d'agitation thermique des atomes qui le constituent. En première approximation, on peut considérer que la variation élémentaire  $(dU)_{\text{th}}$  d'énergie interne d'un corps consécutive à une variation de température  $dT$  est directement proportionnelle à celle-ci. On appelle alors **capacité calorifique** du corps considéré le coefficient de proportionnalité  $C$  tel que :

$$(dU)_{\text{th}} = C dT$$

L'USI de  $C$  est donc le joule (de variation d'énergie interne) par kelvin (de variation de température) :  $\text{J.K}^{-1}$ .

Pour un corps de capacité calorifique constante, on peut facilement intégrer la relation précédente en  $\Delta U = C \Delta T$ . La chose se complique pour des corps où  $C$  dépend de la température, puisqu'il faut alors intégrer  $dU = C(T)dT$  sur le domaine de variation de température considéré.

**Remarque :** on peut aussi donner la capacité calorifique massique du corps, exprimée en joule par kelvin et par kilogramme, égale au rapport de la capacité calorifique du corps à sa masse.

Attention, la capacité calorifique massique d'un corps dépend non seulement de sa nature chimique, mais également de son état. Celle de l'eau liquide, par exemple, est de  $C_e = 4,19 \text{ kJ.kg.K}^{-1}$ , tandis que celle de la glace a pour valeur  $C_g = 2,06 \text{ kJ.kg.K}^{-1}$ , soit environ deux fois moins. Par conséquent si l'on étudie une transformation au cours de laquelle se produit un changement d'état, on doit impérativement prendre en compte l'enchaînement :

- variation de température dans le premier état jusqu'à la température de changement d'état, avec la capacité calorifique du corps dans ce premier état ;
- changement d'état (cf. pages suivantes) ;
- variation de température dans le second état à partir de la température de changement d'état, avec la capacité calorifique du corps dans ce second état.

## — Comment mesure-t-on la capacité calorifique d'un corps ?

On utilise les techniques de la **calorimétrie**, et plus particulièrement la **méthode des mélanges**. Cette méthode consiste à mettre en contact, dans une enceinte adiabatique appelée **calorimètre**, deux corps ou davantage, initialement portés à des températures différentes. Les transferts thermiques ne peuvent se faire qu'entre les corps enfermés dans le calorimètre et ne s'interrompent qu'une fois l'équilibre thermique atteint, c'est-à-dire lorsque tous les corps ont la même température. La connaissance des températures initiales et la mesure de la température d'équilibre permettent de déterminer la capacité calorifique recherchée.

## — Comment l'énergie interne d'un système varie-t-elle lorsqu'il change d'état physique ?

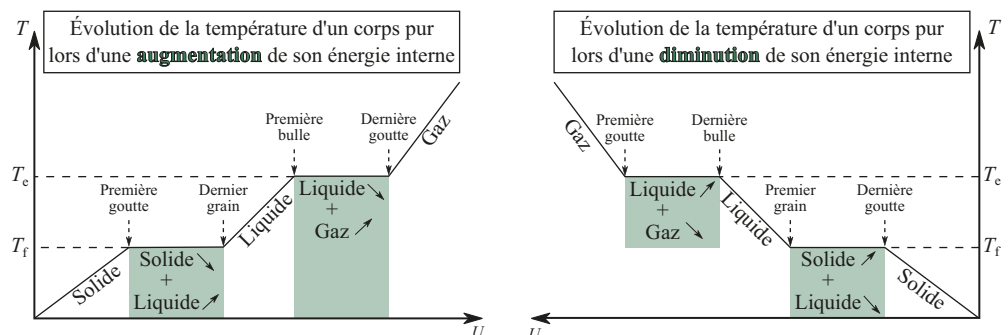
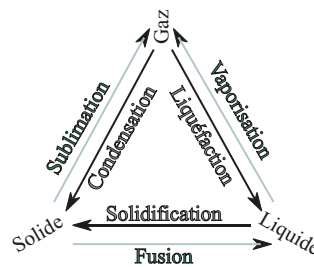
Nous avons vu dans la première partie de cet ouvrage qu'un corps pur pouvait exister sous 3 états, distinguables les uns des autres notamment par leurs propriétés macroscopiques de forme et de volume.

Lorsque l'on soumet un corps à un transfert thermique, on observe dans un premier temps une variation de sa température, témoignant de celle de son énergie interne. Mais au bout d'un moment, deux phénomènes se produisent simultanément :

- la température cesse de croître malgré le transfert thermique ;
- l'état physique du corps commence à changer.

On appelle alors **changement d'état** cette situation transitoire, où le corps existe dans plusieurs états simultanément, et **température de changement d'état** la température à laquelle cette coexistence a lieu.

Les changements d'état répondent à une nomenclature précise (cf. schéma ci-dessous), et les températures de changements d'état dépendent de la nature chimique du corps pur considéré, dont elles constituent une caractéristique couramment utilisée pour la reconnaissance d'espèces en chimie.



Durant le changement d'état d'un corps pur, celui-ci est présent **sous deux phases** (c'est-à-dire dans deux états physiques) à la fois. Le fait d'augmenter le nombre de phases simultanément présentes d'un même corps a pour effet de diminuer la **variance** du système, c'est-à-dire le nombre de paramètres physiques dont on peut librement fixer la valeur.

En pratique, pour un corps pur sous une seule phase, on peut, dans le domaine d'existence de la phase, choisir à la fois la pression et la température auxquelles est soumis ce corps. Ce n'est plus possible pour un corps pur sous deux phases : sitôt qu'un paramètre ( $P$  ou  $T$ ) est fixé, la nature ajuste l'autre automatiquement. Ainsi par exemple la température d'ébullition de l'eau (où celle-ci existe sous deux phases à la fois, donc) est-elle entièrement déterminée dès lors que la pression est fixée. C'est ce qui fait qu'elle bout à 100 °C au niveau de la mer mais plutôt vers 70 °C à 2 000 m d'altitude, posant des problèmes gastronomiques bien connus des alpinistes.

**Remarque :** on peut faire coexister les trois phases d'un corps pur, dans un état appelé **point triple** de ce corps pur. Au point triple, la variance est nulle et l'on ne peut alors plus choisir aucun paramètre : pression et température sont toutes les deux fixées par la nature. Ainsi par exemple l'échelle de température Celsius est-elle définie par la température du point triple de l'eau (273,16 K, et accessoirement une pression de 611 Pa), considérée comme valant 0,01 °C, à quoi l'on ajoute qu'une variation de 1 K égale exactement une variation de 1 °C. Ce faisant, on retrouve bien que 0 °C = 273,15 K.

Durant cette transformation, l'énergie interne du système augmente (respectivement diminue) bien, mais cette augmentation/diminution n'est plus utilisée pour augmenter/diminuer l'énergie d'agitation thermique des molécules : elle entraîne la rupture/création de liaisons intermoléculaires (type Van der Waals), entraînant leur dissociation/association et amenant l'échantillon de matière dans un état physique moins/plus ordonné qu'il ne l'était.

La variation d'énergie interne  $(dU)_{\text{CdE}}$  qui accompagne le changement d'état d'un échantillon de corps pur est d'autant plus importante que celle de la masse  $dm_{\text{ord-}}$  de corps pur dans le moins ordonné des deux états l'est également. En modélisant la relation par une loi de proportionnalité, on appelle alors **chaleur latente de changement d'état** du corps pur considéré, le coefficient de proportionnalité  $L_{\text{CdE}} > 0$  tel que :

$$(dU)_{\text{CdE}} = L_{\text{CdE}} \times dm_{\text{ord-}}$$

On peut si l'on préfère travailler avec la masse  $\delta m_{\text{CdE}}$  d'échantillon qui a changé d'état :

| Sens du changement d'état                                                                                                                     | Lien entre $dm_{\text{ord-}}$ et $\delta m_{\text{CdE}}$                                                       | Lien avec l'énergie interne                                       |
|-----------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------|
| Plus ordonné → Moins ordonné<br>(fusion, vaporisation, sublimation)<br>La masse de corps dans l'état le moins ordonné <b>augmente</b>         | $m_{\text{ord-,f}} = m_{\text{ord-,i}} + \delta m_{\text{CdE}}$<br>$dm_{\text{ord-}} = \delta m_{\text{CdE}}$  | $(\Delta U)_{\text{CdE}} = L_{\text{CdE}} \times m_{\text{CdE}}$  |
| Moins ordonné → Plus ordonné<br>(solidification, liquéfaction, condensation)<br>La masse de corps dans l'état le moins ordonné <b>diminue</b> | $m_{\text{ord-,f}} = m_{\text{ord-,i}} - \delta m_{\text{CdE}}$<br>$dm_{\text{ord-}} = -\delta m_{\text{CdE}}$ | $(\Delta U)_{\text{CdE}} = -L_{\text{CdE}} \times m_{\text{CdE}}$ |

Bilan : la variation d'énergie interne  $(\Delta U)_{\text{CdE}}$  d'un système constitué d'un corps pur, consécutive à un changement d'état d'une masse  $m_{\text{CdE}}$  de ce corps pur, a pour expression :

$$(\Delta U)_{\text{CdE}} = \pm L_{\text{CdE}} \times m_{\text{CdE}}$$

avec :

- le signe « + » (garantissant  $(\Delta U)_{\text{CdE}} > 0$ ) si le désordre **augmente** ;
- le signe « - » (garantissant  $(\Delta U)_{\text{CdE}} < 0$ ) si le désordre **diminue**.

**Remarque :** vous verrez en CPGE que cette notion de « désordre » (ou plus exactement de manque d'information sur le détail de l'état microscopique des éléments constitutifs du système) est au cœur du concept d'entropie que nous évoquions concernant le deuxième principe de la thermodynamique.

## — Comment l'énergie interne d'un système varie-t-elle lorsqu'il subit une transformation chimique ?

Formellement, on peut représenter une transformation chimique comme la succession de deux étapes :

- **la destruction des liaisons de covalence des molécules de réactif :** en notant  $E_{\text{l,mol}}$  les énergies molaires de liaison associées aux différentes liaisons présentes dans une molécule, cette opération va donc consister à fournir aux molécules de réactif une énergie dépassant celle dont elles avaient abaissé leur énergie interne lorsqu'elles ont établi ces liaisons de covalence à partir d'atomes dispersés. La variation  $(\Delta U)_{\text{ch,R}}$  d'énergie interne s'écrira donc :

$$(\Delta U)_{\text{ch,R}} = \sum_{\text{Réactifs}} n_{\text{l,R}} \times E_{\text{l,mol}} > 0$$

où  $n_{\text{l,R}}$  représente la quantité de matière de liaisons d'un certain type à briser, égale à la quantité de matière consommée du réactif contenant cette liaison, que multiplie le nombre de liaisons de ce type dans l'espèce en question ;

- **l'établissement des liaisons de covalence des molécules de produit :** c'est cette fois le système qui va libérer de l'énergie, suite à la stabilisation résultant de la formation de nouvelles liaisons de covalence (les molécules n'existent, rappelons-le, que parce qu'elles permettent la stabilisation des atomes qu'elles engagent).

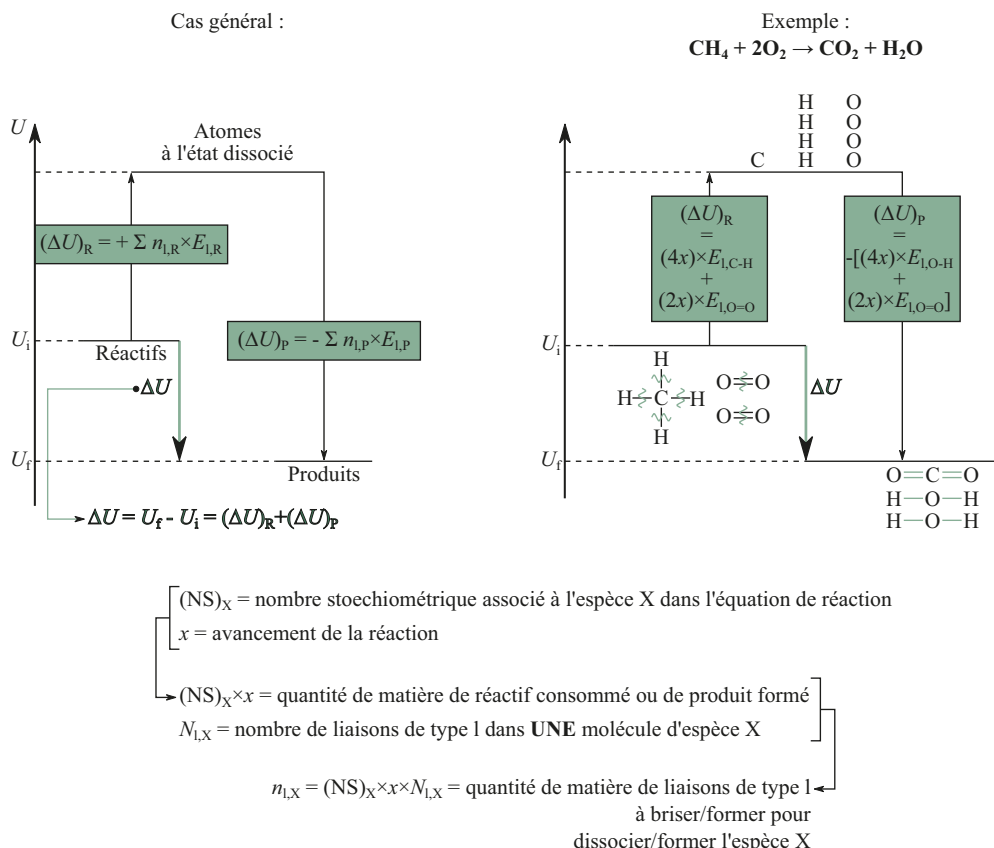
Le système va donc cette fois perdre de l'énergie, et la variation  $(\Delta U)_{\text{ch,P}}$  d'énergie interne qui va s'ensuivre s'écrira cette fois :

$$(\Delta U)_{\text{ch,P}} = - \sum_{\text{Produits}} n_{\text{l,P}} \times E_{\text{l,mol}} < 0$$

où  $n_{\text{l,P}}$  représente la quantité de matière de liaisons d'un certain type à former, égale à la quantité de matière du produit contenant cette liaison, que multiplie le nombre de liaisons de ce type dans l'espèce en question.

En fin de compte, la variation totale d'énergie interne au cours d'une transformation chimique pourra donc s'écrire :

$$(\Delta U)_{\text{ch}} = (\Delta U)_{\text{ch,R}} + (\Delta U)_{\text{ch,P}} \Leftrightarrow (\Delta U)_{\text{ch}} = \sum_{\text{Réactifs}} n_{\text{l,R}} \times E_{\text{l,mol}} - \sum_{\text{Produits}} n_{\text{l,P}} \times E_{\text{l,mol}}$$



Nous savons par ailleurs que les quantités de matière consommées en réactifs et celles formées en produits sont toutes proportionnelles à l'**avancement**  $x$  de la réaction. On vous donnera donc parfois une grandeur appelée **énergie interne molaire de réaction**  $\Delta_r U$ , définie comme :

$$\Delta_r U = \frac{(dU)_{\text{ch}}}{dx}$$

Attention, d'habitude nous utilisons la notation «  $\Delta$  » pour signifier une simple variation, et  $\Delta U$  par exemple est bien homogène à une énergie, en joules. L'opérateur «  $\Delta_r$  » que vous voyez ci-dessus symbolise une opération de dérivation par rapport à l'avancement. Aussi  $\Delta_r U$  ne s'exprimera-t-elle pas simplement en joules, mais en joules (d'énergie échangée par le système chimique avec le milieu extérieur) par mole (d'avancement de réaction).  $\Delta_r U$  constitue donc une version synthétique de la relation vue plus haut avec les énergies de liaisons, incluant directement les énergies molaires des différentes liaisons et nous livrant simplement le bilan de l'énergie échangée par la réaction considérée, par mole d'avancement.

On alors :

$$(\Delta U)_{\text{ch}} = \Delta_r U \times \Delta x$$

## Comment l'énergie interne d'un système varie-t-elle lorsqu'il subit une transformation nucléaire ?

Lors d'une transformation nucléaire, le fond est exactement le même : on peut considérer qu'une désintégration spontanée, une fission ou une fusion résulte d'une série de démantèlement de

noyaux et/ou association de nucléons. La première opération nécessite un apport d'énergie, la seconde en dégage. La différence est que cette fois, on ne vous donnera en général pas les énergies de liaison, ni même l'énergie dégagée par mole d'avancement, mais les masses des noyaux et particules de départ, et celles des noyaux et particules d'arrivée.

En effet nous savons que tout corps massif possède une **énergie de masse**, dont l'expression est sans doute à ce jour la plus célèbre de la physique :

$$E = mc^2$$

Lorsqu'un corps gagne ou perd de l'énergie, donc, techniquement il gagne ou perd de la masse en même temps. Cependant dans les phénomènes précédents, les énergies échangées par unité de masse de système étaient trop faibles pour être facilement mesurées, aussi nous sommes-nous cantonnés à des considérations strictement énergétiques.

Mais dans le cas de transformations nucléaires, les énergies engagées deviennent suffisamment élevées pour affecter sensiblement la masse des noyaux. On vous donnera donc souvent les masses des noyaux et particules de départ, sur la base desquelles vous pourrez calculer l'énergie de masse du système avant sa transformation :

$$E_i = \sum_{\text{N\&P initiaux}} m_{\text{N\&P}} c^2$$

Disposant des grandeurs analogues pour les noyaux et particules obtenus à l'arrivée, vous pourrez également exprimer, à l'état final :

$$E_f = \sum_{\text{N\&P finaux}} m_{\text{N\&P}} c^2$$

On peut alors déterminer la variation d'énergie interne consécutive à une transformation nucléaire simplement comme :

$$(\Delta U)_{\text{nucl}} = \left( \sum_{\text{N\&P finaux}} m_{\text{N\&P}} - \sum_{\text{N\&P initiaux}} m_{\text{N\&P}} \right) \times c^2$$

## — Comment traite-t-on un problème de thermodynamique ?

Tout comme en mécanique, le traitement rigoureux d'un problème de thermodynamique demande en premier lieu de bien cadrer ce à quoi l'on s'intéresse et que l'on appelle **système**. Il s'agit d'une portion d'univers, généralement matérialisée par un objet, et dont nous allons étudier les variations d'énergie dont il est le siège d'une part, et les échanges énergétiques auxquels il procède avec le milieu extérieur d'autre part.

Une fois le système bien défini, donc, la logique sera d'inventorier pour ce système, selon la situation décrite :

- tous les phénomènes susceptibles d'entraîner une **variation d'énergie interne** du système :
  - variation de température,
  - changement d'état physique,
  - transformation chimique,
  - transformation nucléaire ;
- les éventuels **échanges énergétiques** auxquels le système se livre avec le milieu extérieur (rayonnement, conduction, convection, inscription dans un circuit électrique, etc.).

Les variations répondent aux expressions que nous venons de voir dans les pages qui précèdent. Les éventuels échanges avec le milieu extérieur sont en général donnés sous forme de flux de puissance. Le rayonnement solaire reçu sur Terre, si par exemple il intervient dans la situation étudiée, est de l'ordre de  $700 \text{ W.m}^{-2}$  : un système s'y trouvant exposé et en mesure de l'absorber recevra donc  $700 \text{ J}$  d'énergie interne par mètre carré de surface exposée à ce rayonnement, et par seconde de temps écoulée.

Une fois ceci fait, nous n'aurons plus qu'à reporter ces différents termes dans l'équation de conservation de l'énergie. Nous obtiendrons alors une égalité engageant les différents paramètres caractéristiques du système (température, masse, composition chimique ou nucléaire, etc.). Parmi ceux-ci tous seront fixés sauf un, qu'il nous faudra alors isoler et calculer à partir des autres.

**Remarque :** attention, dans le cas d'un flux de puissance échangé avec l'extérieur, à ne pas oublier s'il y a lieu le **rendement** de l'instrument captant cette puissance. Le rendement d'un panneau solaire, par exemple, est seulement de 15 à 20 %, de quoi dévier sérieusement des résultats attendus.

## 11.5 Un mot concernant les gaz

L'étude de la thermodynamique en CPGE commence en général par les gaz, que leur caractère compressible rend particulièrement intéressants. Nous terminerons donc ce chapitre en rappelant et complétant quelques bases de lycée sur le sujet.

### Quelles sont les caractéristiques d'un corps à l'état gazeux ?

Un corps à l'état gazeux est **expansible** : il occupe tout le volume qui lui est offert et n'a donc pas de volume propre, contrairement aux corps à l'état liquide ou solide. Par ailleurs, les gaz sont **compressibles** dans des proportions bien plus importantes que les liquides et les solides. Enfin, la **masse volumique** d'un corps à l'état gazeux est généralement nettement plus faible que celle d'un corps à l'état liquide ou solide.

### Quels sont les paramètres pertinents pour mener l'étude d'un corps à l'état gazeux ?

Outre la quantité de matière  $n$  de molécules de gaz, paramètre pertinent que le système soit gazeux ou non, on peut citer trois paramètres fondamentaux :

- la **volume**  $V$  occupé par le gaz, dont l'unité SI est le mètre cube ( $\text{m}^3$ ) ;
- la **température**  $T$  à laquelle est porté le gaz, dont l'unité SI est le kelvin (K) ;
- la **pression**  $P$  exercée par le gaz, dont l'unité SI est le pascal (Pa).

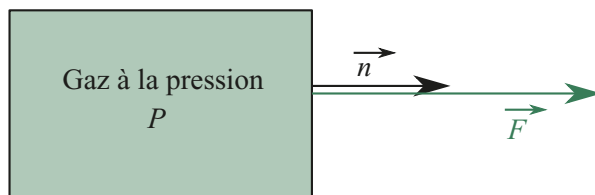
Les trois grandeurs précédentes seront souvent fournies en unité non-SI (par exemple le volume en L, la pression en bar et la température en  $^{\circ}\text{C}$ ) ; il faudra alors les convertir avant de les utiliser dans les relations. Rappelons les unités non-SI les plus couramment utilisées :

- pour le volume,  $1,0 \text{ L} = 1,0 \cdot 10^{-3} \text{ m}^3$  ;
- pour la température,  $T_{(\text{en K})} = T_{(\text{en } ^{\circ}\text{C})} + 273,15$  ;
- pour la pression,  $1,0 \text{ bar} = 1,0 \cdot 10^5 \text{ Pa}$ .

## — Que représente la pression d'un gaz ?

Les molécules du corps à l'état gazeux sont en mouvement désordonné et rapide. Elles heurtent la paroi du récipient contenant le gaz et rebondissent. Chacun de ces chocs microscopiques exerce brièvement une force  $\vec{F}$  sur la paroi. Ces chocs étant très nombreux, et les directions suivant lesquelles les molécules heurtent la paroi très variées, ce phénomène se manifeste à l'échelle macroscopique par l'apparition d'une **force de pression**  $\vec{F}$  :

- normale à la paroi ;
- dirigée du gaz vers la paroi ;
- de valeur proportionnelle à la surface de la paroi.



|                                       |     |                   |
|---------------------------------------|-----|-------------------|
| $\vec{F} = P \times S \times \vec{n}$ | $F$ | en N              |
|                                       | $P$ | en Pa             |
|                                       | $S$ | en m <sup>2</sup> |

Enfin, la pression au sein d'un gaz représente un champ scalaire, en ce sens qu'elle dépend de la position de l'espace dans laquelle on se place. Si la pression est la même en tout point du volume occupé par le gaz, elle est dite **uniforme**.

**Remarque :** on définit de manière plus générale la pression exercée par un fluide (les états fluides regroupant les états gazeux et liquide) ; c'est pourquoi, dans les exercices qui suivent, il sera parfois fait mention de la pression régnant à une certaine profondeur dans l'eau...

## — Que représente la température d'un gaz ?

La première notion que l'on peut en avoir est liée aux perceptions physiologiques du chaud et du froid, fort peu précises. Heureusement, il existe de nombreuses propriétés physiques mesurables, par exemple la dilatation d'un liquide ou son point d'ébullition, qui permettent de réaliser des **échelles de température** ou **thermomètres**.

La **température** d'un corps est donc un paramètre physique permettant de rendre compte, à l'échelle macroscopique, de l'état d'agitation thermique des constituants microscopiques de ce corps. Si ces constituants microscopiques ont le même niveau d'agitation thermique pour deux corps distincts, ceux-ci ont même température. Lors de la mise en contact de ces corps, les constituants de l'un pas plus que ceux de l'autre ne recevront davantage de chocs qu'ils n'en recevaient déjà, et les températures ne varieront pas : on dit de chacun des corps qu'il se trouve à l'**équilibre thermique** avec l'autre. Les thermomètres sont fondés sur le principe suivant, appelé **principe zéro** de la thermodynamique : deux corps en équilibre thermique avec un même troisième sont en équilibre thermique entre eux.

**Remarque :** il faut rester prudent avec le concept de température et veiller notamment à ne pas le confondre avec celui de chaleur ou transfert thermique.

## — Sur quelles hypothèses le modèle du gaz parfait s'appuie-t-il ?

Le **modèle du gaz parfait** est un modèle simple décrivant en bonne approximation le comportement de la plupart des gaz. Dans le cadre de ce modèle, on suppose les molécules de gaz ponctuelles et sans autre interaction entre elles que leurs collisions les unes avec les autres (aucune interaction à distance, en particulier). Les hypothèses qui sous-tendent ce modèle sont alors les suivantes :

- on néglige le volume propre des molécules devant le volume occupé par le gaz ;
- les molécules de gaz ont, entre deux chocs, des mouvements rectilignes uniformes ;
- l'énergie microscopique des molécules est purement cinétique, donc ne dépendra que de la température.

**Remarque :** ce modèle correspond finalement au comportement d'un gaz dans la limite d'une pression nulle, pour assurer l'absence d'interactions entre molécules de gaz. Même s'il apparaît assez restrictif à première vue, un grand nombre de gaz réels sont très bien décrits par ce modèle, et ce dans une large gamme de pression.

## — Comment l'équation d'état du gaz parfait s'écrit-elle ?

Historiquement, on a constaté expérimentalement qu'aux faibles pressions (inférieures ou égales à un bar) les gaz (à quantité de matière fixée) satisfont à un certain nombre de lois :

- La **loi de Boyle-Mariotte** : à température constante, le produit de la pression exercée par un gaz par le volume qu'il occupe est constant.
- La **loi de Gay-Lussac** : à volume et quantité de matière identiques, le rapport de la pression exercée par un gaz à sa température absolue est indépendant de la nature de ce gaz.
- La **loi d'Avogadro-Ampère** : à température constante, le produit de la pression d'un gaz par son volume est proportionnel au nombre de molécules de gaz, et indépendant de sa nature.

Dans le cadre du modèle du gaz parfait, on rassemble ces trois lois en une seule : **l'équation d'état du gaz parfait** qui fait intervenir une constante  $R$ , appelée constante du gaz parfait.

|            |     |                                      |
|------------|-----|--------------------------------------|
| $PV = nRT$ | $P$ | en Pa                                |
|            | $V$ | en $\text{m}^3$                      |
|            | $n$ | en mol                               |
|            | $R$ | en $\text{J.K}^{-1}.\text{mol}^{-1}$ |
|            | $T$ | en K                                 |

## — De quel(s) paramètre(s) le volume molaire du gaz parfait dépend-il ?

D'après la loi d'Avogadro-Ampère, une mole de gaz occupe, quelle que soit sa nature chimique, toujours le même volume dans des conditions données de température et de pression. Le **volume molaire du gaz parfait**  $V_m$  n'est donc fonction que de la pression et de la température du gaz.

Par exemple, dans les conditions usuelles de température et de pression (sous 1,013 bar et à 20 °C), on calcule  $V_m = \frac{RT}{P} = 24,0 \text{ L.mol}^{-1}$ .

## Halte aux idées reçues

Expliquer, pour chacune des affirmations suivantes, pourquoi elle est fausse :

1. Le travail d'une force entre deux positions  $A$  et  $B$  de son point d'application est égal au produit scalaire du vecteur représentatif de cette force, par le vecteur  $\overrightarrow{AB}$ .
2. L'énergie potentielle d'un système à l'équilibre est forcément nulle.

3. Le travail d'une force de frottement est toujours strictement négatif.

4. L'énergie nucléaire peut être convertie en énergie électrique, mais pas l'inverse.

5. La masse d'un système est indépendante de sa température.

## Du Tac au Tac

(Exercices sans calculatrice)

1. Une skieuse de masse  $m = 80$  kg (avec son équipement) descend une pente longue de  $l = 500$  m et inclinée d'un angle  $\theta = 6,0^\circ$  par rapport à l'horizontale. Partie du sommet  $A$  sans vitesse initiale, elle arrive au bas de la piste en  $B$ , avec une vitesse  $v_B = 20$  m.s<sup>-1</sup>.

Déterminer quel pourcentage de son énergie mécanique a été dissipé en frottements entre le haut et le bas de la piste. On fera l'approximation que  $1 \text{ rad} = 60^\circ$ , et l'on prendra  $g = 1.10^1$  N.

2. On considère une voiture, lancée à une vitesse de valeur  $v = 72$  km.h<sup>-1</sup> sur une route droite et horizontale. La masse totale du véhicule et de ses passagers est de  $M = 1\,000$  kg, incluant notamment une masse  $m = 10$  kg de pièces frottant les unes sur les autres pour assurer le freinage du véhicule (disques et plaquettes de frein).

Sachant que la capacité thermique massique moyenne des pièces de freinage a pour valeur  $c = 800$  J.kg<sup>-1</sup>.K<sup>-1</sup>, déterminer leur augmentation de température consécutivement à un freinage complet du véhicule (sans dérapage).

3. Déterminer de quelle hauteur on pourrait élever une voiture, de masse  $M = 660$  kg, avec l'énergie nécessaire à la fusion d'un glaçon de masse  $m = 10$  g.

On donne la chaleur latente de fusion de la glace :  $L_f = 3,3.10^2$  kJ.kg<sup>-1</sup>.

4. Un barrage hydroélectrique possède un débit volumique  $D_v = 8,0$  m<sup>3</sup>.s<sup>-1</sup>, qui alimente une turbine située  $h = 50$  m plus bas (on négligera la vitesse initiale de l'eau).

Déterminer la puissance électrique qu'il délivre, en supposant un rendement  $\eta = 75\%$  et en supposant que l'eau ressort du barrage avec une vitesse pratiquement nulle.

5. Les stimulateurs cardiaques, ou *pacemakers*, servent à prévenir les malaises que peuvent provoquer certaines arythmies cardiaques. Ce sont de petits dispositifs de quelques centimètres de taille, qui délivrent régulièrement des impulsions électriques à destination du cœur. Ils sont implantés chirurgicalement sous la peau et doivent donc être dotés d'une source d'énergie de très grande longévité, évitant au patient les interventions qu'entraînerait le besoin de changer la pile de l'appareil. Dans ce but, on utilise des piles au plutonium 238.

On donne les ordres de grandeur suivants :

- Activité molaire du  $^{238}\text{Pu}$  :  $\frac{A}{n} = 10^{14}$  Bq.mol<sup>-1</sup>
- Demi-vie du  $^{238}\text{Pu}$  : environ un siècle.
- Énergie du rayonnement délivré par la désintégration d'un noyau de  $^{238}\text{Pu}$  :  $E_\gamma = 10^{-12}$  J.
- Puissance électrique consommée par le stimulateur cardiaque :  $\mathcal{P} = 10^{-1}$  W.

Déterminer la masse de plutonium nécessaire au fonctionnement de cette pile, et justifier que cette puissance demeure sensiblement constante.

 Vers la prépa

On considère un satellite modélisé par un point matériel  $M$  de masse  $m$ , en orbite (pas forcément circulaire) à une distance  $r$  du centre de la Terre. On l'étudiera par rapport au référentiel géocentrique supposé galiléen.

**1.** Donner l'expression générique de son énergie mécanique en fonction de  $r$  et de la valeur  $v$  de sa vitesse.

**2.** Montrer, en vous fondant uniquement sur des considérations énergétiques, que l'orbite du satellite est circulaire si et seulement si le mouvement est uniforme.

**3.** On considère un satellite lancé depuis la surface de la Terre avec une vitesse  $v_1$ . Déterminer l'expression du carré de sa vitesse, en fonction de la distance  $r$  qu'il a atteinte.

Tracer la courbe représentative  $v^2(r)$ , et montrer que  $r$  ne peut excéder une certaine valeur dont on donnera l'expression, à moins que  $v_1$  n'atteigne une valeur particulière que l'on détaillera également.

**4.** On suppose cette fois que le satellite est en orbite circulaire autour de la Terre. Montrer que si le satellite subit une force de frottement, alors le rayon de son orbite va diminuer.

## Halte aux idées reçues

1. Le travail d'une force peut s'exprimer sous la forme proposée par l'énoncé, mais seulement sur un chemin en chacun des points duquel la force en question conserve les mêmes caractéristiques. Notons au passage que l'on parle parfois d'une force qui doit rester constante sur le chemin suivi, terme qui prête à confusion :

- La constance d'une grandeur physique désigne généralement, en physique, une indépendance de cette grandeur vis-à-vis du **temps**.
- Dans le cas où l'on veut parler d'une indépendance vis-à-vis de la position, on parle plutôt d'**uniformité** de la grandeur en question.

On peut ainsi parfaitement envisager un champ constant, qui ne soit pas uniforme : il suffit que ses irrégularités spatiales demeurent les mêmes en tout instant. C'est par exemple le cas du champ gravitationnel d'un astre, dont la direction dépend de la position autour de cet astre, et la valeur dépend de la distance au centre de cet astre.

Réciproquement, on peut envisager un champ uniforme mais non constant : par exemple le champ électrique entre deux plaques d'un condensateur, aux bornes duquel on ferait progressivement croître la tension électrique.

Dans le cas présent, il s'agit de dire que la force subie par le système présente les mêmes caractéristiques en tout point du chemin suivi par ce système. Un champ constant mais non uniforme ne suffit pas, non plus qu'un champ uniforme mais non constant (puisque dans ce cas la force subie en une position différerait de celle subie à la suivante). Cependant répétons-le : si l'expression est peut-être maladroite, elle est malgré tout consacrée par l'usage et l'on parle couramment de *force constante le long du chemin suivi*.

À partir de là, deux situations peuvent se présenter.

Si la force est effectivement constante sur le chemin suivi, on peut alors écrire :

$$W_A^B(\vec{F}) = \vec{F} \cdot \overline{AB} = \|\vec{F}\| \times AB \times \cos(\vec{F}, \overline{AB})$$

Dans le cas contraire, la relation ci-dessus reviendrait à calculer un travail moyen, tout en sachant que la force n'étant pas constante, elle ne serait plus nécessairement conservative, et qu'il serait alors peut-être nécessaire de préciser non seulement les positions entre lesquelles a été calculé ce travail, mais encore le chemin emprunté par le système pour se

rendre de la première à la seconde. Dans les faits, on procède comme à l'habitude dans ce genre de situation : on décompose l'intervalle sur lequel la grandeur considérée n'est pas constante en intervalles infinitésimaux sur chacun desquels on peut considérer qu'elle l'est. Cette approche un peu cavalière peut faire frémir, mais elle n'est que le reflet vulgarisé de tout l'arsenal développé par les Mathématiques sous le nom de calcul différentiel, et qui bénéficie de démonstrations parfaitement rigoureuses mais sortant totalement du cadre de cet ouvrage.

On va ainsi définir un travail élémentaire  $\delta W$ , sur un déplacement élémentaire  $d\vec{l}$ , selon la relation :

$$\delta W = \vec{F} \cdot d\vec{l}$$

Puis, arguant de ce que le chemin  $\mathcal{C}_A^B$  allant de A à B n'est jamais que la somme vectorielle des différents déplacements élémentaires réalisés pour se rendre de A à B, réaliser la sommation continue des travaux élémentaires en question, c'est-à-dire leur intégration :

$$W_{\mathcal{C}_A^B} = \int_{\mathcal{C}_A^B} \vec{F} \cdot d\vec{l}$$

On note au passage que si  $\vec{F}$  est constante le long du chemin, on peut alors écrire :

$$W_{\mathcal{C}_A^B} = \vec{F} \cdot \underbrace{\int_{\mathcal{C}_A^B} d\vec{l}}_{\overline{AB}}$$

On retrouve ainsi l'expression précédente dans l'hypothèse d'une force constante sur le chemin suivi.

Il est également intéressant de noter que si la force forme avec le déplacement élémentaire un angle constant (cas d'une force de frottement par exemple, systématiquement opposée au déplacement et formant donc avec  $d\vec{l}$  un angle de  $180^\circ$ ) et est de valeur constante, on a cette fois, si l'on note  $\mathcal{C}_A^B$  la longueur du chemin parcouru entre A et B :

$$W_{\mathcal{C}_A^B} = F \times \cos(\vec{F}, \overline{AB}) \times \mathcal{C}_A^B$$

Au final, il importe donc surtout de conserver à l'esprit que la relation simple vue au lycée pour calculer le travail d'une force entre deux positions de son point d'application possède un champ d'application très restreint, et ne doit pas être librement adaptée à une force non constante sur le chemin suivi, sans d'abondantes précautions.

2. On sait que les énergies potentielles en général ne sont définies qu'à une constante additive près, puisque seules les variations de ces énergies entre deux positions du système (par le biais du travail de la force qui leur est associée), sont accessibles à la mesure. Ainsi, si par exemple on fixe l'origine des énergies potentielles de pesanteur au niveau de la mer ( $E_{p,p}(z) = mgz$ ), et que l'on considère un système posé au sommet de l'Everest, on constate que l'énergie potentielle de pesanteur est loin d'être nulle ( $z = 8\,848$  m, tout de même), ce qui n'empêche nullement ce système d'être au repos.

Il est en revanche vrai que l'équilibre d'un système s'accompagne de la nullité d'un terme. En effet, on appelle position d'équilibre d'un système ponctuel par rapport à un référentiel donné, toute position de ce système en laquelle, abandonné à lui-même sans vitesse initiale, il persévère dans cet état d'immobilité. On comprend ainsi qu'une position d'équilibre d'un système considéré par rapport à un référentiel galiléen est simplement une position de ce système, en laquelle les forces qu'il subit se compensent entre elles. Or, comme nous l'avons vu à la question précédente, les forces conservatives dérivent d'une énergie potentielle (les non-conservatives travaillent uniquement si le système est en mouvement, donc pas lorsqu'il est au repos). Ainsi, une position d'équilibre est une position en laquelle la somme des dérivées par rapport aux coordonnées d'espace, des énergies potentielles emmagasinées par le système (donc la dérivée de l'énergie potentielle totale emmagasinée par un système) est nulle.

$$x_0 \text{ position d'équilibre} \Leftrightarrow \left( \frac{dE_p}{dx} \right)_{x=x_0} = 0$$

Une dernière manière de l'exprimer est alors de dire qu'une position d'un système constitue une position d'équilibre par rapport à un référentiel galiléen, si et seulement si elle correspond à un extremum (minimum ou maximum) d'énergie potentielle de ce système.

Un exemple simple est celui du pendule simple avec une tige rigide, qui possède en tout et pour tout deux positions d'équilibre :

- Celle où le pendule est pendant, donc possède une altitude et une énergie potentielle de pesanteur minimales.
- Celle où le pendule se tient à la verticale au-dessus de son point de fixation, et qui quant à elle correspond à une altitude et une énergie potentielle maximales.

On peut alors ajouter au critère d'équilibre un critère de stabilité :

- Une position d'équilibre est dite **stable** si le système tend à y revenir lorsqu'il s'en trouve écarté. On peut montrer que ceci correspond au cas où la **dérivée seconde** de l'énergie potentielle par rapport aux coordonnées d'espace est **positive**.
- Une position d'équilibre est dite **instable** si le système tend à s'en éloigner d'autant plus qu'il s'en trouve écarté. On peut montrer que ceci correspond au cas où la dérivée seconde de l'énergie potentielle par rapport aux coordonnées d'espace est **négative**.

3. (NB : les explications qui suivent font référence au schéma figurant à la page suivante). Une force de frottement peut ne pas travailler, et même servir au déplacement d'un système. Nous n'irons pas jusqu'à dire que son travail peut être positif, mais il peut être nul. On peut notamment citer le cas des forces de frottement s'exerçant sur un pneumatique de véhicule. En effet, si l'on considère par exemple une voiture reposant sur une plaque de verglas et cherchant à avancer, l'absence de forces de frottements fait qu'elle va patiner sur place. Le déplacement de cette voiture (mais également celui d'un vélo, d'une personne se déplaçant à pied...) repose sur la réaction du sol sur ses roues : lorsque la roue (ou le pied, pour la marche) entre en mouvement par rapport au véhicule dont elle est tributaire, elle accroche le sol et exerce donc une force sur ce dernier, qui tendrait à le faire filer vers l'arrière. Réciproquement le sol exerce une force de frottement sur la roue, qui pour sa part sera dirigée en sens inverse, vers l'avant du véhicule (troisième loi de Newton oblige). La roue ne pouvant patiner sur place, sa rotation autour de l'essieu va entraîner la propulsion en avant du véhicule qu'elle supporte. Enfin, étant ronde, elle va rouler et accompagner le véhicule dans son déplacement, et pouvoir reproduire la même mécanique de propulsion à chaque nouvelle position du véhicule.

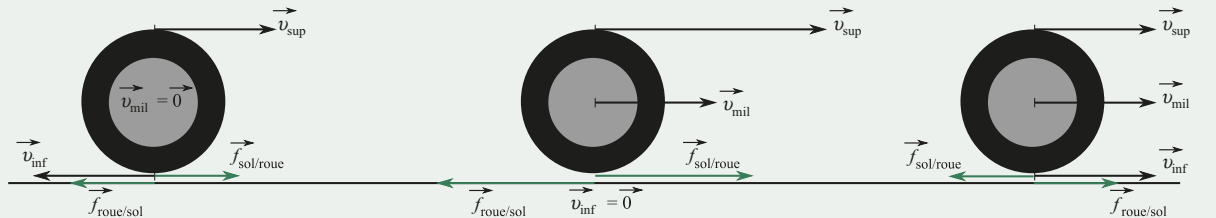
Il nous reste à expliquer pourquoi cette force de frottement ne travaille pas. Commençons par aborder une situation dans laquelle elle travaille : si les frottements ne sont pas suffisants pour assurer la propulsion et que la roue patine, ou que suite à un freinage brutal, le véhicule dérape, la roue glisse sur la route. Son point de contact avec le support se déplace, la force de frottement s'exerçant au niveau de ce point de contact va donc travailler. Étant dans ce cas de figure opposée au déplacement du point de contact, ce travail sera donc strictement négatif, ainsi que l'affirme l'énoncé.

Qu'en est-il à présent du cas où le point de contact ne glisse pas sur la route ? Ne glissant pas sur la route, on est tenté de dire que sa vitesse par rapport à la route est nulle, mais la chose semble paradoxale si l'on imagine une voiture lancée ne serait-ce qu'à quelques dizaines de kilomètres par heure. Et pourtant, si l'on considère les vitesses respectives du point le plus bas, et du point le plus haut de la roue, chacun est animé d'une vitesse totale par rapport au référentiel terrestre, égale à la somme vectorielle de la vitesse du centre de la roue (donc celle du véhicule) par rapport au référentiel terrestre, et de celle dont ce point est animé par rapport à l'essieu. Or le point le plus bas se déplace vers l'arrière de la voiture, tandis que le point le plus haut se déplace vers l'avant. Si la voiture ne glisse pas (on parle alors de *roulement sans glissement*), la longueur de l'arc de cercle dont tourne la roue en une durée donnée est égale à la longueur dont se déplace le véhicule sur la même durée. En d'autres termes, la valeur de la vitesse par rapport à la voiture d'un point périphérique de la roue est égale dans ce cas à la valeur de la vitesse de la voiture par rapport au référentiel terrestre.

Il s'ensuit que, si la voiture se déplace à une vitesse de valeur  $v$  par rapport au référentiel terrestre, alors, par rapport à ce même référentiel :

- Son point le plus haut est animé d'une vitesse de  $2v$ .
- Son point le plus bas est animé d'une vitesse nulle.

Sens de la marche



- |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <ul style="list-style-type: none"> <li>- Les frottements de la roue sur le sol projettent les éléments du sol (graviers, neige, boue, ...) vers l'arrière</li> <li>- Les frottements du sol sur la roue sont insuffisants pour permettre à la roue de se propulser</li> <li>- Le véhicule ne se déplace pas</li> <li>- Le point de contact est animé d'une vitesse vers l'arrière</li> <li>- Il est constamment renouvelé, mais les frottements travaillent : la gomme et le sol s'usent</li> </ul> | <ul style="list-style-type: none"> <li>- Les frottements de la roue sur le sol le poussent vers l'arrière, il ne bouge pas</li> <li>- Les frottements du sol sur la roue la propulsent vers l'avant ; aux vitesses précédentes s'ajoute la vitesse de la voiture</li> <li>- Le véhicule avance, suivant l'impulsion communiquée par le sol, via les roues</li> <li>- Le point de contact, animé de la vitesse du véhicule et de sa vitesse par rapport à celui-ci, a une vitesse globale nulle</li> <li>- Il est constamment renouvelé, et les frottements ne travaillent pas</li> </ul> | <ul style="list-style-type: none"> <li>- Les frottements de la roue sur le sol projettent les éléments du sol (graviers, neige, boue, ...) vers l'avant</li> <li>- Les frottements du sol sur la roue freinent sa progression vers l'avant</li> <li>- Le véhicule se déplace du fait de son inertie de départ</li> <li>- Le point de contact est animé d'une vitesse vers l'avant</li> <li>- Il n'est pas renouvelé, et les frottements travaillent : la gomme et le sol s'usent</li> </ul> |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|

En l'absence de déplacement du point de contact, la force de frottement, dont le point d'application est justement ledit point de contact, ne travaille donc pas. Tout le génie de la roue est de remplacer un point de contact animé de la vitesse du véhicule en train de déraiper, par un point de contact continuellement renouvelé, qui est alors toujours animé d'une vitesse nulle et de ce fait ne travaille pas. Ce système est tellement commun aujourd'hui que son ingéniosité nous échappe parfois, mais l'invention de la roue demeure l'un des prodiges de l'esprit humain.

4. Indépendamment de la question de la conservation de l'énergie, se pose la question de la possibilité de la convertir, c'est-à-dire celle du dispositif alimenté en énergie sous une certaine forme, et qui la restitue sous une ou plusieurs autre(s) forme(s). Un barrage hydroélectrique utilise l'énergie mécanique de l'eau (énergie véhiculée par le mouvement spontané de l'eau des hautes vers les basses altitudes) pour faire tourner une turbine. Cette turbine est conçue pour tourner sous l'effet de cette arrivée d'énergie mécanique et, solidarisée avec un alternateur, délivrer une énergie électrique (énergie véhiculée par le mouvement de porteurs de charges libres de se déplacer).

Les exemples sont nombreux et peuvent être couplés les uns avec les autres. Ainsi par exemple un moteur à explosion peut-

il être couplé avec un alternateur (petite turbine) pour fournir de l'électricité : c'est ainsi que la batterie d'un véhicule se recharge. De même, une brosse à dents électrique utilisée dans une maison dotée de panneaux solaires convertit indirectement l'énergie solaire en énergie mécanique de rotation de la tête brossante.

La conversion de l'énergie responsable de l'instabilité des noyaux radioactifs (énergie nucléaire, donc) peut être utilisée pour fournir de l'électricité. Basiquement, le principe est le même que pour une centrale thermique : une source de chaleur (carburants fossiles dans un cas, noyaux radioactifs dans l'autre) est utilisée pour faire évaporer de l'eau (conversion d'énergie chimique ou nucléaire, en énergie de changement d'état physique). On fait ensuite circuler cette vapeur dans une turbine, où elle va se refroidir partiellement, et également fournir de l'énergie mécanique (conversion d'énergie chimique en énergies mécanique et thermique). Enfin, un alternateur couplé à cette turbine permet de délivrer du courant électrique (conversion d'énergie mécanique en énergie électrique).

Ceci confirme la première partie de l'affirmation fournie par l'énoncé : il est effectivement possible de convertir de l'énergie nucléaire en énergie électrique. L'erreur doit donc se situer au niveau de la deuxième partie, qui affirme que la conversion

| Énergie en entrée | Convertisseur                    | Énergie en sortie      |
|-------------------|----------------------------------|------------------------|
| Liaison chimique  | Batterie                         | Électrique             |
| Électrique        | Électrolyseur                    | Chimique               |
| Liaison chimique  | Moteur à explosion, muscle       | Thermique et mécanique |
| Thermique         | Transformation chimique, cuisson | Liaison chimique       |
| Mécanique         | Éolienne                         | Électrique             |
| Électrique        | Moteur électrique                | Mécanique              |
| Thermique         | Eau + turbine à vapeur           | Électrique             |
| Électrique        | Radiateur électrique             | Thermique              |
| Lumineuse         | Panneau solaire                  | Électrique             |
| Électrique        | Ampoule électrique               | Lumineuse              |
| Acoustique        | Microphone                       | Électrique             |
| Électrique        | Haut-parleur                     | Acoustique             |
| Nucléaire         | Désintégration                   | Rayonnement $\gamma$   |
| Électrique        | ?                                | Nucléaire              |

inverse est impossible. Nous devons donc trouver un dispositif alimenté en énergie par du courant électrique, et qui permette de transformer des noyaux atomiques stables, en noyaux atomiques instables. À cette fin, on peut simplement penser à réaliser l'opération inverse d'une désintégration nucléaire : au lieu de voir un noyau radioactif se désintégrer en éjectant une particule de haute énergie, ne peut-on pas retourner cette particule de haute énergie vers un noyau stable, en espérant qu'il s'y insère, augmentant au passage le nombre de nucléons du noyau cible et rendant donc possiblement celui-ci instable ? C'est précisément ce que réalisent les accélérateurs de particules du monde entier, qui accélèrent des particules diverses (protons, électrons...) au moyen de champs électromagnétiques, pour les précipiter sur des noyaux cibles. Or la création de champs électromagnétiques (qui plus est passablement puissants) se fait grâce à des électroaimants eux-mêmes alimentés par des circuits électriques. Le but n'est plus en général de déterminer de nouveaux éléments, mais d'analyser les diverses particules issues des événements ainsi provoqués (les nucléons n'étaient que le commencement : la relation  $E = mc^2$  se lit dans les deux sens, et il n'est pas rare que l'énergie engagée dans un tel processus soit en partie utilisée à la création de nouvelles particules, pour peu que cette création respecte certaines lois de conservation).

5. Cette affirmation, qui au premier abord peut sembler évidente, est pourtant fausse en toute rigueur. On sait en effet que pour augmenter la température d'un corps, il est nécessaire d'augmenter son énergie interne. Or qui dit masse dit énergie, d'après la relation d'Einstein, mais la réciproque est également vraie : qui dit augmentation de l'énergie dit augmentation de masse. Les valeurs engagées sont certes très faibles : le facteur  $c^2$  dans  $E = mc^2$  vaut  $9,00 \cdot 10^{16} \text{ m}^2 \cdot \text{s}^{-2}$ . Prenons par exemple le cas d'une masse  $m = 1,0 \text{ kg}$  d'eau à

$20^\circ \text{C}$ , dont la capacité thermique massique a pour valeur  $c_{\text{eau}} = 4,2 \cdot 10^3 \text{ J} \cdot \text{kg}^{-1} \cdot \text{K}^{-1}$ . Si l'on amène cette eau liquide à  $70^\circ \text{C}$ , sa température a augmenté de  $\Delta T = 50 \text{ K}$ , et ce faisant son énergie a augmenté de :

$$\Delta E = \Delta U = m_{\text{eau}} \Delta T = 2,1 \cdot 10^5 \text{ J}$$

À cet accroissement d'énergie correspond un accroissement équivalent de masse  $\Delta m = \frac{\Delta E}{c^2} = 2,3 \cdot 10^{-12} \text{ kg}$ , soit donc  $2,3 \text{ ng}$ .

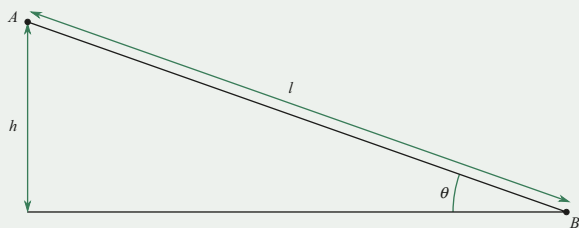
## Du Tac au Tac

1. Commençons par définir le système : la skieuse avec son équipement, que nous étudierons par rapport au référentiel terrestre, supposé galiléen. La skieuse est soumise à :

- Son poids  $\vec{P}$ , force conservative à laquelle on peut associer l'énergie potentielle de pesanteur  $E_{p,P} = mgz$  pour un axe vertical orienté vers le haut.
- La réaction normale  $\vec{R}_n$  de la piste, qui ne travaille jamais et ne peut donc pas plus générer de gain que de perte d'énergie mécanique pour le système.
- La résultante  $\vec{f}$  des forces de frottement qu'elle subit.

L'énergie mécanique initiale du système se résume, en l'absence de vitesse initiale, à son énergie potentielle de pesanteur à la position de départ  $A$  :  $E_{p,A} = mgh$ ,  $h$  désignant l'altitude de cette position par rapport au bas de la piste. On peut exprimer celle-ci à partir de  $l$  et de  $\theta$ , selon :

$$\sin \theta = \frac{h}{l} \Leftrightarrow h = l \sin \theta = l \theta_{\text{rad}}$$



puisque  $\theta_{\text{rad}} = 0,1 \text{ rad} \ll 1$  et que l'on peut donc assimiler la valeur de  $\theta$  en radians à celle de son sinus.

Nous en déduisons une énergie mécanique initiale  $E_{m,A} = mgl\theta_{\text{rad}} = 40 \text{ kJ}$ .

Nous savons par ailleurs, d'après le théorème de l'énergie mécanique, que la variation d'énergie mécanique d'un système par rapport à un référentiel galiléen, est égale au travail des forces non conservatives s'exerçant sur ce système, soit ici :

$$W_A^B(\vec{f}) = \Delta E_m = \Delta E_c + \Delta E_p$$

Explicitons alors les variations d'énergie potentielle et d'énergie cinétique :

$$\Delta E_c = \frac{1}{2}mv_B^2 - 0 = \frac{1}{2}mv_B^2$$

$$\text{et } \Delta E_p = -mgh = -mgl\theta_{\text{rad}}$$

où  $h$  est le dénivelé parcouru par la skieuse entre A et B. En injectant ces deux expressions dans l'expression du travail de  $\vec{f}$ , nous obtenons ainsi un travail des forces de frottement :

$$W_A^B(\vec{f}) = \frac{1}{2}mv_B^2 - mgl\theta_{\text{rad}} = -24 \text{ kJ}$$

En ramenant la valeur absolue de cette perte énergétique, à la valeur initiale de l'énergie mécanique de la skieuse, nous constatons qu'elle a perdu, au cours de la descente :

$$\frac{|\Delta E_m|}{E_{m,A}} \times 100 = \left(1 - \frac{v_B^2}{2gl\theta_{\text{rad}}}\right) \times 100 = 60 \%$$

de son énergie mécanique de départ, c'est-à-dire que toute son énergie potentielle de départ ne se retrouve pas dans l'énergie cinétique acquise à la fin : une (grande) partie a servi à réchauffer insensiblement l'air et la skieuse, mais aussi à réchauffer et faire fondre la neige, formant ainsi la mince couche d'eau permettant au ski de glisser.

**2.** Le véhicule freinant sans déraper, on peut supposer que la perte d'énergie mécanique du véhicule n'est pas due aux frottements de la route sur les pneumatiques (cf. à ce sujet l'étude sur le roulement de la roue, dans la section « Halte aux idées reçues »). Le principe des freins est de provoquer, au sein même du véhicule, la mise au contact de deux pièces (disque et plaquette, par exemple), qui vont frotter l'une sur l'autre. Le travail résultant va permettre une importante conversion d'énergie mécanique (d'où le freinage du véhicule), en éner-

gie interne des pièces de freinage (d'où leur élévation de température).

Si nous supposons que la voiture n'échange aucune énergie avec le milieu extérieur (en particulier si nous négligeons les frottements de l'air et supposons que les pneumatiques ne dérapent pas), nous pouvons écrire :

$$\Delta E = \Delta E_m + \Delta U = 0 \quad \Leftrightarrow \quad \Delta U = -\Delta E_m$$

La voiture évoluant sur une route horizontale, son énergie potentielle ne varie pas, et sa variation d'énergie mécanique se résume donc à sa variation d'énergie cinétique :

$$\Delta E_m = \Delta E_c + \underbrace{\Delta E_p}_0 = \Delta E_c = -\frac{1}{2}Mv^2$$

En exprimant par ailleurs la variation d'énergie interne du dispositif de freinage, nous obtenons :

$$\Delta U = mc\Delta T$$

où  $\Delta T$  désigne la variation de température de ce dispositif, consécutive au freinage.

En injectant ces deux égalités dans l'expression du bilan d'énergie nous obtenons :

$$mc\Delta T = \frac{1}{2}Mv^2 \quad \Leftrightarrow \quad \Delta T = \frac{Mv^2}{2mc} = 25 \text{ K}$$

Les disques et plaquettes s'échauffent donc dans ce cas de 25 °C (rappelons que les échelles Celsius et Kelvin ne diffèrent que par une constante additive de 273,15, le retranchement de cette constante lors d'une différence fait qu'un **écart** de température a même valeur en °C et en K). Ceci explique le fait que les dispositifs de freinage sont particulièrement chauds après utilisation.

On peut à cet égard noter que cette variation de température varie comme le carré de la vitesse initiale du véhicule. Ainsi, pour un véhicule qui sur circuit freine avec une vitesse initiale de 144 km.h<sup>-1</sup> (soit le double de la valeur proposée dans cet exercice), la variation de température aurait été quatre fois supérieure, soit donc de 100 °C. Ce point prend des proportions qui réclament toute l'attention des ingénieurs lorsqu'ils conçoivent des véhicules dédiés à la course : ceux-ci montant parfois au voisinage de 300 km.h<sup>-1</sup> (prenons 288 km.h<sup>-1</sup>), nous constatons cette fois une montée de 400 °C, qui peut altérer certains matériaux. Les dispositifs de freinage sont alors conçus dans des matériaux particuliers, dotés d'une capacité thermique massique plus importante, et qui présentent une résistance à la chaleur particulièrement éprouvée.

**3.** Exprimons pour commencer la variation d'énergie potentielle de la voiture, consécutive à son élévation d'une hauteur  $h$  :

$$\Delta E_p = Mgh$$

Exprimons ensuite la variation d'énergie interne nécessaire à la fusion du glaçon :

$$\Delta U = mL$$

Si nous égalisons ces deux termes entre eux (hypothèse de base de l'exercice), nous obtenons :

$$Mgh = mL \Leftrightarrow h = \frac{mL}{Mg} = 0,50 \text{ m}$$

On constate donc que l'énergie de fusion d'un glaçon permet de faire monter un véhicule de type petite voiture citadine, d'un demi-mètre. Pour rendre la chose plus parlante, on peut dire qu'avec l'énergie de fusion de 10 glaçons, on ferait monter cette voiture à 5,0 m, et qu'avec 1 kg de glace, on pourrait catapulte verticalement ce véhicule à une hauteur de 50 m. Au final, il suffirait donc de l'énergie de fusion d'environ 65 kg de glace pour la lancer au sommet de la tour Eiffel.

Ceci peut paraître surprenant au premier abord, mais il ne s'agit que de la conséquence logique d'un autre fait surprenant, à savoir que les capacités thermiques massiques et les chaleurs latentes de changement d'état, exprimées pour 1 kg de matière (et 1 °C de variation de température pour les capacités calorifiques), se comptent déjà en kJ, voire en dizaines ou centaines de kJ, quand élever 1 kg de matière de 1 m à la surface de la Terre coûte en tout et pour tout 10 J seulement.

On peut ainsi retenir que les énergies engagées dans les processus mécaniques sont en général très faibles devant celles engagées dans les processus thermodynamiques.

**4.** La légère difficulté de cet exercice provient du fait que la présence du paramètre temps réclame un bilan de puissance, plutôt qu'un bilan énergétique. Un bilan de puissance revient en substance au même qu'un bilan énergétique, mais chaque échange d'énergie doit être ramené à une unité de temps.

Dans le cas présent, par exemple, l'énoncé fournit un débit volumique d'eau au lieu d'un volume d'eau : il ne nous parle pas de 8,0 m<sup>3</sup> d'eau qui dévaleraient 50 m d'altitude pour transférer leur énergie mécanique à une turbine qui la convertirait, avec 75 % d'efficacité, en énergie électrique. Il nous parle de 8,0 m<sup>3</sup> d'eau **par seconde**, qui vont transférer tant d'énergie mécanique **par seconde** à une turbine, qui va délivrer sous cet effet tant d'énergie électrique **par seconde**. Toutes les grandeurs doivent être considérées par unité de temps, ce qui pour les énergies revient à basculer des joules aux joules par seconde, c'est-à-dire aux watts.

Ainsi, la puissance mécanique apportée par l'eau à la turbine (énergie mécanique apportée par l'eau, par unité de temps) peut s'exprimer comme :

$$\mathcal{P}_{\text{méca}} = \frac{m_{\text{eau}} gh}{\Delta t} = D_m gh = \rho_{\text{eau}} D_v gh$$

L'énoncé précisant que la puissance électrique délivrée par cette turbine est égale à  $\eta = 75 \%$  de la puissance mécanique reçue, nous pouvons ainsi écrire :

$$\mathcal{P}_{\text{élec}} = \eta \mathcal{P}_{\text{méca}} = \eta \rho_{\text{eau}} D_v gh = 3,0 \text{ MW}$$

**5.** L'énoncé nous fournit, entre autres choses, l'activité molaire du plutonium 238, c'est-à-dire le nombre de désintégrations par seconde et par mole de noyaux de cet iso-

tope (rappelons que 1 Bq = 1 désintégration par seconde). Connaissant l'énergie individuelle dégagée par une seule de ces désintégrations, nous pouvons déterminer l'énergie délivrée par seconde (autrement dit la puissance) et par mole d'atomes de cet isotope.

Exprimons donc la puissance molaire de cet isotope, sous la forme :

$$\frac{\mathcal{P}}{n} = \frac{A \times E_\gamma}{n} = 10^2 \text{ W.mol}^{-1}$$

c'est-à-dire que la désintégration d'une mole de noyaux de plutonium 238 délivre une puissance de 10<sup>2</sup> W (10<sup>2</sup> J par seconde).

On nous précise que la pile du stimulateur doit fournir une puissance électrique  $\mathcal{P}_{\text{stim}} = 10^{-1} \text{ W}$ .

Nous en déduisons, compte tenu de ce qui précède, la quantité de matière de plutonium 238 nécessaire à la libération de cette puissance :

$$n = \frac{\mathcal{P}_{\text{stim}}}{\frac{\mathcal{P}}{n}} = \frac{\mathcal{P}_{\text{stim}}}{\frac{A \times E_\gamma}{n}} = 10^{-3} \text{ mol}$$

En effet, à 10<sup>2</sup> W par mole, un besoin 10<sup>3</sup> fois plus faible (10<sup>-1</sup> W) réclame une quantité de matière 10<sup>3</sup> fois plus faible également, soit donc 10<sup>-3</sup> mol.

Il nous reste à déterminer quelle masse de plutonium 238 cette quantité de matière représente :

$$m = nM_{238\text{Pu}} = 0,2 \text{ g}$$

si l'on prend  $M_{238\text{Pu}} = 2.10^2 \text{ g.mol}^{-1}$ , ce qui pour un isotope à 238 nucléons semble cohérent (1 g.mol<sup>-1</sup>, et par nucléon).

Par ailleurs, l'énoncé nous informe que la demi-vie du plutonium est de près d'un siècle, c'est-à-dire qu'au bout de cette durée la population de noyaux radioactifs au sein d'un échantillon donné aura diminué de moitié. Ce faisant, son activité aura également diminué de moitié, et avec elle la puissance délivrée au stimulateur. Cependant, même si cette baisse provoquait un mauvais fonctionnement du stimulateur, ceci supposerait que son porteur l'ait conservé un siècle durant, longévité plutôt rare, d'autant que les personnes s'en trouvant équipées sont souvent déjà relativement âgées au moment de la pose.

Ainsi, ce faible rythme de décroissance garantit une relative stabilité de la puissance délivrée, à l'échelle d'une vie humaine.

Il est d'ailleurs intéressant de noter, sans chercher à minimiser les risques liés à l'énergie nucléaire, le fait que les isotopes possédant une activité très élevée sont certes potentiellement très agressifs pour l'être humain, mais qu'un échantillon les portant verra son activité diminuer rapidement. À l'inverse, un isotope de demi-vie élevée aura un rythme de désintégration moins important et, à nombre égal de noyaux au départ, sera moins agressif mais durera plus longtemps. Ces considérations sur le rythme de désin-

tégration ne doivent cependant pas faire oublier que même si un isotope peut sembler lent à la désintégration, le nombre de désintégrations peut malgré tout être important à l'échelle humaine, pour peu que l'échantillon en comporte une quantité suffisante.

## Vers la prépa

**1.** L'énergie mécanique d'un système est la somme de son énergie cinétique et des énergies potentielles dont dérivent les forces auxquelles il est soumis. En l'occurrence, il n'est soumis qu'à la force gravitationnelle :

$$\vec{F}_{g,T/M} = -\frac{GM_T m}{(TM)^2} \vec{e}_{TM}$$

où  $M_T$  est la masse de la Terre.

En notant  $r = TM$  la distance séparant le centre de la Terre et le satellite, et  $\vec{e}_r = \vec{e}_{TM}$  le vecteur radial des coordonnées polaires, nous constatons que cette force se résume à une coordonnée radiale en coordonnées polaires, et que cette coordonnée dépend en outre de la seule coordonnée radiale  $r$ .

Nous pouvons ainsi écrire que :

$$\begin{aligned} \frac{dE_{p,g}}{dr} &= -\vec{F}_{g,T/M} \cdot \vec{e}_r = -\frac{GM_T m}{r^2} \\ \Leftrightarrow E_{p,g}(r) &= -\frac{GM_T m}{r} (+cte) \end{aligned}$$

En prenant pour convention que l'énergie potentielle d'interaction gravitationnelle soit nulle lorsque le satellite et la Terre se trouvent à l'infini l'un de l'autre, nous trouvons que la constante d'intégration est nulle et pouvons alors associer à la force gravitationnelle exercée par la Terre sur le satellite

$$\text{l'énergie potentielle } E_p(r) = -\frac{GM_T m}{r}.$$

$$\text{Nous en déduisons que } E_m = -\frac{GM_T m}{r} + \frac{1}{2}mv^2.$$

**2.** La seule force s'exerçant sur le satellite étant conservative, nous pouvons affirmer que son énergie mécanique est conservée au cours du mouvement :  $E_m = cte$ . Donc

$$-\frac{GM_T m}{r} + \frac{1}{2}mv^2 = cte.$$

Si nous ajoutons l'hypothèse d'une trajectoire circulaire, c'est-à-dire que  $r = cte$ , il vient alors immédiatement  $v = cte$  : l'orbite est décrite avec une vitesse de valeur constante, donc le mouvement est uniforme.

**3.** Au départ de la surface de la Terre, le satellite se trouve à une distance  $r_1 = R_T$  (rayon de la Terre) du centre de la Terre. Si en outre on lui communique une vitesse  $v_1$ , alors son énergie mécanique a pour expression  $-\frac{GM_T m}{R_T} + \frac{1}{2}mv_1^2$ .

Considérons alors le satellite en une autre position  $r$ , où il est animé d'une vitesse de valeur  $v$ . L'énergie mécanique étant conservée, nous pouvons écrire :

$$\begin{aligned} -\frac{GM_T m}{R_T} + \frac{1}{2}mv_1^2 &= -\frac{GM_T m}{r} + \frac{1}{2}mv^2 \\ \Leftrightarrow v^2 &= v_1^2 - 2GM_T \left( \frac{1}{R_T} - \frac{1}{r} \right) \end{aligned}$$

À mesure que  $r$  va augmenter,  $\frac{1}{r}$  va diminuer, donc  $\frac{1}{R_T} - \frac{1}{r}$  va augmenter, donc  $v^2$  va diminuer. Or  $v^2$  ne peut être négative. On doit donc nécessairement avoir :

$$v_1^2 \geq 2GM_T \left( \frac{1}{R_T} - \frac{1}{r} \right) \Leftrightarrow \frac{1}{r} \geq \frac{1}{R_T} - \frac{v_1^2}{2GM_T}$$

À partir de là, deux possibilités :

- Soit  $\frac{1}{R_T} - \frac{v_1^2}{2GM_T} \geq 0$ , et l'inégalité se résume à :

$$r \leq \left( \frac{1}{R_T} - \frac{v_1^2}{2GM_T} \right)^{-1} \quad \text{d'où } r_{\max} = \left( \frac{1}{R_T} - \frac{v_1^2}{2GM_T} \right)^{-1}$$

- Soit  $\frac{1}{R_T} - \frac{v_1^2}{2GM_T} \leq 0$  et dans ce cas l'inégalité est vérifiée quel que soit  $r$ . Cette seconde option peut également s'écrire :

$$v_1 \geq \sqrt{\frac{2GM_T}{R_T}}$$

Cette vitesse, appelée vitesse de libération de la Terre, correspond à une énergie cinétique suffisante pour que la vitesse du satellite tombe à 0, précisément lorsqu'il atteint une distance infinie de la Terre.

**4.** Nous savons que dans le cas où travaille une force non conservative, telle qu'une force de frottement, l'énergie mécanique du système varie. Nous savons encore que cette variation s'identifie alors au travail de cette force :

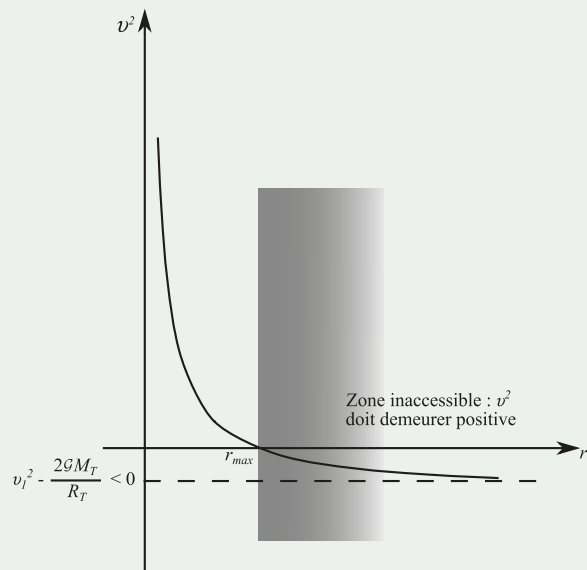
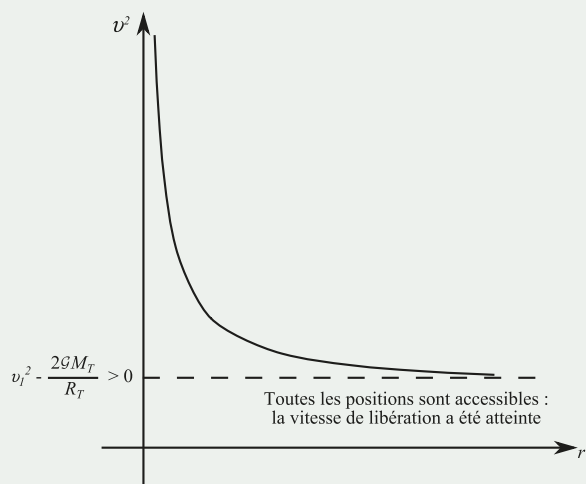
$$\Delta E_m = W(\vec{F}_{nc})$$

Dans le cas d'une force de frottement, résistante par nature, cette variation est donc négative.

Nous en concluons que l'énergie mécanique du satellite va décroître. Or, en mouvement circulaire,

$$E_m(r) = -\frac{GM_T m}{r} + \frac{mr^2\omega^2}{2} \text{ est une fonction croissante de la variable } r.$$

Donc, si  $E_m(r)$  diminue, alors  $r$  doit également diminuer : le satellite se rapproche de plus en plus de la Terre. Ce faisant, il pénètre de plus en plus dans l'atmosphère subissant du même coup des frottements de plus en plus importants. Ce cercle vicieux va précipiter sa chute, si aucune correction de trajectoire n'est apportée, par exemple au moyen d'un moteur d'appoint.



# Index

## A

absorbance 50, 51  
acide 75  
acide fort 75  
alcalin 6  
alcalino-terreux 7  
alcane 19  
alcène 19, 20  
alcool 21  
alcool primaire 22  
alcool secondaire 22  
alcool tertiaire 22  
aldéhyde 22  
ampholyte 76  
amphotère 76  
analyse IR 30  
analyse spectrale 29  
analyse UV-visible 30  
anode 90, 91  
autoprotolyse de l'eau 77

## B

base 75  
base de Frénet 153  
base forte 75

## C

catalyseur 70, 128  
cathode 90, 91  
célérité 162  
cétone 22  
charge partielle 10  
chiralité 24  
cinétique chimique 127  
classe d'un alcool 21  
conductance 48  
conductimètre 47  
conductivité 48  
conductivité ionique molaire 49  
conformation 18  
conformation décalée 18  
conformation éclipsée 18  
constante de Planck 191  
constante d'équilibre 72, 73  
couple acide/base 75  
couple oxydant/réducteur 85

## D

demi-équation électronique 85, 86  
demi-équivalence 105  
diastéréoisomère 25  
diastéréoisométrie 25  
dispersion 170  
dosage 103  
dosage chimique 42  
dosage conductimétrique 113  
dosage direct 106  
dosage d'oxydoréduction 114  
dosage en retour 108  
dosage indirect 107  
dosage par différence 108  
dosage pH-métrie 110

## E

électrolyse 94  
électronégativité 7  
électron-volt 192  
élément 1  
énantiométrie 25  
énergie potentielle de pesanteur 262  
équivalence d'un dosage 110  
équivalence d'un titrage 104  
étalonnage d'un pH-mètre 46

## F

force d'un acide 79  
formule brute 17  
formule développée 17  
formule semi-développée 17  
formule topologique 18

## G

gaz noble 7  
géométrie d'un édifice 12  
groupe fonctionnel 17

## H

halogène 7

## I

indicateur coloré acido-basique 44, 112  
interaction nucléaire forte 193  
isométrie de constitution 23

isométrie Z-E 24  
isotope 1

## K

Kepler 151, 152

## L

lentille mince 208  
liaison de covalence 9  
liaison polarisée 9  
loi de Beer-Lambert 51

## M

mécanique quantique 191  
mélange racémique 25  
méthode d'analyse physique 42  
méthode de Gillespie 12  
méthode des tangentes 112  
méthode VSEPR 12

## N

Newton  
deuxième loi 147  
première loi 147  
troisième loi 148  
nombre de masse 1  
numéro atomique 1

## O

onde 196  
onde longitudinale 163  
onde transversale 163  
oxydant 84, 86

## P

pH 43  
pH-mètre 45  
pile 89  
pont salin 90  
pression 277  
principe des actions réciproques 148  
principe d'inertie 147  
principe fondamental de la dynamique 147  
prisme 207  
produit ionique de l'eau 77

## Q

quantité de mouvement 147  
quotient de réaction 71

## R

réaction acido-basique 74  
réaction athermique 128  
réaction d'oxydoréduction 84  
réaction endothermique 128  
réaction exothermique 128  
réaction lente 127  
réaction quantitative 73  
réducteur 84, 86  
référentiel galiléen 146

réflexion totale 207  
règles du duet et de l'octet 7  
représentation de Cram 18

## S

spectre d'émission d'un corps dense 28  
spectre de raies d'absorption 28  
spectrophotomètre 52  
squelette carboné 16  
stéréoisomérisation de configuration 23  
système de coordonnées cartésiennes 142

## T

tableau d'avancement 66

température 277  
titrage 42, 103  
transmittance 30  
trempe 127, 128

## V

vecteur position 144  
vecteur vitesse 144

## Z

zone de virage 44, 112