

Bertrand HAUCHECORNE

Les
CONTRE-EXEMPLES
en **MATHÉMATIQUES**

2^e édition

revue et augmentée



ellipses

LES CONTRE-EXEMPLES EN MATHÉMATIQUES

Deuxième édition, revue et augmentée

522 contre-exemples

Bertrand HAUCHECORNE

Agrégé de l'Université
Professeur de mathématiques
en classes préparatoires aux grandes écoles scientifiques
au lycée Pothier (Orléans)



Du même auteur

Des mathématiciens de A à Z

Avec Daniel Suratteau, Ellipses 1996, 1999.

Lexique bilingue du vocabulaire mathématique

Avec Adrian Shaw, Ellipses 2000.

Les mots et les maths

Dictionnaire historique et étymologique du vocabulaire mathématique

Ellipses 2003.

ISBN 978-2-7298-3418-0

© Ellipses Edition Marketing S.A., 2007

32, rue Bargue, 75740 Paris cedex 15.



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 ». (Article L.122-4).

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

www.editions-ellipses.fr

A Sylvie

Préface de la deuxième édition

Contre-exemple, ce terme fait parfois sourire car il rappelle quelque fonction farfelue proposée aux étudiants pour les mettre en garde contre l'emploi abusif d'un théorème, mais il ne semble pas digne d'intérêt car il ne traite pas le cas général, et il est pour cela relégué au registre des divertissements. Sans enlever au contre-exemple cette dernière qualité, nous désirons par cet ouvrage montrer sa valeur mathématique et sa vertu pédagogique.

Qu'est-ce qu'un contre-exemple ?

Le contre-exemple n'est pas l'exception qui confirme la règle du bon sens commun. Un théorème est valable dans tous les cas où les hypothèses qu'il impose sont vérifiées. La négation d'un énoncé, c'est-à-dire l'affirmation qu'il est faux, est démontrée par l'existence d'un cas où ses hypothèses sont vérifiées sans que la conclusion le soit, ce qui s'exprime dans le langage de la logique par le fait que nier la proposition "Pour tout x , $P(x)$ " c'est affirmer l'assertion "Il existe au moins un x tel que $\text{NON}(P(x))$ ". La justification mathématique de la fausseté d'un énoncé est la donnée d'un contre-exemple. La recherche d'une quelconque généralité est inutile et même erronée au niveau de la démarche.

Des motivations pédagogiques amènent aussi à l'étude de contre-exemples. Un théorème nécessite la plupart du temps plusieurs hypothèses. Pour bien comprendre son mécanisme et mieux cerner son champ d'application, il est important d'être convaincu de la nécessité de chacune d'entre elles. On est donc amené à montrer que l'énoncé obtenu en supprimant l'une

des hypothèses est faux. L'étude d'un résultat avec des hypothèses plus générales que celles habituellement utilisées entre dans le même type de préoccupations.

Certaines notions, lors de leur première étude, donnent l'impression qu'un résultat doit être vrai, par exemple que toute fonction continue en zéro est continue sur un voisinage de zéro ; si cette intime conviction surgit, c'est que la notion intuitive que l'on se fait de la continuité est erronée : on envisage en fait des fonctions beaucoup plus régulières que ce qu'énonce la définition. La donnée de quelques contre-exemples à ces idées préconçues permet de rectifier cette mauvaise intuition.

On est souvent amené en mathématiques à considérer plusieurs définitions de plus en plus restrictives pour pouvoir cerner des résultats plus précis ; citons par exemple la continuité et la continuité uniforme. Pour permettre une réflexion plus poussée sur ces notions, il nous a paru intéressant de montrer que certaines définitions sont effectivement plus restrictives que d'autres, en explicitant des exemples vérifiant seulement l'une d'entre elles.

Conception de l'ouvrage

Nous avons refusé la notion de catalogue ; les contre-exemples sont placés dans une problématique pour justifier leur intérêt. Nous rappelons au passage les définitions et les théorèmes utiles à leur compréhension sans nous attarder à leur démonstration, que l'on peut trouver dans les livres courants. Pour certaines notions, nous précisons quelques références d'ouvrages dans lesquels le lecteur pourra compléter ses connaissances. Certains contre-exemples ont une valeur historique, dans la mesure où un mathématicien de renom les a introduits pour contredire des conjectures en vogue à son époque ; si par la suite des contre-exemples plus simples ont été trouvés, nous présentons les deux exemples, agrémentés d'une note historique.

Nous avons placé notre étude, chaque fois que c'est possible, dans le cadre le plus simple, préférant par exemple le corps des nombres réels à un espace topologique très général. Cependant certains contre-exemples nécessitent plus de connaissances et de technicité. Les démonstrations sont, à quelques exceptions près, traitées en détail. Pour les exemples les plus simples, elles sont plus détaillées, parfois un peu lourdes pour éviter l'utilisation de résultats ou de méthodes qui pourraient désarçonner le lecteur débutant. Pour les exemples plus délicats, en principe abordées par un lecteur ayant plus d'expérience, elles vont à l'essentiel pour ne pas noyer la démarche dans un amas de précisions trop banales.

Cette nouvelle édition contient plus de 200 exemples supplémentaires. Les énoncés des définitions et des théorèmes ont été introduits de manière systématique. Pour la clarté de la lecture, un "carré vide" $\square$ a été placé à la fin de l'explication de chaque contre-exemple, afin de la séparer de remarques éventuelles et de l'introduction du suivant.

A qui s'adresse cet ouvrage ?

Un grand nombre d'exemples sont accessibles aux étudiants des classes préparatoires aux grandes écoles scientifiques ou des deux premières années de licence de mathématiques ; d'autres manient des notions plus délicates et intéresseront les étudiants de troisième année de licence et de master en mathématique, ou ceux qui préparent les concours du professorat.

Cet ouvrage permettra aux étudiants d'approfondir l'enseignement de mathématiques qu'ils reçoivent, à ceux qui préparent les concours menant à l'enseignement des mathématiques, CAPES et Agrégation, d'enrichir une leçon et aux enseignants de trouver des thèmes d'exercices ou de problèmes. Plus généralement, il intéressera tous ceux qui veulent approfondir leur réflexion sur les notions de définition, d'hypothèse ou de théorème. Il apportera surtout, je l'espère, bien du plaisir à ceux qui partagent avec moi cette curiosité permanente pour les mathématiques.

Remerciements

Je tiens à remercier ceux de mes étudiants qui, par leurs remarques parfois naïves, souvent pertinentes, m'ont permis d'enrichir cet ouvrage.

Mes remerciements vont également aux Editions Ellipses, qui m'ont vivement encouragé à rédiger cette nouvelle édition des Contre-exemples en mathématiques, et à Maryse Facy pour la relecture du dernier chapitre.

Cependant, ma plus profonde gratitude va à mon ami et ancien collègue Daniel Suratteau. Sa relecture toujours avisée a permis de débusquer de nombreuses erreurs dans la première édition et de donner plus de rigueur au texte. Ses suggestions ont enrichi cet ouvrage. Il m'a surtout conseillé d'énoncer plus de définitions et de théorèmes. Par ailleurs sa grande compétence dans le maniement du logiciel $\text{T}_{\text{E}}\text{X}$ de Donald E. Knuth donne à ce livre sa qualité d'impression et a permis d'ajouter de nombreux croquis de sa conception. Sans lui, l'ouvrage aurait gardé le côté austère de la première édition. Qu'il en soit remercié !

Bertrand Hauchecorne

Mareau-aux-Prés, février 1988 et mars 2007

Notations

Le lecteur trouvera ci-dessous une liste des notations « de base » utilisées dans le texte de cet ouvrage, qui lui sont pour la plupart familières. Celles qui n'y figurent pas sont définies dans les chapitres où elles sont utilisées.

► Nombres et ensembles de nombres.

$0, 1, e, \pi, i$	Les nombres les plus célèbres des mathématiques.
$n!, \binom{n}{p}$	Factorielle d'un entier naturel, coefficient binomial.
$\mathbb{N}$	Ensemble des nombres entiers naturels.
$\mathbb{N}^*$	Ensemble des entiers naturels non nuls.
$\llbracket p, q \rrbracket$	Si $p, q \in \mathbb{N}$, ensemble des entiers i tels que $p \leq i \leq q$.
$\mathbb{Z}$	Anneau des nombres entiers relatifs.
$m \mid n, m \nmid n$	m divise n , m ne divise pas n .
$\mathbb{Q}$	Corps des nombres rationnels.
$\mathbb{Q}^*$	Ensemble des nombres rationnels non nuls.
$\mathbb{R}$	Corps des nombres réels.
$\mathbb{R}^*$	Ensemble des nombres réels non nuls.
$\overline{\mathbb{R}}$	Droite numérique achevée $\{-\infty\} \cup \mathbb{R} \cup \{+\infty\}$.
$ x $	Valeur absolue d'un nombre réel x .
$\mathbb{C}$	Corps des nombres complexes.
$\mathbb{C}^*$	Ensemble des nombres complexes non nuls.
$\mathbb{U}$	Groupe multiplicatif des complexes de module 1.
$ z , \bar{z}$	Module, conjugué d'un nombre complexe z .
$\operatorname{Re}(z), \operatorname{Im}(z)$	Partie réelle, partie imaginaire d'un complexe z .

► Ensembles, désignés par les lettres E, F, A, I, A_i pour $i \in I$.

$\emptyset$	L'ensemble vide.
$\{x \in E \mid P(x)\}$	Ensemble des éléments de E tels que $P(x)$ est vrai.
$E \cup F$ et $E \cap F$	Réunion et intersection des ensembles E et F .
$E \setminus F$	Différence $\{x \in E \mid x \notin F\}$ entre E et F .
$\complement_E A$	Si $A \subset E$, complémentaire $E \setminus A$ de A dans E .
$\bigcup_{i \in I} A_i, \bigcap_{i \in I} A_i$	Réunion, intersection d'une famille $(A_i)_{i \in I}$ d'ensembles.
Id_E	Application identique de l'ensemble E .
$\mathcal{P}(E)$	Ensemble des parties de l'ensemble E .
$\mathcal{F}(E, F)$	Ensemble des applications de E dans F .
$f : E \rightarrow F$	f est une application de E dans F .
$\left \begin{array}{l} f : E \longrightarrow F \\ x \longmapsto f(x) = \dots \end{array} \right.$	f est l'application de E dans F telle que, pour tout élément x de E , $f(x) = \dots$.
$f _A$	Si $f : E \rightarrow F$ et si $A \subset E$, restriction de f à A .
$f(A)$	Si $f : E \rightarrow F$ et si $A \subset E$, image directe $\{f(x) \mid x \in A\}$ de A par f .

$\text{Im}(f)$ ou $\text{Im } f$	Si $f : E \rightarrow F$, image $f(E)$ de l'application f .
$f^{-1}(M)$	Si $f : E \rightarrow F$ et si $M \subset F$, image réciproque $\{x \in E \mid f(x) \in M\}$ de M par f .
$E/\mathcal{R}$	Ensemble quotient de E par une relation d'équivalence $\mathcal{R}$ sur E .
$\text{card}(E)$ ou $\text{card } E$	Cardinal d'un ensemble fini E .

► **Algèbre linéaire.**

Soit K un corps (ou un anneau) commutatif et $n, p \in \mathbb{N}^*$.

I_n	Matrice unité d'ordre n .
$\mathcal{M}_{n,p}(K)$	Espace vectoriel des (n, p) -matrices à coefficients dans K .
$\mathcal{M}_n(K)$	Algèbre des matrices carrées d'ordre n à coefficients dans K .

Soit E et F des espaces vectoriels.

$V + W, V \oplus W$	Somme, somme directe des sous-espaces vectoriels V et W de E .
$\mathcal{L}(E, F)$	Espace vectoriel des applications linéaires de E dans F .
$\text{Ker}(f)$ ou $\text{Ker } f$	Noyau de $f \in \mathcal{L}(E, F)$.
$\mathcal{L}(E)$	Algèbre des endomorphismes de E .

► **Polynômes** à coefficients dans un corps (ou un anneau) commutatif K .

$K[X]$	Algèbre des polynômes à coefficients dans K .
$\deg(P)$ ou $\deg P$	Degré du polynôme P .
$\text{val}(P)$ ou $\text{val } P$	Valuation du polynôme P .

► **Relations d'ordre.**

$\text{Min}(A), \text{Max}(A)$	Plus petit, plus grand élément de A .
$\text{Inf}(A), \text{Sup}(A)$	Borne inférieure, borne supérieure de A .

► **Intervalles.**

Soit a et b des réels tels que $a \leq b$.

$[a, b]$	Segment $a b$, c'est-à-dire $\{x \in \mathbb{R} \mid a \leq x \leq b\}$.
$]a, b[$	Intervalle ouvert $\{x \in \mathbb{R} \mid a < x < b\}$.
$[a, b[$	Intervalle $\{x \in \mathbb{R} \mid a \leq x < b\}$.
$]a, b]$	Intervalle $\{x \in \mathbb{R} \mid a < x \leq b\}$.

Soit a un nombre réel.

$[a, +\infty[$	Intervalle fermé $\{x \in \mathbb{R} \mid x \geq a\}$.
$]a, +\infty[$	Intervalle ouvert $\{x \in \mathbb{R} \mid x > a\}$.
$]-\infty, a]$	Intervalle fermé $\{x \in \mathbb{R} \mid x \leq a\}$.
$]-\infty, a[$	Intervalle ouvert $\{x \in \mathbb{R} \mid x < a\}$.
$]-\infty, +\infty[$	L'intervalle $\mathbb{R}$.
$\mathbb{R}_+$	L'intervalle fermé $[0, +\infty[$.
$\mathbb{R}_-$	L'intervalle fermé $]-\infty, 0]$.
$\mathbb{R}_+^*$	L'intervalle ouvert $]0, +\infty[$.
$\mathbb{R}_-^*$	L'intervalle ouvert $]-\infty, 0[$.

Bibliographie

Voici une liste d'ouvrages auxquels le texte fait référence et qui contiennent des précisions sur les notions abordées : définitions, théorèmes et démonstrations détaillées de ces derniers.

- [ARN1] Jean-Marie ARNAUDIÈS & Henri FRAYSSE, *Cours de mathématiques 1, Algèbre*, Dunod 1987.
- [ARN2] Jean-Marie ARNAUDIÈS & Henri FRAYSSE, *Cours de mathématiques 2, Analyse*, Dunod 1988.
- [BARB] BARBE & LEDOUX, *Probabilités*, Belin 1998.
- [BLAN] André BLANCHARD, *Les corps non commutatifs*, PUF 1972.
- [BOUA] Hassan BOUALEM & Robert BROUZET, *La planète $\mathbb{R}$* , Dunod 2002.
- [BRIA] Marc BRIANE & Gilles PAGÈS, *Théorie de l'intégration*, Vuibert 1998.
- [CAU1] Augustin-Louis CAUCHY, *Cours d'Analyse de l'Ecole royale polytechnique, Première partie : Analyse algébrique*, 1821, réédition Jacques Gabay 1989.
- [CAU2] Augustin-Louis CAUCHY, *Résumé des leçons données à l'Ecole royale polytechnique sur le Calcul infinitésimal*, 1823, réédition Ellipses 1994.
- [CHOQ] Gustave CHOQUET, *Cours d'Analyse, tome II : Topologie*, Masson 1964, 1969.
- [DELM] Jean-Pierre DELMAS, *Introduction aux probabilités*, Ellipses 2000.
- [GIAN] GIANELLA, KRUST, TAÏEB & TOSEL, *Problèmes choisis de mathématiques supérieures*, Springer 2001.
- [GUEG] Jean GUÉGAND & Thierry DUGARDIN, *Intégration MP-MP**, Ellipses 1998.
- [GOUR] Xavier GOURDON, *Les maths en tête, Analyse*, Ellipses 1994.
- [HAIR] Ernst HAIRER & Gerhard WANNER, *L'analyse au fil de l'histoire*, Scopos Springer 2000.
- [PERR] Daniel PERRIN, *Cours d'algèbre*, Ellipses 1996.
- [RAM1] RAMIS, DESCHAMPS & ODOUX, *Cours de mathématiques spéciales, tome 1 : Algèbre*, Masson 1974.
- [RAM5] RAMIS, DESCHAMPS & ODOUX, *Cours de mathématiques spéciales, tome 5 : Application de l'analyse à la géométrie*, Masson 1981.
- [SCHW] Lionel SCHWARTZ, *Algèbre 3^e année*, 2^e édition, Dunod 2003.
- [SKAN] Georges SKANDALIS, *Topologie et Analyse, 3^e année*, Dunod 2004.
- [WALT] Wolfgang WALTER, *Analysis 1*, 7^e édition, Springer 2004.

Table des matières

Chapitre 1

Logique, ensembles, arithmétique	1
Logique des prédicats du premier ordre	1
Paradoxes de la théorie « naïve » des ensembles	2
Image directe et image réciproque d'une partie	3
Ensembles équipotents	4
Relations binaires fondamentales	8
Ensembles ordonnés	9
Treillis	11
Bon ordre et axiome du choix	14
Arithmétique	15

Chapitre 2

Groupes	17
Lois de composition interne	17
Axiomes de la structure de groupe	19
Centre d'un groupe	21
Sous-groupes	22
Ordre d'un élément dans un groupe	25
Morphismes et isomorphismes de groupes	28
Groupes simples et groupes résolubles	30

Chapitre 3

Anneaux et corps	33
Propriétés générales	34
Éléments inversibles et diviseurs de zéro	35
Anneau des polynômes à coefficients dans un anneau commutatif	37
Caractéristique d'un anneau non nul	42
Idéaux d'un anneau non nul	43
Divisibilité dans un anneau intègre	47
Autres types d'anneaux	51
Corps	56
Corps ordonnés	57

Chapitre 4

Espaces vectoriels	60
Nécessité des axiomes	61
Sous-espaces vectoriels	62
Endomorphismes d'un espace vectoriel	64
Valeurs propres et vecteurs propres, polynôme caractéristique et polynôme minimal	67

Matrices	73
Modules	76
Formes bilinéaires symétriques et formes quadratiques	79
Chapitre 5	
Nombres réels	83
Ecriture décimale des nombres réels	83
Différents ensembles de nombres réels	84
Ordre canonique de $\mathbb{R}$	87
Topologie de $\mathbb{R}$	88
Distance d'un point à une partie et distance entre deux parties	93
Endomorphismes du groupe additif $(\mathbb{R}, +)$ de $\mathbb{R}$	94
Mesurabilité de parties de $\mathbb{R}$	96
Chapitre 6	
Suites numériques	100
Convergence et divergence	100
Convergence au sens de Cesàro	104
Limite supérieure et limite inférieure	105
Valeurs d'adhérence	107
Chapitre 7	
Séries numériques	110
Convergence et convergence absolue	110
Mise en défaut de certains critères de convergence	114
Séries à termes positifs	117
Règles de convergence	118
Comparaison série-intégrale	122
Modification de l'ordre des termes	122
Séries doubles et produit de Cauchy	124
Équivalence des suites des sommes partielles et équivalence des suites des restes	127
Autres types de convergence	129
Chapitre 8	
Fonctions d'une variable réelle : continuité et limites	133
Continuité	133
Théorème des valeurs intermédiaires	137
Limites	140
Continuité uniforme, continuité absolue et fonctions lipschitziennes	143
Continuité de l'application réciproque d'une fonction continue	149
Continuité et topologie	151
Chapitre 9	
Fonctions d'une variable réelle : dérivabilité	155
Dérivabilité locale et globale	155
Discontinuité de la fonction dérivée	165
Sens de variation d'une fonction dérivable	169
Dérivées et limites	172
Dérivées et extremums	175
Fonctions indéfiniment dérivables	176
Développements limités	178
Dérivées supérieures et inférieures	181
Equations différentielles	183

Chapitre 10

Fonctions d'une variable réelle monotones, périodiques, convexes, bornées	190
Monotonie et continuité	190
Fonctions périodiques	196
Fonctions convexes	197
Fonctions bornées	201
Fonctions à variation bornée	202

Chapitre 11

Intégration	207
Intégrale de Riemann	207
Convergence des intégrales	213
Primitives et intégrales	219
Intégrales dépendant d'un paramètre	223
Sommes de Riemann	226
Intégration des relations de comparaison	228

Chapitre 12

Suites de fonctions	231
Convergence simple	231
Convergence uniforme	235
Dérivation	240
Intégration	243
Convergence en moyenne	245

Chapitre 13

Séries de fonctions	250
Différents types de convergence	250
Discontinuité de la somme d'une série de fonctions	256
Interversion de sommations et de limites	258
Séries entières	260
Série de Taylor	262
Séries de Fourier	264

Chapitre 14

Fonctions de plusieurs variables	268
Continuité	269
Différentiabilité	270
Théorèmes d'inversion	277
Dérivées partielles secondes	279
Extremums	281
Intégration	282

Chapitre 15

Topologie générale	288
Séparation	289
Suites convergentes, limites de suites, continuité	291
Connexité	295
Compacité	301

Chapitre 16	
Espaces métriques	305
Boules	306
Equivalences de distances	310
Complétude	312
Distance d'un point à une partie et distance entre deux parties	315
Chapitre 17	
Espaces vectoriels normés	317
Nécessité des axiomes	318
Somme de deux parties	320
Comparaison des normes	321
Continuité des applications linéaires	322
Parties compactes et parties fermées et bornées	325
Parties convexes d'un espace vectoriel normé	327
Points internes à une partie	329
Isométries	330
Espaces vectoriels euclidiens et hilbertiens	332
Chapitre 18	
Courbes planes	338
Tangentes et points d'inflexion	339
Directions asymptotiques et asymptotes	344
Longueur d'une courbe	347
Courbes remarquables	350
Chapitre 19	
Probabilités	355
Evénements indépendants	356
Espérance mathématique et variance	357
Convergence des suites de variables aléatoires	360

Chapitre 1

Logique, ensembles, arithmétique

La logique mathématique prend naissance avec George Boole au milieu du XIX^e siècle et se formalise aux alentours de 1900 alors qu'apparaît la théorie des ensembles sous l'impulsion de Georg Cantor (1845-1918) et Richard Dedekind (1831-1916). Elle permet de fonder rigoureusement les mathématiques, tâche qui devient nécessaire lorsque les mathématiques se détachent de la physique. Quant à l'arithmétique, elle est aussi ancienne que les mathématiques puisqu'elle débute dès l'Antiquité ; certaines conjectures en arithmétique se sont révélées fausses, comme nous le verrons grâce à quelques contre-exemples.

■ Logique des prédicats du premier ordre

Les mathématiques se fondent en particulier sur la logique des prédicats du premier ordre, construite à l'aide des connecteurs propositionnels « non », « et », « ou », « implique » et « équivalent », de variables $x, y, z, \dots$, de propositions, de prédicats $P(x), Q(x, y), \dots$ et des quantificateurs existentiel $\exists$ et universel $\forall$.

Si l'on souhaite utiliser des quantificateurs dans un texte mathématique — ce n'est jamais indispensable... —, leur maniement doit se faire avec soin. On peut intervertir deux quantificateurs de même nature se suivant dans une formule, mais il n'en est pas de même pour deux quantificateurs de natures différentes.

1.1. Exemple où la proposition $\forall x \exists y (P(x, y))$ est vraie alors que l'assertion $\exists y \forall x (P(x, y))$ est fausse.

Nous prenons $\mathbb{N}$ comme ensemble de référence et nous symbolisons par $P(x, y)$ l'assertion : x est inférieur ou égal à y . Alors, pour tout entier naturel x , si l'on choisit un entier naturel y supérieur ou égal à x , l'assertion $P(x, y)$ est vraie ; la proposition $\forall x \exists y (P(x, y))$ est donc vraie. Par contre, $\mathbb{N}$ n'ayant pas de plus grand élément, l'assertion $\exists y \forall x (P(x, y))$ est fausse. $\square$

Remarquons que si la deuxième des formules de l'exemple 1.1 est vraie, alors la première l'est, c'est-à-dire que l'implication $\exists y \forall x (P(x, y)) \implies \forall x \exists y (P(x, y))$ est universellement valide ; intuitivement nous voyons que dans le terme de gauche y ne dépend pas de x , alors que dans celui de droite il peut en dépendre : la condition est donc moins forte. C'est ainsi qu'une utilisation de l'intervention de quantificateurs de natures différentes est faite pour passer de la continuité à la

continuité uniforme¹ : $\mathcal{A}$ étant une partie de $\mathbb{R}$ et f une fonction de $\mathbb{R}$ dans $\mathbb{R}$ définie sur $\mathcal{A}$, la continuité de f sur $\mathcal{A}$ s'exprime par la proposition :

$$(\forall \varepsilon > 0)(\forall x \in \mathcal{A})(\exists \eta > 0)(\forall y \in \mathcal{A})(|x - y| < \eta \implies |f(x) - f(y)| < \varepsilon)$$

et la continuité uniforme de f sur $\mathcal{A}$ par l'assertion :

$$(\forall \varepsilon > 0)(\exists \eta > 0)(\forall x \in \mathcal{A})(\forall y \in \mathcal{A})(|x - y| < \eta \implies |f(x) - f(y)| < \varepsilon).$$

Les connecteurs «et» et «ou» étant symbolisés respectivement par $\wedge$ et $\vee$, les propositions $\forall x (P(x) \wedge Q(x))$ et $\forall x (P(x)) \wedge \forall x (Q(x))$ sont équivalentes, de même que les assertions $\exists x (P(x) \vee Q(x))$ et $\exists x (P(x)) \vee \exists x (Q(x))$.

1.2. Exemple où l'assertion $\forall x (P(x) \vee Q(x))$ est vraie alors que la proposition $\forall x (P(x)) \vee \forall x (Q(x))$ est fausse.

Nous prenons de nouveau $\mathbb{N}$ comme ensemble de référence et nous symbolisons, pour tout entier naturel x , par $P(x)$ l'assertion : x est pair, et par $Q(x)$ la proposition : x est impair. Alors $\forall x (P(x) \vee Q(x))$ exprime que tout entier naturel est pair ou impair, ce qui est vrai ; par contre $\forall x (P(x)) \vee \forall x (Q(x))$ exprime que soit tous les entiers sont pairs, soit ils sont tous impairs, ce qui est bien sûr faux. $\square$

1.3. Exemple où la proposition $\exists x (P(x)) \wedge \exists x (Q(x))$ est vraie alors que l'assertion $\exists x (P(x) \wedge Q(x))$ est fausse.

Les hypothèses sont celles l'exemple 1.2. Alors $\exists x (P(x)) \wedge \exists x (Q(x))$ exprime qu'il existe (au moins) un entier pair et (au moins) un entier impair, alors que $\exists x (P(x) \wedge Q(x))$ exprime qu'il existe un entier simultanément pair et impair. $\square$

Remarquons que dans la première des propositions de l'exemple précédent 1.3, on peut remplacer la deuxième occurrence de x par y , car x est ici une variable muette, et que donc un élément qui vérifie $\exists x (P(x))$ n'est pas forcément le même qu'un élément vérifiant $\exists x (Q(x))$; par contre, dans la deuxième proposition, c'est le même élément qui doit remplir les deux conditions.

■ Paradoxes de la théorie “naïve” des ensembles

La théorie des ensembles, œuvre des mathématiciens allemands Georg Cantor et Richard Dedekind, apparaît à la fin du XIX^e siècle. Dans cette première approche, que l'on qualifie maintenant de «naïve», on appelle ensemble n'importe quelle collection d'objets², ce qui conduit à des paradoxes.

1.4. Paradoxe de Russel³.

Nous considérons l'ensemble $A = \{ X \mid X \notin X \}$. Alors l'objet A appartient ou bien n'appartient pas à l'ensemble A . Si A appartient à A , alors, par définition de A , on obtient que A n'appartient pas à A ; si A n'appartient pas à A , cette même

1. Voir le chapitre 8, page 144.

2. Georg Cantor débute son mémoire de mars 1895, publié dans les *Mathematische Annalen*, ainsi : « Nous appelons “ensemble” toute réunion M d'objets de notre conception m , déterminés et bien distincts, et que nous appelons “éléments” de M . »

3. Il est énoncé en 1905 par le mathématicien et philosophe anglais Bertrand Russel (1872-1970).

définition nous permet d'affirmer que A appartient à A . Ainsi les deux assertions « A appartient à A » et « A n'appartient pas à A » sont simultanément vraies... $\square$

Nous rappelons que l'ensemble des parties d'un ensemble E est noté $\mathcal{P}(E)$.

THÉORÈME 1.1. — Théorème de Cantor.

Si E est un ensemble, il n'existe aucune injection de $\mathcal{P}(E)$ dans E .

1.5. Paradoxe de Cantor.

En notant E l'ensemble de tous les ensembles, l'ensemble $\mathcal{P}(E)$ est inclus dans E donc l'injection canonique de $\mathcal{P}(E)$ dans E contredit le théorème de Cantor. $\square$

Pour remédier à de tels paradoxes, une théorie axiomatique des ensembles a été élaborée. C'est la théorie des ensembles Zermelo-Fraenkel, en abrégé ZF, ainsi dénommée car elle a été conçue par Ernst Zermelo⁴ en 1908 — il l'expose dans ses *Recherches sur les fondements de la théorie des ensembles* — et modifiée par Abraham Fraenkel⁵ en 1921 et 1922. La théorie ZF sert depuis cette époque de fondement aux mathématiques, dont elle permet une construction rigoureuse, même si sa consistance, c'est-à-dire l'absence de paradoxes, ne pourra jamais être prouvée, comme le montre Kurt Gödel⁶ en 1931. Les mathématiciens d'aujourd'hui font le pari de cette consistance et fondent les mathématiques sur la théorie ZF.

■ Image directe et image réciproque d'une partie

Ayant surtout en vue de l'étude de la notion de cardinal, Georg Cantor utilise essentiellement des correspondances biunivoques entre les ensembles, ce que nous appelons maintenant des bijections, alors que Richard Dedekind introduit dans toute sa généralité la notion d'application d'un ensemble dans un autre.

Soit f une application d'un ensemble E dans un ensemble F .

Si A est une partie de l'ensemble E , l'image directe de A par f est la partie $f(A) = \{f(x) \mid x \in A\}$ de F et, si M est une partie de F , l'image réciproque de M par f est la partie $f^{-1}(M) = \{x \in E \mid f(x) \in M\}$ de E .

Si A et B sont des parties de E , alors $f(A \cup B) = f(A) \cup f(B)$, mais en général on a seulement l'inclusion $f(A \cap B) \subset f(A) \cap f(B)$. Cependant, si l'application f est injective, cette inclusion devient une égalité.

1.6. Parties A et B de E telles que $f(A \cap B) \neq f(A) \cap f(B)$.

Nous introduisons l'application $f : n \mapsto f(n) = |n|$ de $\mathbb{Z}$ dans $\mathbb{Z}$ et les parties $A = \mathbb{N}$ et $B = -\mathbb{N}$ de $\mathbb{Z}$. Alors $A \cap B = \{0\}$ donc $f(A \cap B) = \{0\}$. Par contre $f(A) = f(B) = \mathbb{N}$ donc $f(A) \cap f(B) = \mathbb{N}$. $\square$

Pour une partie A de E , une partie B de F et une application f de E dans F , on a les inclusions $A \subset f^{-1}(f(A))$ et $f(f^{-1}(B)) \subset B$. La première est une égalité si f est injective et la seconde si f est une surjection de E sur F .

4. Mathématicien et logicien allemand (1871-1953).

5. Mathématicien et logicien israélien d'origine allemande (1891-1965).

6. Mathématicien et logicien autrichien (1906-1978).

1.7. Partie A de E telle que $A \neq f^{-1}(f(A))$.

Nous utilisons de nouveau l'application $f : n \mapsto f(n) = |n|$ de $\mathbb{Z}$ dans $\mathbb{Z}$ et nous posons $A = \mathbb{N}$. Alors $f(A) = \mathbb{N} = f(\mathbb{Z})$ donc $f^{-1}(f(A)) = \mathbb{Z} \neq A$. $\square$

1.8. Partie M de F telle que $M \neq f(f^{-1}(M))$.

Pour l'application $f : n \mapsto f(n) = |n|$ de $\mathbb{Z}$ dans $\mathbb{Z}$ et $M = -\mathbb{N}$, $f^{-1}(M) = \{0\}$ — seul 0 s'envoie sur un élément de $-\mathbb{N}$ —, donc $f(f^{-1}(M)) = \{0\} \neq M$. $\square$

On a, pour toute partie M de F , $f^{-1}(\mathbb{C}_F M) = \mathbb{C}_E f^{-1}(M)$. Ceci est faux pour l'image directe et il n'y a dans ce cas d'inclusion ni dans un sens, ni dans l'autre.

1.9. Partie A de E telle que $f(\mathbb{C}_E A) \neq \mathbb{C}_F f(A)$.

Nous considérons de nouveau l'application $f : n \mapsto f(n) = |n|$ de $\mathbb{Z}$ dans $\mathbb{Z}$ et nous posons $A = \mathbb{N}$. Alors $\mathbb{C}_E A = -\mathbb{N}^*$, ensemble des entiers strictement négatifs, donc $f(\mathbb{C}_E A) = \mathbb{N}^*$. Par ailleurs $f(A) = \mathbb{N}$, donc $\mathbb{C}_F f(A) = -\mathbb{N}^*$; non seulement les deux ensembles diffèrent mais ils sont disjoints. $\square$

■ Ensembles équipotents

DÉFINITION 1.1. — Si E et F sont des ensembles, E est équipotent à F s'il existe une bijection de E sur F .

La relation d'équipotence, symbolisée par $\mathcal{E}q$, est une relation réflexive, transitive et symétrique sur la « classe » — et non l'ensemble, voir l'exemple 1.5 — de tous les ensembles, ce qui signifie que, quels que soient les ensembles X , Y et Z , $X \mathcal{E}q X$ (réflexivité), $X \mathcal{E}q Y$ et $Y \mathcal{E}q Z$ entraînent $X \mathcal{E}q Z$ (transitivité), et $X \mathcal{E}q Y$ entraîne $Y \mathcal{E}q X$ (symétrie). Compte tenu de la symétrie, on dit que les ensembles E et F sont équipotents pour exprimer que E (ou F) est équipotent à F (ou à E).

Intuitivement, des ensembles E et F sont équipotents s'ils ont « le même nombre d'éléments », ce qui conduit à la notion de cardinal d'un ensemble, développée par Cantor parallèlement à la notion d'ordinal — les cardinaux servent à « compter » le nombre des éléments d'un ensemble et les ordinaux à les « numéroter ».

Tous les ensembles ont un cardinal, des ensembles sont équipotents si, et seulement si, ils ont le même cardinal et le cardinal d'un ensemble fini est le nombre de ses éléments au sens traditionnel. Le cardinal d'un ensemble E se note $\text{card}(E)$.

Euclide cite parmi ses axiomes : « *La partie est toujours plus petite que le tout* ». Ce n'est en fait vrai que pour les ensembles finis : dans un ensemble infini E , il existe une partie stricte de même cardinal que E , c'est-à-dire équipotente à E .

1.10. Partie stricte de $\mathbb{N}$ équipotente à $\mathbb{N}$.

L'ensemble $\mathbb{P}$ des nombres entiers naturels pairs est une partie stricte de $\mathbb{N}$ et l'application $f : n \mapsto f(n) = 2n$ une bijection de $\mathbb{N}$ sur $\mathbb{P}$, donc $\mathbb{N}$ et $\mathbb{P}$ sont des ensembles équipotents⁷. $\square$

7. Galilée avait déjà remarqué que l'application $\varphi : n \mapsto \varphi(n) = n^2$ de $\mathbb{N}$ dans $\mathbb{N}$ est une bijection de $\mathbb{N}$ sur l'ensemble $\mathcal{C}$ des carrés d'entiers naturels, partie stricte de $\mathbb{N}$.

1.11. Partie stricte de $\mathbb{R}$ équipotente à $\mathbb{R}$.

Pour tout réel x , $\left| \frac{x}{1+|x|} \right| = \frac{|x|}{1+|x|} < 1$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow]-1, 1[\\ x \longmapsto f(x) = \frac{x}{1+|x|} \end{array} \right.$$

Nous remarquons que, pour tout réel x , les réels x et $f(x)$ sont de même signe. Si x et y sont des nombres réels et si $f(x) = f(y)$, x et y sont de même signe donc, si $x \geq 0$, $x/(1+x) = y/(1+y)$ et, si $x < 0$, $x/(1-x) = y/(1-y)$, d'où l'on déduit que $x = y$. Si c appartient à l'intervalle ouvert $] -1, 1[$, alors, en notant a la solution de l'équation $x/(1+x) = c$ si $c \geq 0$ et la solution de $x/(1-x) = c$ si $c < 0$, a est un antécédent de c par f . Ainsi f est une bijection de $\mathbb{R}$ sur $] -1, 1[$. En conclusion, la partie stricte $] -1, 1[$ de $\mathbb{R}$ est équipotente à $\mathbb{R}$. $\square$

Remarquons que si a et b sont des réels et si $a < b$, l'intervalle ouvert $]a, b[$ est équipotent à $\mathbb{R}$. En effet, en notant φ l'unique application affine de $\mathbb{R}$ dans $\mathbb{R}$ telle que $\varphi(-1) = a$ et $\varphi(1) = b$ — c'est l'application $\varphi : x \mapsto \alpha x + \beta$ où $\alpha = (b-a)/2$ et $\beta = (b+a)/2$ — la restriction de φ à $] -1, 1[$ est une bijection de $] -1, 1[$ sur $]a, b[$, donc $]a, b[$ est équipotent à $] -1, 1[$ et, comme l'intervalle $] -1, 1[$ est équipotent à $\mathbb{R}$ — voir l'exemple précédent 1.11 —, les ensembles $]a, b[$ et $\mathbb{R}$ sont équipotents.

On peut remplacer l'application f de l'exemple 1.11 par la fonction « tangente hyperbolique » ou l'application $g : x \mapsto g(x) = (2/\pi) \arctan x$ de $\mathbb{R}$ dans $] -1, 1[$.

L'ensemble $\mathbb{N}$ paraît avoir « moins d'éléments » que $\mathbb{N}^2$. En fait il n'en est rien : ces ensembles sont équipotents.

1.12. Bijection de $\mathbb{N}^2$ sur $\mathbb{N}$.

Nous introduisons la suite $(a_n)_{n \in \mathbb{N}}$ d'entiers naturels, de terme général :

$$a_n = \sum_{i=0}^n i = 1 + 2 + \cdots + n = \frac{n(n+1)}{2} \quad (\text{quotient exact}).$$

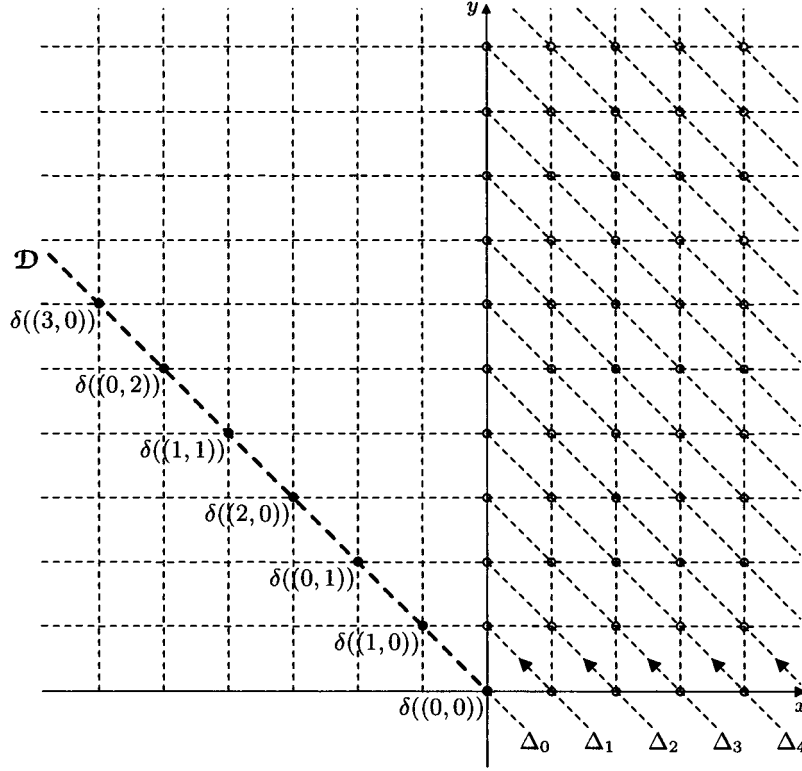
On a, pour tout entier naturel n , $a_{n+1} = a_n + (n+1)$; en particulier, $(a_n)_{n \in \mathbb{N}}$ est strictement croissante. Nous déduisons de la suite $(a_n)_{n \in \mathbb{N}}$ l'application :

$$\left| \begin{array}{l} f : \mathbb{N}^2 \longrightarrow \mathbb{N} \\ (p, q) \longmapsto f((p, q)) = a_{p+q} + q. \end{array} \right.$$

Soit $m \in \mathbb{N}$. Comme $a_0 = 0 \leq m \leq a_m$, l'ensemble $M = \{k \in \mathbb{N} \mid a_k \leq m\}$ est une partie non vide de $\mathbb{N}$ majorée dans $\mathbb{N}$ (par m) donc M admet un plus grand élément n , et n est l'unique entier naturel tel que $a_n \leq m < a_{n+1}$. Si (p, q) est un antécédent de m par f , alors $m = f((p, q)) = a_{p+q} + q$ et $a_{p+q+1} = a_{p+q} + (p+q+1)$, donc $a_{p+q} \leq m < a_{p+q+1}$, ce qui montre que $p+q = n$ et $q = m - a_n$, d'où l'unicité du couple (p, q) . Inversement, $m - a_n$ et $n - m + a_n$ sont des entiers naturels de somme n et $f(n - m + a_n, m - a_n) = a_n + (m - a_n) = m$, donc $(n - m + a_n, m - a_n)$ est un antécédent de m par f . En conclusion, f est une bijection de $\mathbb{N}^2$ sur $\mathbb{N}$. $\square$

Nous gardons les définitions et les notations de l'exemple précédent 1.12 et nous posons, pour tout entier naturel n , $\Delta_n = \{(p, q) \mid p, q \in \mathbb{N} \text{ et } p + q = n\}$,

c'est-à-dire $\Delta_n = \{(n, 0), (n-1, 1), \dots, (n-k, k), \dots, (1, n-1), (0, n)\}$, ensemble fini de cardinal $n+1$ ordonné par les deuxièmes projections des couples — voir les flèches du dessin ci-dessous sur $\Delta_1, \Delta_2, \Delta_3 \dots$. Sur ce dessin, $\mathbb{N}^2$ est représenté dans le plan $\mathbb{R}^2$, les couples (p, q) pour $p, q \in \mathbb{N}$ étant figurés par de petits cercles.



La famille $(\Delta_n)_{n \in \mathbb{N}}$ est une partition de $\mathbb{N}^2$ et la famille $([a_n, a_{n+1} - 1])_{n \in \mathbb{N}}$ une partition de $\mathbb{N}$, et, pour tout $n \in \mathbb{N}$, la restriction f_n de f à Δ_n est la bijection de Δ_n sur $[a_n, a_{n+1} - 1]$ définie par $f_n((n-k, k)) = a_n + k$ pour tout $k \in [0, n]$. Nous notons, pour tout entier naturel n , δ_n la restriction à Δ_n de la translation τ_n de $\mathbb{R}^2$ de vecteur $(-n - a_n, a_n)$, qui transforme $(n, 0)$ en $(-a_n, a_n)$, et nous posons $\mathcal{D} = \{(-n, n) \mid n \in \mathbb{N}\}$, ensemble évidemment équipotent à $\mathbb{N}$. Alors, en «recollant» la famille $(\delta_n)_{n \in \mathbb{N}}$, on obtient une bijection δ de $\mathbb{N}^2$ sur $\mathcal{D}$ qui illustre de façon beaucoup plus visuelle que f l'équipotence entre $\mathbb{N}^2$ et $\mathbb{N}$.

Ces considérations montrent que l'équipotence entre $\mathbb{N}^2$ et $\mathbb{N}$ est bien moins étonnante qu'il n'y paraît au premier abord : nous avons transformé facilement, par une bijection, $\mathbb{N}^2$ en une partie $\mathcal{D}$ de $\mathbb{R}^2$ équipotente à $\mathbb{N}$. Il paraît en revanche impossible de transformer par une bijection le plan $\mathbb{R}^2$ (une plaque d'aire infinie) en une droite de $\mathbb{R}^2$ (un fil de fer). Pourtant Cantor démontre, en juin 1877, que $\mathbb{R}^2$ est équipotent à $\mathbb{R}$ et il en est stupéfait⁸. On sait maintenant que tout ensemble infini E est équipotent à son carré $E^2 = E \times E$.

8. Dans une lettre à Dedekind datée du 29 juin 1877, il écrit, en français dans le texte : *Je le vois, mais je ne le crois pas.*

THÉORÈME 1.2. — Théorème de Cantor-Bernstein⁹.

Si E et F sont des ensembles et s'il existe une injection de E dans F et une injection de F dans E , alors les ensembles E et F sont équipotents.

DÉFINITION 1.2. — Un ensemble est dénombrable s'il est équipotent à $\mathbb{N}$.

L'ensemble E est donc dénombrable si, et seulement si, il existe une suite $(a_n)_{n \in \mathbb{N}}$ d'éléments deux à deux distincts de E telle que $E = \{a_n \mid n \in \mathbb{N}\} = \{a_0, a_1, a_2, \dots\}$, ce qui signifie que l'on peut « numérotter » les éléments de E par les entiers naturels.

Les ensembles $\mathbb{Z}$ et $\mathbb{N}^2$, l'ensemble $\mathbb{Q}$ des nombres rationnels et l'ensemble $\mathbb{A}$ des nombre algébriques¹⁰ sont dénombrables. Pour $\mathbb{N}^2$ une preuve en est faite dans l'exemple précédent 1.12, mais on peut aussi raisonner de la manière suivante : l'application $n \mapsto (n, 0)$ est une injection de $\mathbb{N}$ dans $\mathbb{N}^2$ et, 2 et 3 étant des nombres premiers, l'application $(k, \ell) \mapsto 2^k 3^\ell$ est une injection de $\mathbb{N}^2$ dans $\mathbb{N}$, donc on déduit du théorème de Cantor-Bernstein que $\mathbb{N}$ et $\mathbb{N}^2$ sont équipotents. Un raisonnement analogue, utilisant n nombres premiers $p_1, \dots, p_n$ deux à deux distincts, permet de prouver que, pour tout entier $n \geq 2$, les ensembles $\mathbb{N}$ et $\mathbb{N}^n$ sont équipotents.

Les ensembles $\mathbb{Z}$, puis $\mathbb{Q}$, les $\mathbb{N}^n$ pour $n \geq 2$ et $\mathbb{A}$, sont « de plus en plus grands » mais ils sont tous dénombrables. On pourrait donc penser qu'il en est de même pour tous les ensembles infinis, ce qui est faux comme le montre l'exemple suivant.

1.13. Ensemble infini qui n'est pas dénombrable.

Nous démontrons que l'ensemble infini $\mathbb{R}$ n'est pas dénombrable. La preuve la plus souvent donnée de ce résultat utilise les développements décimaux illimités. Cependant la première démonstration de cette assertion par Cantor est plus simple et elle met bien en évidence la propriété fondamentale du corps des nombres réels qui est mise en œuvre.

Nous débutons par une remarque : si a, b et x sont des réels et si $a < b$, on peut construire un couple (a', b') de nombres réels tel que $a \leq a' < b' \leq b$ et $x \notin [a', b']$ — en posant par exemple $c = a + (b - a)/2$, $d = a + (b - a)/3$ et : $a' = c$ et $b' = b$ si $x < c$; $a' = a$ et $b' = d$ si $x \geq c$.

Supposons $\mathbb{R}$ dénombrable. Alors $\mathbb{R} = \{\omega_n \mid n \in \mathbb{N}\} = \{\omega_0, \omega_1, \omega_2, \dots, \omega_n, \dots\}$ où $(\omega_n)_{n \in \mathbb{N}}$ est une suite de réels deux à deux distincts. On construit les suites $(a_n)_{n \in \mathbb{N}}$ et $(b_n)_{n \in \mathbb{N}}$ de réels, par récurrence sur $n \in \mathbb{N}$, en posant $a_0 = 0$ et $b_0 = 1$ et en notant, pour tout entier naturel n , (a_{n+1}, b_{n+1}) le couple de réels tel que $a_n \leq a_{n+1} < b_{n+1} \leq b_n$ et $\omega_n \notin [a_{n+1}, b_{n+1}]$ obtenu dans la remarque ci-dessus. Les suites (a_n) et (b_n) convergent dans $\mathbb{R}$ et, en notant λ la limite de (a_n) et μ celle de (b_n) , on a $\lambda \leq \mu$ et, pour tout entier naturel n , $a_n \leq \lambda \leq \mu \leq b_n$. De plus il existe $p \in \mathbb{N}$ tel que $\lambda = \omega_p$; or on a $a_{p+1} \leq \lambda \leq b_{p+1}$, ce qui contredit l'assertion $\omega_p \notin [a_{p+1}, b_{p+1}]$. En conclusion, $\mathbb{R}$ n'est pas dénombrable. $\square$

9. Ce théorème est énoncé par Cantor qui ne réussit pas à le démontrer. Indépendamment l'un de l'autre, Ernst Schröder (1841-1902) en 1896 et Felix Bernstein (1878-1956) en 1897 en donnent une démonstration. Dedekind en expose une preuve en 1899, dans une lettre à Cantor, preuve rédigée en 1887 mais qu'il ne publie pas, en raison des doutes qui l'assaillent de plus en plus quant à la validité de la théorie des ensembles telle que créée par lui-même et Cantor, et qu'il développe dans son ouvrage « *Was sind und was sollen die Zahlen ?* » — en français « *Que sont et que représentent les nombres ?* » — élaboré entre 1872 et 1878 et publié seulement en 1888.

10. Voir la définition 5.1, page 84.

1.14. Autre ensemble infini qui n'est pas dénombrable.

Supposons que l'ensemble infini $\mathcal{A} = \mathcal{F}(\mathbb{N}, \mathbb{N})$ des applications de $\mathbb{N}$ dans $\mathbb{N}$ soit dénombrable. Alors $\mathcal{A} = \{\varphi_n \mid n \in \mathbb{N}\} = \{\varphi_0, \varphi_1, \varphi_2, \dots, \varphi_n, \dots\}$ où $(\varphi_n)_{n \in \mathbb{N}}$ est une suite d'applications deux à deux distinctes de $\mathbb{N}$ dans $\mathbb{N}$. Nous introduisons l'application $f : n \mapsto f(n) = 1 + \varphi_n(n)$ de $\mathbb{N}$ dans $\mathbb{N}$. Comme $f \in \mathcal{A}$, il existe un entier naturel p tel que $f = \varphi_p$. On a, pour tout $n \in \mathbb{N}$, $\varphi_p(n) = f(n)$, donc en particulier $\varphi_p(p) = f(p) = 1 + \varphi_p(p)$, ce qui est évidemment faux. Par conséquent l'ensemble $\mathcal{A} = \mathcal{F}(\mathbb{N}, \mathbb{N})$ n'est pas dénombrable. $\square$

La méthode utilisée dans l'exemple précédent 1.14 est *le procédé diagonal de Cantor* et c'est avec ce procédé, appliqué aux suites des décimales des points de $]0, 1[$, supposé dénombrable, donc s'écrivant $\{a_n \mid n \in \mathbb{N}\}$ où $(a_n)_{n \in \mathbb{N}}$ est une suite de réels deux à deux distincts, que l'on peut prouver que $\mathbb{R}$ n'est pas dénombrable.

■ Relations binaires fondamentales

Parmi les relations binaires sur un ensemble¹¹, les relations d'équivalence sont essentielles, car elles permettent d'identifier les éléments jouant un même rôle dans une structure, ainsi que les relations d'ordre, indispensables à l'étude des nombres réels et plus généralement à l'analyse, mais aussi, par l'intermédiaire de la notion de bon ordre — voir dans la suite de ce chapitre — à la théorie des ensembles.

DÉFINITION 1.3. — Une relation d'équivalence sur un ensemble E est une relation binaire réflexive, transitive et symétrique sur E .

Montrons que chacune des trois propriétés est nécessaire, car deux d'entre elles n'entraînent pas nécessairement la troisième. Remarquons tout d'abord que la relation de divisibilité sur $\mathbb{N}$ est réflexive et transitive mais n'est pas symétrique. Remarquons aussi qu'un raisonnement fallacieux pourrait faire penser qu'une relation $\mathcal{R}$ symétrique et transitive est nécessairement réflexive, car $a \mathcal{R} b$ entraîne $b \mathcal{R} a$ par symétrie et $a \mathcal{R} a$ par transitivité ; cependant un élément peut n'être en relation avec aucun autre, comme dans l'exemple qui suit.

1.15. Relation symétrique et transitive qui n'est pas réflexive.

Définissons sur l'ensemble $\mathbb{Z}$ des entiers relatifs la relation binaire $\mathcal{R}$ par : $a \mathcal{R} b$ si $ab > 0$. La relation $\mathcal{R}$ est clairement symétrique. Si $a, b, c \in \mathbb{Z}$ et si $a \mathcal{R} b$ et $b \mathcal{R} c$, alors $ab > 0$ et $bc > 0$ donc, en multipliant membre à membre, $ab^2c > 0$, et comme $b^2 > 0$ — en effet, $ab > 0$, donc $b \neq 0$ —, on a $ac > 0$, donc $a \mathcal{R} c$: $\mathcal{R}$ est transitive. Cependant l'assertion $0 \mathcal{R} 0$ est fausse, donc $\mathcal{R}$ n'est pas réflexive. $\square$

1.16. Relation réflexive et symétrique qui n'est pas transitive.

Nous posons $N = \{n \in \mathbb{N} \mid n \geq 2\}$ et nous introduisons la relation binaire $\mathcal{R}$ définie sur N par : $a \mathcal{R} b$ si a et b ont en commun un diviseur premier. Tout élément de N possédant au moins un diviseur premier, $\mathcal{R}$ est réflexive ; la symétrie de $\mathcal{R}$ est

11. Pour les définitions, voir [RAM1], chapitre 1, §1.3.

claire. Cependant, 2 divise 2 et 12 donc $2 \mathcal{R} 12$, et 3 divise 12 et 3 donc $12 \mathcal{R} 3$, alors que l'assertion $2 \mathcal{R} 3$ est fausse, ce qui montre que $\mathcal{R}$ n'est pas transitive. $\square$

DÉFINITION 1.4. — Une relation d'ordre sur un ensemble E est une relation binaire réflexive, transitive et antisymétrique sur E .

Nous prouvons, comme dans le cas des relations d'équivalence, que chacune des trois propriétés est nécessaire. Remarquons d'abord que la relation de divisibilité sur $\mathbb{Z}$ est réflexive et transitive mais n'est pas antisymétrique.

1.17. Relation réflexive et antisymétrique qui n'est pas transitive.

Nous définissons la relation binaire $\mathcal{R}$ sur $\mathbb{N}^*$ par : $a \mathcal{R} b$ si a divise b et b/a est 1 ou un nombre premier. Pour tout $a \in \mathbb{N}^*$, $a/a = 1$ donc $a \mathcal{R} a$: la relation $\mathcal{R}$ est réflexive. Si $a, b \in \mathbb{N}^*$ et si $a \mathcal{R} b$ et $b \mathcal{R} a$, alors a divise b et b divise a donc $a = b$, ce qui montre l'antisymétrie de $\mathcal{R}$. Par ailleurs, $6/2 = 3$ et $12/6 = 2$ sont des nombres premiers, donc $2 \mathcal{R} 6$ et $6 \mathcal{R} 12$, alors que l'assertion $2 \mathcal{R} 12$ est fausse — en effet, $12/2 = 6$ n'est ni 1 ni un nombre premier — d'où l'on déduit que la relation $\mathcal{R}$ n'est pas transitive. $\square$

1.18. Relation antisymétrique et transitive qui n'est pas réflexive.

Nous définissons la relation binaire $\mathcal{R}$ sur $\mathbb{N}$ par : $a \mathcal{R} b$ si a est pair et divise b . Si $a, b \in \mathbb{N}$ et si a divise b et b divise a , alors $a = b$: $\mathcal{R}$ est antisymétrique. Si $a, b, c \in \mathbb{N}$ et si $a \mathcal{R} b$ et $b \mathcal{R} c$, a et b sont pairs et, par transitivité de la relation de divisibilité, a divise c , donc $a \mathcal{R} c$, ce qui prouve la transitivité de $\mathcal{R}$. Cependant, si a est impair, la proposition $a \mathcal{R} a$ est fausse ; par suite $\mathcal{R}$ n'est pas réflexive. $\square$

■ Ensembles ordonnés

DÉFINITION 1.5. — Un ensemble ordonné est un couple $(E, \leq)$ où E est un ensemble et $\leq$ une relation d'ordre sur E .

Dans les définitions qui suivent, $(E, \leq)$ est un ensemble ordonné.

DÉFINITION 1.6. — La relation $\leq$ est une relation d'ordre total sur E , ou $(E, \leq)$ est un ensemble totalement ordonné, si, quels que soient les éléments x et y de E , $x \leq y$ ou $y \leq x$. Dans le cas contraire, on dit que $\leq$ est une relation d'ordre partiel sur E ou que $(E, \leq)$ est un ensemble partiellement ordonné.

DÉFINITION 1.7. — Soit A une partie de E .

a) Un majorant de A est un élément ν de E tel que $x \leq \nu$ pour tout élément x de A , un minorant de A un élément μ de E tel que $x \geq \mu$ pour tout élément x de A et on dit qu'un élément λ de E majore A (resp. minore A) si λ est un majorant (resp. un minorant) de A .

b) Le plus grand élément de A , appelé aussi l'élément maximum de A , est, s'il existe, l'unique majorant de A appartenant à A , et le plus petit élément de A , appelé aussi l'élément minimum de A , est, s'il existe, l'unique minorant de A qui appartient à A .

Si une partie A de E possède un plus grand élément (resp. un plus petit élément), celui-ci est le plus petit des majorants (resp. le plus grand des minorants) de A .

DÉFINITION 1.8. — Si A est une partie de E , la borne supérieure de A est, à condition qu'il existe, le plus petit des majorants de A , et la borne inférieure de A est, s'il existe, le plus grand des minorants de A .

Par conséquent, si une partie A de E admet un plus grand (resp. un plus petit) élément a , a est la borne supérieure (resp. la borne inférieure) de A .

1.19. Partie majorée admettant une borne supérieure mais n'ayant pas plus grand élément.

Nous munissons $E = \{(0, 1)\} \times \mathbb{N}$ de l'ordre lexicographique $\leq_L$, c'est-à-dire de la relation d'ordre définie sur E par : $(k, n) \leq_L (\ell, m)$ si $k < \ell$ ou $(k = \ell \text{ et } n \leq m)$. Nous posons $A = \{(0, n) \mid n \in \mathbb{N}\}$. Si $p \in \mathbb{N}$, $(0, p)$ et $(0, p + 1)$ appartiennent à A et $(0, p) <_L (0, p + 1)$, donc $(0, p)$ ne majore pas A ; en particulier A n'admet pas de plus grand élément. Quels que soient $p, q \in \mathbb{N}$, $(0, p) <_L (1, q)$, donc l'ensemble des majorants de A est $M = \{(1, m) \mid m \in \mathbb{N}\}$. Or $(1, 0)$ est clairement le plus petit élément de M , donc $(1, 0)$ est la borne supérieure de A . $\square$

1.20. Partie majorée n'ayant pas de borne supérieure.

Nous munissons $E = \{(0, 1)\} \times \mathbb{Z}$ de l'ordre lexicographique $\leq_L$ — voir l'exemple précédent — ce qui fait de E un ensemble ordonné. Posons $A = \{(0, n) \mid n \in \mathbb{Z}\}$. Si $p \in \mathbb{Z}$, $(1, p)$ majore A , et comme $(0, p) <_L (0, p + 1)$ et que $(0, p + 1) \in A$, $(0, p)$ ne majore pas A . Les majorants de A sont donc les $(1, n)$ pour $n \in \mathbb{Z}$. Si $q \in \mathbb{Z}$, $(1, q - 1)$ majore A et $(1, q - 1) <_L (1, q)$, donc l'ensemble des majorants de A n'admet pas de plus petit élément : A n'a pas de borne supérieure. $\square$

DÉFINITION 1.9. — Un élément maximal de E est un élément β de E tel que, pour tout élément x de E , $x \geq \beta$ entraîne $x = \beta$, et un élément minimal de E un élément α de E tel que, pour tout élément x de E , $x \leq \alpha$ entraîne $x = \alpha$.

Les éléments maximaux de E sont donc ceux qui ne sont strictement inférieurs à aucun élément de E , et les éléments minimaux de E ceux qui ne sont strictement supérieurs à aucun élément de E . S'il existe, le plus grand élément de E est un élément maximal de E et c'est le seul, et de même, s'il existe, le plus petit élément de E est un élément minimal de E et c'est le seul. Si $(E, \leq)$ est un ensemble totalement ordonné et a un élément de E , a est un élément maximal de E si, et seulement si, a est l'élément maximum de E , et a est un élément minimal de E si, et seulement si, a est l'élément minimum de E .

1.21. Ensemble ordonné ayant plusieurs éléments minimaux¹².

Nous munissons l'ensemble $E = \{n \in \mathbb{N} \mid n \geq 2\}$ de la relation de divisibilité. Les éléments minimaux de E sont les entiers $p \geq 2$ qui n'admettent aucun diviseur autre que 1 ou p , ce sont donc les nombres premiers, et il y en a une infinité. $\square$

12. Voici un autre exemple, de nature géographique : si $\leq$ est la relation d'ordre «en amont de» sur l'ensemble V des villes, villages et hameaux situés au bord d'un fleuve ou de l'un de ses affluents, les éléments minimaux de V sont ceux qui sont situés au plus haut sur le fleuve ou l'un des affluents ; il y en a en général beaucoup...

DÉFINITION 1.10. — Une application f d'un ensemble ordonné $(E, \leq)$ dans un ensemble ordonné $(F, \leq)$ est croissante (resp. décroissante) si, quels que soient les éléments x et y de E , $x \leq y$ entraîne $f(x) \leq f(y)$ (resp. $f(x) \geq f(y)$).

Si f est une application croissante et bijective d'un ensemble ordonné $(E, \leq)$ dans un ensemble ordonné $(F, \leq)$, et si l'ordre est total sur E , son application réciproque f^{-1} est une application croissante de $(F, \leq)$ dans $(E, \leq)$.

1.22. Application f bijective et croissante pour laquelle f^{-1} n'est pas croissante.

Appelons f l'identité de l'ensemble $\mathbb{N}^*$, muni au départ de la relation de divisibilité et à l'arrivée de l'ordre habituel. Si $a, b \in \mathbb{N}^*$ et si a divise b , alors $a \leq b$, ce qui montre que f est croissante ; par contre, $2 \leq 3$ bien que 2 ne divise pas 3, donc f^{-1} n'est pas croissante. $\square$

Lorsque l'ordre sur l'ensemble de départ est un ordre total, une application strictement croissante est injective. Ceci devient faux pour un ordre partiel.

1.23. Application strictement croissante qui n'est pas injective.

Nous choisissons un ensemble fini E de cardinal supérieur ou égal à 2 et nous munissons l'ensemble $\mathcal{P}(E)$ des parties de E de la relation d'ordre définie par l'inclusion. Si A et B sont des parties de E , si $A \subset B$ et si $A \neq B$, alors $\text{card } A < \text{card } B$, donc l'application $f : X \mapsto f(X) = \text{card } X$ de $\mathcal{P}(E)$ dans $\mathbb{N}$ est strictement croissante. Cependant deux parties différentes de E peuvent avoir le même cardinal — par exemple $\{a\}$ et $\{b\}$ où a et b sont des éléments distincts choisis dans E — donc f n'est pas injective. $\square$

■ Treillis

Dans un ensemble ordonné dont l'ordre n'est pas total, certaines parties à deux éléments n'ont pas de plus grand élément ni de plus petit élément ; il est intéressant de rechercher alors si une telle partie à deux éléments admet ou non une borne supérieure et une borne inférieure.

DÉFINITION 1.11. — Un treillis¹³ est un ensemble ordonné E dans lequel toute partie à deux éléments admet une borne supérieure et une borne inférieure.

On note, pour tout couple (a, b) d'éléments d'un treillis E , $a \vee b$ (resp. $a \wedge b$) la borne supérieure (resp. la borne inférieure) de la paire $\{a, b\}$.

Un ensemble totalement ordonné $(E, \leq)$ est évidemment un treillis, puisque si a et b sont des éléments de E , la paire $\{a, b\}$ admet un plus grand élément et un plus petit élément, qui sont respectivement sa borne supérieure et sa borne inférieure.

En notant $|$ la relation de divisibilité sur $\mathbb{N}^*$, $(\mathbb{N}^*, |)$ est un treillis ; en effet, si m et n appartiennent à $\mathbb{N}^*$, le ppcm et le pgcd du couple (m, n) sont respectivement la borne supérieure et la borne inférieure de $\{m, n\}$ dans l'ensemble ordonné $(\mathbb{N}^*, |)$.

13. Cette notion a été introduite dans les années 1930 par les mathématiciens américains Garrett Birkhoff (1911-1996) et Philip Hall (1904-1982).

1.24. Ensemble ordonné qui n'est pas un treillis.

Nous posons $E = \{0, 1\} \times \mathbb{Z}$ et nous notons $\leq$ la relation binaire définie sur E par : $(k, n) \leq (\ell, p)$ si $k = \ell$ et $n \leq p$. La réflexivité est immédiate. Si (k, n) , (ℓ, p) et (m, q) sont des éléments de E et si $(k, n) \leq (\ell, p)$ et $(\ell, p) \leq (m, q)$, alors $k = \ell = m$ et $n \leq p \leq q$, donc $(k, n) \leq (m, q)$. Enfin, si (k, n) et (ℓ, p) sont des éléments de E et si $(k, n) \leq (\ell, p)$ et $(\ell, p) \leq (k, n)$, alors $k = \ell$ et $n \leq p \leq n$, donc $(k, n) = (\ell, p)$. Nous avons établi que $(E, \leq)$ est un ensemble ordonné. Comme $(0, 0)$ et $(1, 0)$ n'ont pas de majorant commun, la paire $\{(0, 0), (1, 0)\}$ n'a pas de borne supérieure dans $(E, \leq)$. L'ensemble ordonné $(E, \leq)$ n'est donc pas un treillis. $\square$

Nous donnons un exemple dans lequel toute paire possède des majorants alors que pour certaines d'entre-elles, l'ensemble des majorants n'a pas de plus petit élément.

1.25. Autre ensemble ordonné qui n'est pas un treillis.

Nous notons $\mathcal{C}$ l'ensemble des parties connexes de $\mathbb{R}^2$ muni de sa topologie canonique¹⁴, et nous munissons $\mathcal{C}$ de la relation d'ordre définie par l'inclusion. Nous posons :

$$A = \mathbb{R} \times \mathbb{R}_+^* = \{z = (x, y) \mid x, y \in \mathbb{R} \text{ et } y > 0\}$$

et :

$$B = \mathbb{R} \times \mathbb{R}_-^* = \{z = (x, y) \mid x, y \in \mathbb{R} \text{ et } y < 0\}.$$

Les ensembles A et B sont des parties connexes de $\mathbb{R}^2$ comme produits de connexes. De plus A et B sont des ouverts disjoints et différents de l'ensemble vide, donc $A \cup B$ n'est pas un connexe de $\mathbb{R}^2$.

Soit a un nombre réel. Nous posons $C_a = A \cup B \cup \{(a, 0)\}$ et nous notons D_a la droite d'équation $x = a$. On a $C_a = A \cup B \cup D_a = (A \cup D_a) \cup (B \cup D_a)$; or $A \cup D_a$ et $B \cup D_a$ sont connexes comme réunions de deux connexes d'intersection non vide et de même $C_a = (A \cup D_a) \cup (B \cup D_a)$ est connexe. De plus $A \cup B$ n'appartient pas à $\mathcal{C}$, donc C_a est un élément minimal de l'ensemble des majorants de la paire $\{A, B\}$ dans $(\mathcal{C}, \subset)$.

Par conséquent l'ensemble des majorants de $\{A, B\}$ dans $(\mathcal{C}, \subset)$ admet une infinité d'éléments minimaux, donc il n'admet pas de plus petit élément. Il en résulte que la paire $\{A, B\}$ n'admet pas de borne supérieure dans $(\mathcal{C}, \subset)$.

L'ensemble ordonné $(\mathcal{C}, \subset)$ n'est donc pas un treillis. $\square$

THÉORÈME 1.3. — Dans un treillis $(E, \leq)$, on a, pour tout triplet (a, b, c) d'éléments de E , $a \vee (b \wedge c) \leq (a \vee b) \wedge (a \vee c)$ et $a \wedge (b \vee c) \geq (a \wedge b) \vee (a \wedge c)$.

DÉFINITION 1.12. — Un treillis $(E, \leq)$ est distributif si, pour tout triplet (a, b, c) d'éléments de E , les deux inégalités du théorème 1.3 sont des égalités.

1.26. Treillis qui n'est pas distributif¹⁵.

Nous considérons un ensemble $E = \{0, 1, a, b, c\}$ composé de cinq éléments et nous le munissons de la relation d'ordre $\leq$ possédant les propriétés suivantes : 0 est le plus petit et 1 le plus grand élément de $(E, \leq)$, et a , b et c ne sont pas comparables entre eux — on obtient clairement ainsi une relation d'ordre sur E .

14. Voir dans le chapitre 15, pages 295 à 301, la définition et quelques propriétés des connexes.

15. Voir aussi l'exemple 4.8, page 64.

L'ensemble ordonné $(E, \leq)$ est un treillis ; en effet, une partie de E à deux éléments appartenant à $\{a, b, c\}$ admet 0 pour borne inférieure et 1 pour borne supérieure, la partie $\{0, 1\}$ admet 0 pour borne inférieure et 1 pour borne supérieure, la partie $\{0, x\}$ où x appartient à $E \setminus \{0\}$ admet 0 pour borne inférieure et x pour borne supérieure, et la partie $\{y, 1\}$ où y appartient à $E \setminus \{1\}$ admet y pour borne inférieure et 1 pour borne supérieure.

De plus $a \vee b = 1$, donc $(a \vee b) \wedge c = c$, alors que $(a \wedge c) \vee (b \wedge c) = 0 \vee 0 = 0$. Il en résulte que $(E, \leq)$ n'est pas un treillis distributif. $\square$

Si $(E, \leq)$ est un treillis, alors, pour tout triplet (a, b, c) d'éléments de E , l'assertion $a \leq c$ entraîne l'inégalité $a \vee (b \wedge c) \leq (a \vee b) \wedge c$.

DÉFINITION 1.13. — Un treillis modulaire est un treillis tel que, pour tout triplet (a, b, c) d'éléments de E , $a \leq c$ implique $a \vee (b \wedge c) = (a \vee b) \wedge c$.

1.27. Treillis qui n'est pas modulaire.

Nous considérons le groupe $\mathfrak{S}_4$ des permutations de l'ensemble $\llbracket 1, 4 \rrbracket = \{1, 2, 3, 4\}$ et nous désignons, quels que soient les éléments distincts i et j de $\{1, 2, 3, 4\}$, par (i/j) la transposition de $\{1, 2, 3, 4\}$ qui échange i et j .

Nous notons $\mathfrak{A}_4$ le groupe alterné — c'est-à-dire le sous-groupe de $\mathfrak{S}_4$ formé des permutations paires —, $\mathcal{R}$ le sous-groupe d'ordre 3 engendré par le 3-cycle $(1/2/3)$, $\mathcal{V}$ le groupe à quatre éléments $\{\text{Id}, (1/2) \circ (3/4), (1/3) \circ (2/4), (1/4) \circ (2/3)\}$ et $\mathcal{B} = \{\text{Id}, (1/2) \circ (3/4)\}$. On a le schéma d'inclusions suivant :

$$\begin{array}{ccccccc} \{\text{Id}\} & \longrightarrow & \mathcal{R} & \longrightarrow & \mathfrak{A}_4 & \longrightarrow & \mathfrak{S}_4 \\ & & \searrow & & \nearrow & & \\ & & \mathcal{B} & \longrightarrow & \mathcal{V} & & \end{array}$$

Il résulte des propriétés des sous-groupes que l'ensemble des sous-groupes de $\mathfrak{S}_4$ est un treillis pour l'inclusion, la borne inférieure d'une paire de sous-groupes étant leur intersection et la borne supérieure le sous-groupe engendré par leur réunion. On a $\mathcal{B} \subset \mathcal{V}$ et $\mathcal{R} \cap \mathcal{V} = \{\text{Id}\}$, donc $\mathcal{B} \vee (\mathcal{R} \wedge \mathcal{V}) = \{\text{Id}\}$. Soit $\mathcal{K}$ un sous-groupe de $\mathfrak{A}_4$ contenant $\mathcal{R}$ et $\mathcal{B}$; son ordre est un multiple des ordres de $\mathcal{R}$ et $\mathcal{B}$, donc un multiple de 6, ce qui montre que l'ordre de $\mathcal{K}$ est 6 ou 12. Cependant l'exemple 2.16 (page 23) nous montre que $\mathfrak{A}_4$ n'a aucun sous-groupe d'ordre 6, donc $\mathcal{K} = \mathfrak{A}_4$. Par conséquent $\mathcal{B} \vee \mathcal{R} = \mathfrak{A}_4$ donc $(\mathcal{B} \vee \mathcal{R}) \wedge \mathcal{V} = \mathcal{V}$. En conclusion le treillis des sous-groupes de $\mathfrak{A}_4$ n'est pas modulaire. $\square$

Pour obtenir un treillis non modulaire, il était en fait suffisant de définir un ensemble à six éléments ordonné par le diagramme ci-dessus, mais nous avons préféré cet exemple par les sous-groupes, car il montre que le théorème affirmant que l'ensemble des sous-groupes distingués d'un groupe est un treillis modulaire ne se généralise pas aux sous-groupes quelconques.

DÉFINITION 1.14. — Un morphisme de treillis est une application d'un treillis dans un autre qui conserve la borne supérieure et la borne inférieure de toutes les parties à deux éléments.

Un morphisme de treillis est une application croissante ; cependant la réciproque est en général fausse.

1.28. Application croissante d'un treillis dans un autre qui n'est pas un morphisme de treillis.

Nous notons f l'identité de $\mathbb{N}^*$, muni au départ de la relation de divisibilité et à l'arrivée de l'ordre habituel ; nous savons qu'il s'agit, au départ comme à l'arrivée, d'un treillis. Si $m, n \in \mathbb{N}^*$ et si m divise n , alors $m \leq n$, donc f est croissante. Or la borne supérieure de $\{2, 3\}$ dans le treillis de départ est 6, alors que c'est 3 dans le treillis d'arrivée, donc f n'est pas un morphisme de treillis. $\square$

■ Bon ordre et axiome du choix

DÉFINITION 1.15. — Un bon ordre sur un ensemble E est une relation d'ordre sur E pour laquelle toute partie non vide de E admet un plus petit élément.

L'exemple fondamental de bon ordre est celui de l'ordre canonique de $\mathbb{N}$.

Un bon ordre sur un ensemble E est un ordre total sur E — en effet, si $\leq$ est un bon ordre sur E , alors, quels que soient les éléments x et y de E , la paire $\{x, y\}$ admet un plus petit élément z et, comme $z = x$ ou $z = y$, on a $x \leq y$ ou $y \leq x$. La réciproque est fautive comme le prouve l'exemple suivant.

1.29. Ordre total qui n'est pas un bon ordre.

Nous munissons $\mathbb{Z}$ de son ordre canonique ; c'est un ordre total sur $\mathbb{Z}$. Pour tout entier relatif n , $n - 1 < n$, ce qui montre que la partie $\mathbb{Z}$ de $\mathbb{Z}$ n'est pas minorée dans $\mathbb{Z}$, donc cette partie non vide de $\mathbb{Z}$ n'admet pas de plus petit élément. Par conséquent $\leq$ n'est pas un bon ordre sur $\mathbb{Z}$. $\square$

En fait, une partie non vide de $\mathbb{Z}$ admet un plus petit élément si, et seulement si, elle est minorée dans $\mathbb{Z}$. Comme pour $\mathbb{Z}$, l'ordre canonique de $\mathbb{R}$ est un ordre total sur $\mathbb{R}$ mais n'est pas un bon ordre. En effet, une partie non vide de $\mathbb{R}$ peut ne pas être minorée dans $\mathbb{R}$ — c'est le cas de $\mathbb{R}$ lui-même, de $\mathbb{Q}$, de $\mathbb{Z} \dots$ — et si elle l'est, elle admet une borne inférieure qui n'a aucune raison d'être son plus petit élément ; ainsi 0 est la borne inférieure de $A =]0, 1]$ ou de $A = \{1/n \mid n \in \mathbb{N}^*\}$ mais n'est pas le plus petit élément de A .

Parmi les axiomes de la théorie des ensembles ZF^{16} figure l'axiome du choix.

DÉFINITION 1.16. — Une fonction de choix sur un ensemble E est une application γ de $\mathcal{P}(E) \setminus \{\emptyset\}$ dans E telle que, pour toute partie non vide X de E , l'objet $\gamma(X)$ appartient à l'ensemble X .

Une fonction de choix γ sur E associe donc à toute partie non vide X de E un objet $\gamma(X)$ choisi dans X .

Remarquons que si $\leq$ est un bon ordre sur un ensemble E , on construit facilement une fonction de choix γ sur E en notant, pour toute partie non vide X de E , $\gamma(X)$ le plus petit élément de X .

Axiome du choix.

Sur tout ensemble il existe une fonction de choix.

16. Voir à la page 3 ce qui suit l'exemple 1.5.

L'axiome du choix est nécessaire à la démonstration de nombreux résultats, dans tous les domaines des mathématiques, comme par exemple le théorème qui affirme que si f est une surjection d'un ensemble E sur un ensemble F , il existe une injection g de F dans E telle que $g \circ f = \text{Id}_E$ — pour le démontrer, on introduit une fonction de choix γ sur E et l'application $g : y \mapsto g(y) = \gamma(A_y)$ de F dans E où, pour tout élément y de F , A_y est l'ensemble des antécédents de y par f , qui n'est pas vide puisque f est une surjection de E sur F .

C'est l'axiome du choix qui permet à Zermelo de prouver, en 1904, son théorème du bon ordre¹⁷ : *Sur tout ensemble il existe au moins un bon ordre.*

Plusieurs assertions sont équivalentes à l'axiome du choix et en simplifient l'application, comme le théorème de Zermelo¹⁸ et surtout le lemme de Zorn¹⁹.

■ Arithmétique

Certaines conjectures en arithmétique ont été longues à justifier et il a fallu parfois attendre plusieurs siècles pour les démontrer. La plus célèbre est bien sûr celle de Fermat, démontrée seulement en 1993-1994 par le mathématicien anglais Andrew Wiles et connue désormais sous le nom de théorème de Fermat-Wiles : *Pour tout entier naturel $p \geq 3$, l'équation $x^p + y^p = z^p$ n'admet aucune solution (x, y, z) où x, y et z sont des entiers naturels non nuls.*

S'inspirant de ce (futur) résultat, Leonhard Euler conjecture plus généralement qu'une puissance n -ième ne peut s'écrire comme la somme de k puissances n -ième avec $k < n$. Ainsi d'après lui on ne peut écrire $q^4 = a^4 + b^4 + c^4$ ni $q^5 = a^5 + b^5 + c^5 + d^5$ où a, b, c, d et q sont des entiers naturels non nuls, ce qui est faux comme le prouvent les exemples suivants.

1.30. Puissance quatrième somme de trois puissances quatrième.

Il «suffit» de remarquer que :

$$2682440^4 + 15365639^4 + 18796760^4 = 20615673^4.$$

Si notre lecteur en doute, nous lui laissons le soin de la vérification²⁰... □

1.31. Puissance cinquième somme de quatre puissances cinquième.

Nous avons en effet l'égalité $27^5 + 84^5 + 110^5 + 133^5 = 144^5$. □

On peut démontrer que tout entier naturel s'écrit comme la somme d'au plus neuf cubes d'éléments de $\mathbb{N}^*$, mais que l'on ne peut réduire à huit cubes.

17. Pour ce faire, Zermelo part du «principe que, même pour une totalité infinie d'ensembles, il y a toujours des applications qui associent à chaque ensemble l'un de ses éléments.»

18. Sur tout ensemble il existe au moins un bon ordre.

19. Dans un ensemble ordonné inductif — ce qui signifie que toute partie non vide et totalement ordonnée est majorée —, tout élément est majoré par au moins un élément maximal. Ce théorème est énoncé, en 1934, par le mathématicien américain Max Zorn (1907-1993).

20. Ce résultat a été démontré en 1966 par le mathématicien américain d'origine grecque Noam Elkies, ce qui a mis fin à la conjecture d'Euler.

1.32. Entier qui ne s'écrit pas comme la somme d'au plus huit cubes.

Si l'on décompose 23 en somme de cubes, ceux-ci sont des cubes de 1 ou de 2 puisque $3^3 = 27$ est plus grand que 23. Or $3 \times 2^3 = 24 > 23$, donc au plus deux d'entre eux sont égaux à 8. Pour compléter alors à 23, il faut utiliser 7 fois $1^3 = 1$. Ainsi 23 ne peut s'écrire comme la somme de 8 cubes ou moins. $\square$

On peut de même prouver qu'on ne peut pas écrire 239 comme la somme de moins de 9 cubes.

Les nombres de Fermat, ainsi appelés car ils ont été introduits par Fermat, sont les entiers naturels $F_n = 2^{2^n} + 1$ pour n parcourant $\mathbb{N}$. On a $F_0 = 3$, $F_1 = 5$, $F_2 = 17$ et $F_3 = 257$, donc F_n est un nombre premier pour $n = 0, 1, 2$ et 3.

1.33. Nombre de Fermat qui n'est pas un nombre premier.

Nous démontrons que le nombre de Fermat $F_5 = 2^{32} + 1$ est divisible par 641. De $640 = 5 \times 2^7$ on déduit que $5 \times 2^7 \equiv -1 \pmod{641}$; en élevant à la puissance 4, on obtient que $5^4 \times 2^{28} \equiv 1 \pmod{641}$. Or $5^4 = 625 \equiv -16 \pmod{641}$ donc, en multipliant ces deux congruences, il vient : $-2^{32} = (-2^4) \times 2^{28} \equiv 1 \pmod{641}$. Ainsi 641 divise F_5 et, comme $641 \neq F_5$, F_5 n'est pas un nombre premier. $\square$

Fermat affirme et croit démontrer en 1658 que F_n est un nombre premier pour tout $n \in \mathbb{N}$, ce qui est faux comme le montre l'exemple précédent 1.33. En fait F_5 est le produit des deux nombres premiers 641 et 6 700 417 et Landry a par ailleurs prouvé, en 1880, que $F_6 = 274 177 \times 67 280 421 310 721$. L'entier $F_4 = 65 537$ est, comme les F_n pour $n = 0, 1, 2$ et 3, un nombre premier mais on n'a pour le moment trouvé aucun entier $n \geq 5$ tel que F_n soit un nombre premier. La conjecture « Pour tout entier naturel n , F_n est un nombre premier » de Fermat semble donc se transformer en : « Pour tout entier $n \geq 5$, F_n n'est pas un nombre premier »...

THÉORÈME 1.4. — Petit théorème de Fermat.

Si p est un nombre premier, si a est un entier naturel et si p est premier avec a , alors p divise $a^{p-1} - 1$.

Cette propriété caractérise-t-elle les nombres premiers ? Il n'en est rien.

1.34. Entier $p \geq 2$ qui n'est pas un nombre premier²¹ et tel que, pour tout entier naturel a premier avec p , p divise $a^{p-1} - 1$.

Nous posons $p = 3 \times 11 \times 17 = 561$. Soit a un entier naturel non nul premier avec p . Par le petit théorème de Fermat, nous savons que $a^2 \equiv 1 \pmod{3}$. En élevant cette égalité à la puissance 280, on obtient que $a^{560} \equiv 1 \pmod{3}$. De même $a^{10} \equiv 1 \pmod{11}$, et en élevant à la puissance 56, on obtient que $a^{560} \equiv 1 \pmod{11}$. Enfin $a^{16} \equiv 1 \pmod{17}$ donc, en remarquant que $560 = 35 \times 16$ et en élevant à la puissance 35, il vient : $a^{560} \equiv 1 \pmod{17}$. Comme 3, 11 et 17 sont des nombres premiers et qu'ils divisent tous trois $a^{560} - 1$, leur produit 561 divise $a^{560} - 1$. $\square$

21. Les entiers $p \geq 2$ tels que, pour tout entier $a \geq 1$ premier avec p , p divise $a^{p-1} - 1$, sont les nombres de Carmichael, du nom du mathématicien américain Robert Carmichael (1879-1967).

Chapitre 2

Groupes

Evariste Galois introduit vers 1830 la notion de groupe. Il meurt peu de temps après, au cours d'un duel, à l'âge de vingt-et-un an. Ses notes sont confuses et ce n'est qu'une quinzaine d'années plus tard que l'on comprend leur importance.

Deux types de problèmes amènent à la formalisation de cette théorie. La première, issue des travaux de Galois, est l'étude des permutations des lettres, on dirait de nos jours l'étude du groupe des bijections d'un ensemble fini sur lui même. La seconde est celle des groupes de transformation en géométrie, en d'autres termes de groupes de bijections du plan ou de l'espace conservant certaines propriétés (isométries, similitudes, etc.).

Plusieurs mathématiciens, parmi lesquels Camille Jordan et Leopold Kronecker saisissent alors l'importance de fonder cette théorie sur des axiomes pour lui donner toute sa généralité. La formalisation définitive n'est faite qu'en 1893.

■ Lois de composition interne

DÉFINITION 2.1. — Une loi de composition interne sur un ensemble E est une application de $E \times E$ dans E .

Dans ce qui suit, une loi de composition interne sur E est symbolisée par $*$ et on note $a * b$ l'image d'un couple (a, b) d'éléments de E par cette loi de composition interne, que l'on appelle le composé de a et de b par $*$.

DÉFINITION 2.2. — La loi de composition interne $*$ est associative si, quels que soient les éléments a, b et c de E , $(a * b) * c = a * (b * c)$, et commutative si, pour tout couple (a, b) d'éléments de E , $a * b = b * a$.

Parmi les lois de composition interne fréquemment utilisées sur les ensembles de nombres, rares sont celles qui ne sont pas associatives.

2.1. Loi de composition interne sur $\mathbb{N}^*$ qui n'est ni associative ni commutative.

Nous définissons la loi de composition interne $*$ sur $\mathbb{N}$ par $a * b = a^b$. Comme $2 * 3 = 2^3 = 8 \neq 9 = 3^2 = 3 * 2$, la loi $*$ n'est pas commutative. Elle n'est pas non plus associative ; en effet, on a $(2 * 2) * 3 = (2^2)^3 = 4^3 = 64$ alors que $2 * (2 * 3) = 2^{2^3} = 2^8 = 256$. $\square$

2.2. Loi de composition interne sur $\mathbb{Z}$ qui n'est ni associative ni commutative.

Nous définissons la loi de composition interne $*$ sur $\mathbb{Z}$ par $a * b = a - b$. Si a est différent de b , alors $a - b \neq b - a$. De plus, si $c \neq 0$, $a - (b - c)$ est différent de $(a - b) - c$. La loi $*$ n'est donc ni commutative ni associative sur $\mathbb{Z}$. $\square$

Un élément e de E est un élément neutre à gauche (resp. à droite) pour la loi de composition interne $*$ sur E si $e * x = x$ (resp. $x * e = x$) pour tous les éléments x de E . Un élément neutre pour $*$ est un élément neutre simultanément à droite et à gauche. Si une loi de composition interne possède un élément neutre à gauche et un élément neutre à droite, ils sont égaux, ce qui assure l'existence et l'unicité de l'élément neutre. Cependant une loi peut avoir une infinité d'éléments neutres d'un côté ; compte tenu de ce qui précède, elle n'en possède aucun de l'autre côté.

2.3. Loi possédant une infinité d'éléments neutres à gauche.

Nous définissons la loi de composition interne $*$ sur $\mathbb{N}$ par $a * b = b$. Cette loi est clairement associative et tout entier naturel en est un élément neutre à gauche. $\square$

Si une loi de composition interne $*$ sur E possède un élément neutre e , ou un unique élément neutre e d'un côté, un élément x de E possède un symétrique à gauche x' pour cette loi si $x' * x = e$ et de même x possède un symétrique x' à droite pour $*$ si $x * x' = e$. Si $*$ est associative et si x possède un symétrique à gauche et à droite, ils sont égaux, ce qui garantit l'existence et l'unicité *du* symétrique de x .

2.4. Loi possédant un élément neutre à droite et dont les éléments ont une infinité de symétrique de chaque côté.

Nous désignons par E la fonction partie entière et nous définissons sur $\mathbb{N}$ la loi de composition interne $*$ par :

$$a * b = E\left(\frac{a}{b+1}\right)$$

Pour tout entier naturel a , $a * 0$ est la partie entière de $a/1 = a$, donc $a * 0 = a$: 0 est un élément neutre à droite de $*$. D'autre part, si $a, b \in \mathbb{N}$ et si $b \geq a$, on a $b+1 > a \geq 0$, ce qui montre que :

$$0 \leq \frac{a}{b+1} < 1, \text{ donc } a * b = E\left(\frac{a}{b+1}\right) = 0.$$

Par suite, si a est un entier naturel, tout entier naturel $b \geq a$ est un symétrique de a à droite. De même, si $a, b \in \mathbb{N}$ et si $b \leq a$, alors $b * a = 0$. Il en résulte que si a est un entier naturel, tout entier naturel $b \leq a$ est un symétrique de a à gauche. En particulier a est symétrique de lui-même à droite et à gauche. $\square$

2.5. Loi de composition interne possédant un élément neutre et dont les éléments ont une infinité de symétriques.

Nous définissons la loi de composition interne $*$ sur $\mathbb{N}$ par : $a * 0 = 0 * a = a$, et $a * b = 0$ si $a \neq 0$ et $b \neq 0$. Clairement 0 est un élément neutre de $*$ et, si $a \in \mathbb{N}^*$, tout entier naturel $b \neq 0$ est un symétrique de a à droite et à gauche. $\square$

2.6. Élément possédant une infinité de symétriques à gauche pour une loi de composition interne associative.

Nous considérons l'ensemble $\mathcal{A}_{\mathbb{N}}$ des applications de $\mathbb{N}$ dans $\mathbb{N}$, que nous munissons de la loi de composition des applications. Ceci définit sur $\mathcal{A}_{\mathbb{N}}$ une loi de composition interne associative dont l'identité Id de $\mathbb{N}$ est l'élément neutre. Nous choisissons une injection non surjective f de $\mathbb{N}$ dans $\mathbb{N}$ — par exemple $f : n \mapsto f(n) = 2n$. Nous associons à tout entier naturel n l'application :

$$g_n : \mathbb{N} \longrightarrow \mathbb{N} \\ y \longmapsto g_n(y) = \begin{cases} x, & \text{unique antécédent de } y \text{ par } f \\ n & \text{si } y \notin f(\mathbb{N}). \end{cases}$$

Soit $n \in \mathbb{N}$. Pour tout $x \in \mathbb{N}$, l'unique antécédent de $f(x)$ par f est bien sûr x , donc $g_n(f(x)) = x$. Par suite $g_n \circ f = \text{Id}$, donc g_n est un symétrique de f à gauche dans $(\mathcal{A}_{\mathbb{N}}, \circ)$. En faisant décrire $\mathbb{N}$ à n , on obtient une infinité de symétriques de f à gauche. De plus f n'est pas une surjection de $\mathbb{N}$ sur $\mathbb{N}$, donc f n'admet aucun symétrique à droite. En choisissant une surjection f non injective de $\mathbb{N}$ sur $\mathbb{N}$, on obtient de même une infinité de symétriques à droite pour f dans $(\mathcal{A}_{\mathbb{N}}, \circ)$. $\square$

■ Axiomes de la structure de groupe

DÉFINITION 2.3. — Un ensemble G muni d'une loi de composition interne notée $*$ est un groupe¹ s'il vérifie les trois axiomes suivants.

- (I) La loi $*$ est associative.
- (II) La loi $*$ possède un élément neutre e .
- (III) Tout élément de G possède un symétrique.

On peut en fait affaiblir les deux derniers axiomes en les remplaçant par :

- (II') La loi $*$ possède un élément neutre à gauche.
- (III') Tout élément de G possède un symétrique à gauche.

On peut bien sûr remplacer dans les deux axiomes (II') et (III') « à gauche » par « à droite ». Cependant on ne peut remplacer « à gauche » par « à droite » seulement dans (III').

2.7. Loi de composition interne associative et régulière à droite et à gauche, mais ne définissant pas un groupe.

Il suffit de considérer l'ensemble $\mathbb{N}$ des entiers naturels muni de l'addition. $\square$

1. La notion de groupe a été axiomatisée en 1893 par le mathématicien Walter von Dyck (1856-1934). On lui devait déjà une axiomatisation des groupes finis parue en 1882 dans le numéro 20 de la revue *Mathematische Annalen*. Curieusement, dans ce même numéro, son compatriote Heinrich Weber (1842-1913) avait donné une axiomatisation équivalente à celle de Dyck. Celle-ci imposait l'associativité de la loi et la régularité à droite et à gauche, ce qui signifie que, quels que soient les éléments a, b et c de G , $a * c = b * c$ entraîne $a = b$ et $c * a = c * b$ entraîne $a = b$. Dans le cas d'un ensemble infini, ces propriétés sont insuffisantes pour obtenir un groupe.

2.8. Loi de composition interne vérifiant les axiomes (I), (II) et (III'), mais ne définissant pas une structure de groupe.

Nous définissons sur $\mathbb{R}^*$ la loi de composition interne $*$ par $a * b = |a|b$. Cette loi est associative ; en effet, quels que soient les éléments a, b et c de $\mathbb{R}^*$:

$$a * (b * c) = a * (|b|c) = |a|(|b|c) = (|a||b|)c = |ab|c = (a * b) * c.$$

Pour tout élément b de $\mathbb{R}^*$, $1 * b = |1|b = b$, donc 1 est un élément neutre à gauche de $*$. Si $b \in \mathbb{R}^*$ et si l'on pose $b' = 1/|b|$, alors $b * b' = |b|(1/|b|) = 1$, donc b' est un symétrique à droite de b . Cependant $*$ ne possède aucun élément neutre ; en effet, si un tel élément e existait, e serait différent de zéro et on aurait, pour tout $b \in \mathbb{R}^*$, $|e|b = e * b = b * e = |b|e$, donc $|e| = e$ pour $b = 1$ et $|e| = -e$ pour $b = -1$, d'où une contradiction. On en déduit que $(\mathbb{R}^*, *)$ n'est pas un groupe. $\square$

Nous revenons aux trois axiomes (I), (II) et (III) de la définition 2.3 (page 19) d'un groupe. Nous allons étudier des ensembles munis d'une loi de composition interne vérifiant deux seulement de ces trois axiomes ; remarquons tout de suite que si (II) n'est pas vérifié, l'axiome (III) n'a pas de sens.

2.9. Loi de composition interne vérifiant seulement les deux premiers axiomes.

Il suffit de nouveau de considérer $\mathbb{N}$ muni de l'addition. $\square$

2.10. Loi de composition interne vérifiant seulement les axiomes (II) et (III).

Nous considérons un ensemble $E = \{e, a, b\}$ à trois éléments et la loi de composition interne $*$ définie sur E par : $e * x = x * e = x$ pour tout élément x de E et $x * y = e$ quels que soient les éléments x et y de $E \setminus \{e\}$. Clairement e est l'élément neutre de $(E, *)$ et tout élément de E est son propre symétrique. Cependant la loi n'est pas associative car $a * (a * b) = a * e = a$ et $(a * a) * b = e * b = b$. $\square$

Les permutations d'un ensemble E sont les bijections de E sur E et, muni de la loi « rond » de composition des applications, l'ensemble $\mathfrak{S}(E)$ des permutations de E est un groupe dont l'élément neutre est l'identité Id_E de E . Si E est un ensemble et a et b des éléments distincts de E , la transposition de E qui intervertit (ou échange) a et b est la permutation f de E définie par : $f(a) = b$, $f(b) = a$ et $f(x) = x$ pour tout élément x de $E \setminus \{a, b\}$.

Pour tout entier naturel n , on note $\mathfrak{S}_n$ le groupe des permutations de l'ensemble $[1, n] = \{1, \dots, n\}$. Soit $n \in \mathbb{N}^*$. La permutation σ de $\{1, \dots, n\}$ se note :

$$\begin{bmatrix} 1 & 2 & \dots & n \\ \sigma(1) & \sigma(2) & \dots & \sigma(n) \end{bmatrix}.$$

Si p est un entier ≥ 2 et $h_1, \dots, h_p$ des éléments deux à deux distincts de $\{1, \dots, n\}$, $\gamma = (h_1 / \dots / h_p)$ est la permutation de $\{1, \dots, n\}$ définie par :

$$\gamma(h_1) = h_2, \dots, \gamma(h_{p-1}) = h_p, \gamma(h_p) = h_1 \text{ et } \gamma(i) = i \text{ pour tout } i \notin \{h_1, \dots, h_p\}.$$

Les $\gamma = (h_1 / \dots / h_p)$ sont les cycles de longueur p , également appelés les p -cycles. Si i et j sont des éléments distincts de $\{1, \dots, n\}$, la transposition qui échange i et j est donc le cycle (i/j) , de longueur 2.

DÉFINITION 2.4. — L'ordre d'un groupe fini est son cardinal.

Par exemple, si n est un entier naturel, le groupe $\mathfrak{S}_n$ est fini et d'ordre $n!$.

Nous rappelons enfin une définition essentielle, celle des sous-groupes.

DÉFINITION 2.5. — Une partie non vide H d'un groupe G est un sous-groupe de G si H est stable par la loi de composition interne de G et si H , muni de la loi induite, est un groupe.

Si H est un sous-groupe d'un groupe G , l'élément neutre du groupe H est alors celui de G . Un sous-groupe de G est une partie H de G contenant l'élément neutre e de G , stable par la loi de composition interne de G et par le passage au symétrique, ou encore une partie H de G contenant e et telle que, quels que soient les éléments x et y de H , le composé de x par le symétrique de y appartient encore à H .

■ Centre d'un groupe

Un groupe est commutatif si sa loi de composition interne est commutative ; on dit aussi que c'est un groupe abélien. Certains groupes finis ne sont pas commutatifs, mais un groupe fini dont l'ordre est un nombre premier est commutatif, ainsi qu'un groupe fini d'ordre ≤ 5 . Un groupe fini non commutatif est donc d'ordre au moins 6.

2.11. Groupe d'ordre 6 qui n'est pas commutatif.

Nous considérons le groupe $\mathfrak{S}_3$ des permutations de l'ensemble $\{1, 2, 3\}$ et nous introduisons les transpositions $\sigma_1 = (2/3)$ et $\sigma_2 = (1/3)$ de $\{1, 2, 3\}$. On voit que $\sigma_1 \circ \sigma_2(1) = \sigma_1(3) = 2$ et que $\sigma_2 \circ \sigma_1(1) = \sigma_2(1) = 3$; par conséquent $\sigma_1 \circ \sigma_2 \neq \sigma_2 \circ \sigma_1$. Le groupe $\mathfrak{S}_3$ n'est donc pas un groupe commutatif. De plus $\mathfrak{S}_3$ est d'ordre $3! = 6$, donc c'est le plus petit groupe non commutatif. $\square$

Dans un groupe, on s'intéresse aux éléments qui commutent avec tous les autres.

DÉFINITION 2.6. — Le centre d'un groupe $(G, *)$ est l'ensemble $Z(G)$ des éléments x de G tels que $x * y = y * x$ pour tout élément y de G .

Le centre $Z(G)$ d'un groupe G est un sous-groupe distingué² de G et il est égal à G si, et seulement si, le groupe G est commutatif.

2.12. Groupe dont le centre est réduit à l'élément neutre.

Nous reprenons le groupe $\mathfrak{S}_3$ des permutations de $\{1, 2, 3\}$ et nous notons, pour $i = 1, 2$ et 3 , σ_i la transposition qui laisse fixe i , ρ_1 la permutation circulaire directe et ρ_2 la permutation circulaire inverse ; ce sont les 3-cycles $\rho_1 = (1/2/3)$ et $\rho_2 = (3/2/1)$. Nous avons vu dans l'exemple 2.11 que σ_1 et σ_2 n'appartiennent pas au centre de $\mathfrak{S}_3$. De plus $\rho_1 \circ \sigma_1(1) = \rho_1(1) = 2$ alors que $\sigma_1 \circ \rho_1(1) = \sigma_1(2) = 3$, donc $\rho_1 \circ \sigma_1 \neq \sigma_1 \circ \rho_1$, ce qui montre que ρ_1 n'appartient pas au centre de $\mathfrak{S}_3$; de même σ_3 et ρ_2 n'appartiennent pas au centre de $\mathfrak{S}_3$. En conclusion, le centre du groupe $\mathfrak{S}_3$ est réduit à son élément neutre $\text{Id} = \text{Id}_{\{1,2,3\}}$. $\square$

2. Voir la définition 2.8, page 24.

2.13. Groupe fini dont le centre est un sous-groupe propre.

Nous considérons le groupe quaternionique $G = \{1, -1, i, -i, j, -j, k, -k\}$ à huit éléments avec les opérations $ij = -ji = k$, $jk = -kj = i$ et $ki = -ik = j$, les éléments avec un signe $-$ se multipliant avec les règles habituelles sur les signes³. Le groupe G n'est pas commutatif et, clairement, -1 et 1 appartiennent au centre de G et aucun autre élément de G n'y appartient ; par suite $Z(G) = \{1, -1\}$. $\square$

2.14. Groupe infini dont le centre est un sous-groupe propre.

Nous considérons l'anneau $\mathcal{M}_2(\mathbb{R})$ des matrices carrées d'ordre 2 à coefficients réels et le groupe multiplicatif $\text{GL}_2(\mathbb{R})$ des éléments inversibles de l'anneau $\mathcal{M}_2(\mathbb{R})$. Nous démontrons que le centre du groupe $\text{GL}_2(\mathbb{R})$ est l'ensemble des matrices représentatives des homothéties, c'est-à-dire des matrices :

$$N = \begin{pmatrix} a & 0 \\ 0 & a \end{pmatrix}$$

où a est un nombre réel non nul. Soit N une telle matrice. Un calcul facile montre que, quels que soient les réels u, v, w et x :

$$\text{si } M = \begin{pmatrix} u & v \\ w & x \end{pmatrix}, \text{ alors } MN = \begin{pmatrix} au & av \\ aw & ax \end{pmatrix} = NM,$$

d'où l'on déduit que N appartient au centre du groupe $\text{GL}_2(\mathbb{R})$.

Soit maintenant un élément $N = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ du centre du groupe $\text{GL}_2(\mathbb{R})$. On pose :

$$A = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \text{ et } B = \begin{pmatrix} 1 & 0 \\ 0 & 2 \end{pmatrix}.$$

Alors :

$$A \times N = \begin{pmatrix} c & d \\ a & b \end{pmatrix}, N \times A = \begin{pmatrix} b & a \\ d & c \end{pmatrix} \text{ et } A \times N = N \times A,$$

donc $b = c$ et $a = d$, et de même :

$$B \times N = \begin{pmatrix} a & b \\ 2b & 2a \end{pmatrix}, N \times B = \begin{pmatrix} a & 2b \\ b & 2a \end{pmatrix} \text{ et } B \times N = N \times B,$$

donc $b = 2b$; par suite $c = b = 0$ et $d = a$, donc $N = \begin{pmatrix} a & 0 \\ 0 & a \end{pmatrix}$. $\square$

■ Sous-groupes

La définition des sous-groupes a été rappelée à la page 21 (définition 2.5).

Si G est un groupe multiplicatif (resp. additif) et a un élément de G , l'ensemble :

$$H = \{a^k \mid k \in \mathbb{Z}\} \text{ (resp. } H = \{k \cdot a \mid k \in \mathbb{Z}\})$$

est le plus petit modulo l'inclusion des sous-groupes de G contenant a ; H est appelé le sous-groupe de G engendré par a et on dit que a engendre le groupe G lorsque $G = H$.

Plus généralement, si A est une partie d'un groupe G , le sous-groupe de G engendré par A est le plus petit modulo l'inclusion des sous-groupes de G contenant A .

3. G est en fait un sous-groupe du groupe multiplicatif des éléments non nuls de l'algèbre $\mathbb{H}$ des quaternions d'Hamilton.

DÉFINITION 2.7. — Un groupe cyclique est un groupe fini engendré par l'un de ses éléments.

Tous les sous-groupes d'un groupe cyclique sont des groupes cycliques.

2.15. Groupe dont tous les sous-groupes stricts sont des groupes cycliques mais qui n'est pas lui-même un groupe cyclique.

Nous considérons le groupe additif $G = \mathbb{Z}/2\mathbb{Z} \times \mathbb{Z}/2\mathbb{Z}$ muni de l'addition produit⁴ : $(x, y) + (x', y') = (x + x', y + y')$, dont l'élément neutre est $(0, 0)$. En examinant le sous-groupe de G engendré par chacun de ses quatre éléments, on voit que G n'est pas un groupe cyclique et on prouve aisément que les sous-groupes stricts de G sont les groupes cycliques $\{(0, 0)\}$ engendré par $(0, 0)$, $\mathbb{Z}/2\mathbb{Z} \times \{0\}$ engendré par $(1, 0)$, $\{0\} \times \mathbb{Z}/2\mathbb{Z}$ engendré par $(0, 1)$ et $\{(0, 0), (1, 1)\}$ engendré par $(1, 1)$. $\square$

THÉORÈME 2.1. — Théorème de Lagrange⁵.

Dans un groupe fini, l'ordre de tout sous-groupe divise l'ordre du groupe.

On peut se poser le problème inverse : étant donné un diviseur p de l'ordre n d'un groupe fini G , existe-t-il un sous-groupe H de G d'ordre p ? Si G est cyclique et si p est un nombre premier, la réponse est affirmative. Cependant ce résultat ne s'étend pas aux groupes finis quelconques.

Si n est un entier ≥ 2 , l'ensemble $\mathfrak{A}_n$ des permutations paires — c'est-à-dire de signature $+1$ — de $\{1, \dots, n\}$ est un sous-groupe du groupe $\mathfrak{S}_n$ des permutations de $\{1, \dots, n\}$. Le groupe $\mathfrak{A}_n$ est appelé le groupe alterné et son ordre est $n!/2$.

2.16. Groupe fini d'ordre 12 n'ayant aucun sous-groupe d'ordre 6.

Nous considérons le groupe alterné $\mathfrak{A}_4$, d'ordre 12, et nous introduisons l'ensemble $\mathcal{V}$ composé de l'identité $\text{Id} = \text{Id}_{\{1,2,3,4\}}$ et des permutations :

$$r_1 = \begin{bmatrix} 1 & 2 & 3 & 4 \\ 2 & 1 & 4 & 3 \end{bmatrix}, \quad r_2 = \begin{bmatrix} 1 & 2 & 3 & 4 \\ 3 & 4 & 1 & 2 \end{bmatrix} \quad \text{et} \quad r_3 = \begin{bmatrix} 1 & 2 & 3 & 4 \\ 4 & 3 & 2 & 1 \end{bmatrix}.$$

Comme r_1 , r_2 et r_3 sont des produits de deux transpositions, leur signature est égale à $+1$, donc $\mathcal{V}$ est inclus dans $\mathfrak{A}_4$. Chaque élément de $\mathcal{V}$ est son propre symétrique et on a, si les indices i , j et k sont deux à deux distincts, $r_i \circ r_j = r_k$. Il en résulte que $\mathcal{V}$ est un sous-groupe du groupe $\mathfrak{A}_4$.

Nous supposons l'existence d'un sous-groupe $\mathcal{H}$ d'ordre 6 de $\mathfrak{A}_4$. L'intersection $\mathcal{H}' = \mathcal{H} \cap \mathcal{V}$ est un sous-groupe de $\mathcal{H}$ et de $\mathcal{V}$, donc son ordre, qui divise 4 et 6, est égal à 1 ou à 2. Comme $\mathfrak{A}_4$ possède, outre les éléments de $\mathcal{V}$, les huit 3-cycles, $\mathcal{H}$ possède quatre ou cinq cycles de longueur 3. Si $\mathcal{H}$ contient $t_1 = (1/2/3)$ et $t_2 = (1/2/4)$, alors $r_2 = t_1 \circ t_2$ et $r_3 = t_2 \circ t_1$ appartiennent à $\mathcal{H}$ donc $\mathcal{H}'$ admet au moins trois éléments, en contradiction avec ce qui précède. On peut en fait choisir t_1 de manière quelconque. Ceci étant fait, les entiers 1, 2 et 3 jouant le même rôle, on peut supposer que 3 est invariant par t_2 et, par définition de t_2 , il ne reste plus qu'à opérer une permutation circulaire sur 1, 2 et 4 ; en remplaçant éventuellement t_2 par t_2^{-1} , les deux 3-cycles envoient 1 sur la même image.

En conclusion, $\mathfrak{A}_4$ ne possède aucun sous-groupe d'ordre 6. $\square$

4. Il s'agit du groupe de Klein, du nom du mathématicien allemand Felix Klein (1849-1925).

5. Pour une démonstration, voir [SCHW], chapitre 1, § 2, théorème 1.

DÉFINITION 2.8. — Un sous-groupe H d'un groupe multiplicatif G est distingué si, pour tout élément x de G et tout élément y de H , xyx^{-1} appartient à H .

L'intérêt de cette notion est la possibilité, lorsque H est un sous-groupe distingué de G , de définir, sur l'ensemble quotient G/H de G par la relation d'équivalence $\mathcal{R}$, définie sur G par : $x\mathcal{R}y$ si $xy^{-1} \in H$, une structure de groupe (unique) telle que la surjection canonique de G sur G/H est un morphisme de groupe⁶ ; le groupe G/H ainsi obtenu est appelé le *groupe quotient de G par H* .

2.17. Sous-groupe qui n'est pas un sous-groupe distingué.

Nous reprenons le groupe $\mathfrak{S}_3$ des permutations de $\{1, 2, 3\}$ et nous notons, pour $i = 1, 2$ et 3 , σ_i la transposition qui laisse fixe i ; remarquons que $\sigma_i \circ \sigma_i$ est l'identité. Il en résulte que $H = \{\text{Id}, \sigma_1\}$ est un sous-groupe du groupe $\mathfrak{S}_3$. On a $\sigma_2 \circ \sigma_1 \circ \sigma_2(2) = \sigma_2 \circ \sigma_1(2) = \sigma_2(3) = 1$; aucun élément de H n'envoyant 2 sur 1, le composé $\sigma_2 \circ \sigma_1 \circ \sigma_2^{-1} = \sigma_2 \circ \sigma_1 \circ \sigma_2$ n'appartient pas à H , donc H n'est pas un sous-groupe distingué de $\mathfrak{S}_3$. $\square$

DÉFINITION 2.9. — Si H et K sont des sous-groupes d'un groupe multiplicatif G , leur produit est l'ensemble $HK = \{xy \mid x \in H \text{ et } y \in K\}$.

Si H et K sont des sous-groupes d'un groupe multiplicatif G et si H est distingué, HK est un sous-groupe de G et c'est le sous-groupe de G engendré par $H \cup K$.

2.18. Produit de deux sous-groupes qui n'est pas un sous-groupe.

Nous reprenons le groupe $\mathfrak{S}_3$ des permutations de $\{1, 2, 3\}$, nous conservons les notations de l'exemple précédent 2.17 et nous posons $H = \{\text{Id}, \sigma_1\}$ et $K = \{\text{Id}, \sigma_2\}$. On a alors $HK = \{\text{Id}, \sigma_1, \sigma_2, \rho\}$ où $\rho = \sigma_1 \circ \sigma_2 = (1/2/3)$. Comme la permutation :

$$\rho \circ \rho = \begin{bmatrix} 1 & 2 & 3 \\ 3 & 1 & 2 \end{bmatrix} = (1/3/2)$$

n'appartient pas à HK , HK n'est pas un sous-groupe de $\mathfrak{S}_3$. $\square$

Les sous-groupes distingués de G sont les sous-groupes de G stables par tous les automorphismes intérieurs de G , c'est-à-dire les :

$$\left| \begin{array}{l} f_a : G \longrightarrow G \\ x \longmapsto f_a(x) = axa^{-1} \end{array} \right.$$

pour a décrivant G . On appelle de manière analogue sous-groupe caractéristique de G un sous-groupe stable par tous les automorphismes. Tout sous-groupe caractéristique de G est donc un sous-groupe distingué de G .

2.19. Sous-groupe distingué qui n'est pas caractéristique.

Nous considérons le groupe de Klein $G = \mathbb{Z}/2\mathbb{Z} \times \mathbb{Z}/2\mathbb{Z}$, groupe additif défini dans l'exemple 2.15 (page 23). Le groupe G étant commutatif, tous ses sous-groupes sont distingués, en particulier le sous-groupe $H = \mathbb{Z}/2\mathbb{Z} \times \{0\}$ de G .

6. Un rappel des définitions des morphismes, isomorphismes, endomorphismes et automorphismes de groupe se trouve à la page 28.

L'application $f : z = (x, y) \mapsto f(z) = (y, x)$ de G dans G est une bijection de G sur G et on a, quels que soient les éléments x, y, x' et y' de $\mathbb{Z}/2\mathbb{Z}$:

$$\begin{aligned} f((x, y) + (x', y')) &= f(x + x', y + y') = (y + y', x + x') \\ &= (y, x) + (y', x') = f(x, y) + f(x', y'), \end{aligned}$$

donc f est un automorphisme du groupe G . Il est clair cependant que $f(H)$ n'est pas inclus dans H , donc H n'est pas un sous-groupe caractéristique de G . $\square$

2.20. Groupe non commutatif dans lequel tout sous-groupe propre est commutatif et distingué.

Nous considérons le groupe quaternionique $G = \{1, -1, i, -i, j, -j, k, -k\}$ — voir l'exemple 2.13 (page 22). L'ordre d'un sous-groupe strict de G est 1, 2 ou 4 (théorème de Lagrange). Tout groupe d'ordre inférieur ou égal à 5 étant commutatif, tous les sous-groupes stricts de G sont commutatifs. Le seul sous-groupe d'ordre 2 est $\{1, -1\}$; il est distingué car 1 et -1 commutent avec tous les éléments de G . Les autres éléments de G sont d'ordre 4 — voir ce qui suit. Les trois sous-groupes d'ordre 4 sont donc ceux qui sont engendrés par i, j ou k . Comme ces éléments jouent le même rôle, nous démontrons que le sous-groupe H engendré par i est distingué. Or les deux produits : $(-j)i(-j)^{-1} = jij^{-1} = (ji)(-j) = (-k)(-j) = i$ et $(-k)i(-k)^{-1} = kik^{-1} = j(-k) = -i$ appartiennent à H , donc H est distingué. $\square$

■ Ordre d'un élément dans un groupe

DÉFINITION 2.10. — Si G est un groupe multiplicatif d'élément neutre e , un élément x de G est d'ordre fini s'il existe au moins un entier $k \geq 1$ tel que $x^k = e$, et l'ordre de x est alors le plus petit des entiers $k \geq 1$ tel que $x^k = e$, et d'ordre infini dans le cas contraire.

Dans le cas où G est un groupe additif d'élément neutre 0 , il suffit dans cette définition de remplacer la proposition $x^k = e$ par l'assertion $k \cdot x = 0$.

Un élément x d'un groupe G est d'ordre fini q si, et seulement si, il engendre un sous-groupe fini d'ordre q de G .

Il résulte donc du théorème de Lagrange (théorème 2.1, page 23) que, dans un groupe fini, l'ordre de tout élément est un diviseur de l'ordre du groupe.

L'ordre du produit ab de deux éléments d'un groupe multiplicatif G est égal à celui de ba . Ceci devient faux pour trois éléments ou plus.

2.21. Trois éléments a, b et c du groupe $\mathfrak{S}_3$ des permutations de $\{1, 2, 3\}$ tels que $a \circ b \circ c$ et $c \circ b \circ a$ n'ont pas le même ordre.

Nous considérons dans le groupe $\mathfrak{S}_3$ des permutations de $\{1, 2, 3\}$ les transpositions $a = (1/2)$ et $b = (2/3)$ et le 3-cycle $c = (1/2/3)$. Alors $a \circ b \circ c = (1/3/2)$, d'ordre 3, et $b \circ c \circ a = \text{Id}$, d'ordre 1, donc $a \circ b \circ c$ et $c \circ b \circ a$ n'ont pas le même ordre. $\square$

Soit G un groupe fini. On lui associe les trois entiers naturels suivants : son ordre $O(G)$ (son cardinal), le plus grand ordre $a(G)$ des éléments de G et le plus petit entier strictement positif $b(G) = k$ tel que $x^k = e$ pour tout élément x de G . Alors $a(G)$ divise $b(G)$ et $b(G)$ divise $O(G)$, on a $a(G) = O(G)$ si, et seulement si, G est cyclique, et $a(G) = b(G)$ dans le cas où le groupe G est commutatif.

2.22. Groupe commutatif dans lequel $b(G) \neq O(G)$.

Dans le groupe de Klein $G = \mathbb{Z}/2\mathbb{Z} \times \mathbb{Z}/2\mathbb{Z}$ (voir les exemples 2.15 page 23 et 2.19 page 24), tout élément est son propre symétrique, donc $b(G) = 2$; par ailleurs G possède quatre éléments donc $b(G) \neq O(G)$. $\square$

2.23. Groupe dans lequel $a(G) \neq b(G)$ et $b(G) \neq O(G)$.

Nous considérons le groupe $\mathfrak{S}_5$ des permutations de $\{1, 2, 3, 4, 5\}$; il est d'ordre $5! = 120$. La décomposition d'une permutation en produit de cycles disjoints prouve que tout élément de $\mathfrak{S}_5$ qui n'est pas un cycle est le composé de deux transpositions disjointes, et il est alors d'ordre 2, ou bien le composé d'une transposition et d'un 3-cycle disjoints, et dans ce cas il est d'ordre 6. Or les k -cycles sont d'ordre $k \in \llbracket 2, 5 \rrbracket$, donc $a(\mathfrak{S}_5) = 6$. De plus $b(\mathfrak{S}_5)$ est le ppcm des ordres des éléments de $\mathfrak{S}_5$, donc $b(\mathfrak{S}_5) = \text{ppcm}(2, 3, 4, 5, 6) = 30$. Il en résulte que $a(\mathfrak{S}_5) < b(\mathfrak{S}_5) < O(\mathfrak{S}_5)$. $\square$

Dans un groupe commutatif G vérifiant $b(G) = O(G)$, l'égalité $a(G) = b(G)$ entraîne $a(G) = O(G)$, ce qui montre que G est cyclique.

2.24. Groupe qui n'est pas cyclique mais dans lequel $b(G) = O(G)$.

Nous choisissons dans le groupe $\mathfrak{S}_3$ des permutations de $\{1, 2, 3\}$ une transposition τ et une permutation circulaire ρ . Alors $\tau^2 = \rho^3 = \text{Id}$, donc $\tau^6 = \rho^6 = \text{Id}$, ce qui montre que $b(\mathfrak{S}_3) = 6 = O(\mathfrak{S}_3)$. $\square$

Dans un groupe commutatif G , l'ordre d'un élément divise $a(G)$. On voit dans l'exemple précédent 2.24 que l'ordre de τ ne divise pas $a(\mathfrak{S}_3) = 3$.

Si deux éléments commutent dans un groupe multiplicatif, l'ordre de leur produit divise le ppcm des ordres de ces deux éléments. Ceci est donc toujours réalisé dans un groupe commutatif.

2.25. Deux éléments d'ordre fini dont l'ordre du produit ne divise pas le ppcm de leurs ordres.

Nous considérons, dans le groupe $\mathfrak{S}_5$ des permutations de $\{1, 2, 3, 4, 5\}$, les 3-cycles $\sigma_1 = (1/2/3)$ et $\sigma_2 = (3/4/5)$, qui sont tous deux d'ordre 3. Leur produit $\sigma_1 \circ \sigma_2 = (1/2/3/4/5)$ est d'ordre 5, entier strictement plus grand que le ppcm des ordres de σ_1 et σ_2 et qui n'en est même pas un multiple. $\square$

2.26. Deux éléments d'ordre fini dont le produit est d'ordre infini.

L'ensemble $G = \text{GL}_2(\mathbb{Z})$ des matrices carrées d'ordre 2 à coefficients dans $\mathbb{Z}$ dont le déterminant vaut 1 ou -1 est un sous-groupe du groupe multiplicatif $\text{GL}_2(\mathbb{Q})$ des éléments inversibles de l'anneau $\mathcal{M}_2(\mathbb{Q})$ des matrices carrées d'ordre 2 à coefficients

rationnels. En effet, la matrice unité I_2 appartient clairement à G , le produit MN de deux matrices M et N à coefficients entiers est à coefficients entiers et, si $\det(M) = \pm 1$ et $\det(N) = \pm 1$, alors $\det(MN) = \pm 1$, et enfin, si une matrice :

$$M = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

appartient à G , son inverse dans $\text{GL}_2(\mathbb{Q})$, qui vaut $M^{-1} = \varepsilon \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}$ où $\varepsilon = \pm 1$ comme inverse de $\det(M) = \pm 1$, appartient à G .

Nous posons $A = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$ et $B = \begin{pmatrix} 0 & -1 \\ 1 & -1 \end{pmatrix}$.

Le calcul donne $A^4 = I_2$ et $B^3 = I_2$, donc A et B sont d'ordre fini dans $G = \text{GL}_2(\mathbb{Z})$.

De plus $AB = \begin{pmatrix} -1 & 1 \\ 0 & -1 \end{pmatrix}$ et un calcul facile donne, pour tout entier naturel n :

$$(AB)^n = (-1)^n \begin{pmatrix} 1 & -n \\ 0 & 1 \end{pmatrix},$$

matrice dont le coefficient d'indice $(1, 2)$ ne s'annule pas pour $n \in \mathbb{N}^*$. La matrice AB est donc d'ordre infini dans le groupe $G = \text{GL}_2(\mathbb{Z})$. $\square$

2.27. Groupe infini dans lequel tout élément est d'ordre fini.

Nous notons G le groupe additif quotient $\mathbb{Q}/\mathbb{Z}$, c'est-à-dire le quotient du groupe additif de $\mathbb{Q}$ par la relation d'équivalence $\mathcal{R}$, définie sur $\mathbb{Q}$ par : $x \mathcal{R} y$ si $x - y$ appartient à $\mathbb{Z}$, et muni de l'addition quotient. Soit r un élément non nul de $\mathbb{Q}$ et (p, q) le représentant irréductible de r de dénominateur positif. Les entiers p et q sont premiers entre eux et le produit $qr = p$ est un entier relatif. Notons $\bar{x}$ la classe d'un rationnel x modulo $\mathcal{R}$. On a $q\bar{r} = \overline{qr} = \bar{p} = \bar{0}$, élément neutre de G , ce qui, puisque p et q sont premiers entre eux, montre que l'ordre de $\bar{r}$ dans le groupe G est égal à q . Par conséquent tout élément de G est d'ordre fini. De plus, si m et n sont des éléments distincts de $\mathbb{N}^*$, les nombres rationnels $1/m$ et $1/n$ ne sont pas équivalents modulo $\mathcal{R}$; en effet :

$$-1 < d = \frac{1}{m} - \frac{1}{n} < 1 \text{ et } d \neq 0,$$

donc d n'est pas un nombre entier, ce qui prouve que leurs classes d'équivalence modulo $\mathcal{R}$ sont différentes. On en déduit que G est infini. $\square$

Dans l'exemple précédent 2.27, les entiers $a(G)$ et $b(G)$ ne sont pas définis, car pour tout entier $q \geq 1$, la classe d'équivalence de $1/q$ modulo $\mathcal{R}$ est d'ordre q dans le groupe G , donc il existe dans G des éléments d'ordre aussi grand que l'on veut.

2.28. Groupe infini G dans lequel $b(G)$ est défini.

Nous choisissons un ensemble infini E et nous munissons l'ensemble $\mathcal{P}(E)$ des parties de E de la différence symétrique Δ , loi de composition interne définie sur $\mathcal{P}(E)$ par :

$$X \Delta Y = (X \setminus Y) \cup (Y \setminus X) = (X \cup Y) \setminus (X \cap Y).$$

La loi Δ est associative, l'ensemble vide est son élément neutre et, pour toute partie A de E , $A \Delta A = \emptyset$ donc A est son propre symétrique. Il en résulte que $(\mathcal{P}(E), \Delta)$ est un groupe dans lequel tout élément est d'ordre 2, donc $a(\mathcal{P}(E)) = b(\mathcal{P}(E)) = 2$. L'ensemble E étant infini, il en est de même de $\mathcal{P}(E)$. $\square$

2.29. Groupe infini dans lequel tout sous-groupe strict est cyclique (donc fini).

Pour tout entier naturel k , l'ensemble $\mathbb{U}_k = \{z \in \mathbb{C} \mid z^k = 1\}$ des racines k -ièmes de l'unité est un sous-groupe du groupe multiplicatif $\mathbb{U}$ des nombres complexes de module 1. Soit p un nombre premier. Nous posons :

$$G = \bigcup_{n \in \mathbb{N}} \mathbb{U}_{p^n}.$$

Nous montrons que G est un sous-groupe de $\mathbb{U}$. Clairement, $G \subset \mathbb{U}$ et 1 appartient à G . Si a et b sont des éléments de G , il existe des entiers naturels m et n tels que $a \in \mathbb{U}_{p^m}$ et $b \in \mathbb{U}_{p^n}$, et, en supposant par exemple que $m \leq n$, p^m divise p^n donc $\mathbb{U}_{p^m} \subset \mathbb{U}_{p^n}$, ce qui montre que a et b appartiennent à $\mathbb{U}_{p^n}$, sous-groupe de $\mathbb{U}$, d'où l'on déduit que ab^{-1} appartient à $\mathbb{U}_{p^n}$ donc à G .

Soit H un sous-groupe strict de G . Nous choisissons un élément y de $G \setminus H$ et nous notons n le plus petit entier naturel tel que $y \in \mathbb{U}_{p^n}$. Soit x un élément de H et q le plus petit entier naturel tel que $x \in \mathbb{U}_{p^q}$. Alors il existe un entier relatif k tel que $x = \exp((2ik\pi)/p^q)$. Comme $x \notin \mathbb{U}_{p^{q-1}}$, les entiers k et p sont premiers entre eux, donc aussi k et p^q . Par suite x engendre le groupe $\mathbb{U}_{p^q}$, donc $\mathbb{U}_{p^q} \subset H$. Si $n \leq q$, alors $y \in \mathbb{U}_{p^n} \subset \mathbb{U}_{p^q}$ donc $y \in H$, en contradiction avec l'hypothèse. On en déduit que $H \subset \mathbb{U}_{p^{n-1}}$; en particulier l'ensemble H est fini. Nous notons m le plus petit entier naturel tel que $H \subset \mathbb{U}_{p^m}$ et nous choisissons un élément z de H appartenant à $\mathbb{U}_{p^m} \setminus \mathbb{U}_{p^{m-1}}$. Par un raisonnement analogue au précédent, on prouve que z engendre $\mathbb{U}_{p^m}$; il en résulte que $H \subset \mathbb{U}_{p^m} \subset H$. Finalement $H = \mathbb{U}_{p^m}$ et ce dernier est un groupe cyclique. $\square$

■ Morphismes et isomorphismes de groupes

Soit (G, γ) et (G', δ) des groupes. Un morphisme de groupe de G dans G' est une application f de G dans G' telle que, quels que soient $x, y \in G$:

$$f(x \gamma y) = f(x) \delta f(y).$$

Le noyau d'un morphisme de groupe f de G dans G' est l'ensemble $\text{Ker } f$ des éléments de G dont l'image par f est l'élément neutre de G' ; c'est un sous-groupe de G et f est injectif si, et seulement si, son noyau est réduit à l'élément neutre du groupe G . Un isomorphisme de groupe de G sur G' est un morphisme de groupe f de G dans G' pour lequel il existe un morphisme de groupe g de G' dans G tel que $g \circ f = \text{Id}_G$ et $f \circ g = \text{Id}_{G'}$. Les isomorphismes de groupe de G sur G' sont les morphismes de groupe bijectifs de G sur G' et les groupes G et G' sont isomorphes s'il existe un isomorphisme de G sur G' .

Si G est un groupe, les endomorphismes de G sont les morphismes de groupe de G dans G et les automorphismes de G sont les isomorphismes de G sur G .

2.30. Groupe isomorphe à l'un de ses sous-groupes propres.

L'ensemble $\mathbb{P}$ des entiers relatifs pairs est un sous-groupe du groupe additif $(\mathbb{Z}, +)$ et le groupe $(\mathbb{P}, +)$ est isomorphe à $(\mathbb{Z}, +)$; en effet l'application $\varphi : n \mapsto \varphi(n) = 2n$ de $\mathbb{Z}$ dans $\mathbb{P}$ est un isomorphisme de groupe de $(\mathbb{Z}, +)$ sur $(\mathbb{P}, +)$. $\square$

2.31. Groupe isomorphe à son quotient par un sous-groupe non réduit à l'élément neutre.

Nous considérons le groupe multiplicatif $\mathbb{U}$ des nombres complexes de module 1. Clairement, $H = \{-1, 1\}$ est un sous-groupe de $\mathbb{U}$.

Nous introduisons l'application $f : z \mapsto f(z) = z^2$ de $\mathbb{U}$ dans $\mathbb{U}$.

Quels que soient les éléments z et z' de $\mathbb{U}$, $f(zz') = zz'zz' = z^2z'^2 = f(z)f(z')$ et, pour tout élément z de $\mathbb{U}$, $z^2 = 1$ si, et seulement si, $z = 1$ ou $z = -1$, donc f est un endomorphisme du groupe $\mathbb{U}$, de noyau H . Tout nombre complexe $z = e^{i\theta}$ de module 1 admet deux racines carrées dans $\mathbb{C}$, à savoir $z_1 = e^{i\theta/2}$ et $z_2 = -z_1$, qui sont aussi des éléments de $\mathbb{U}$, donc f est une surjection de $\mathbb{U}$ sur $\mathbb{U}$. On note $\hat{z}$ la classe d'équivalence modulo H d'un élément quelconque z de $\mathbb{U}$. La décomposition canonique de f nous assure l'existence d'un isomorphisme de groupe g du groupe quotient $\mathbb{U}/H$ sur le groupe $\mathbb{U}$ — et que $g(\hat{z}) = z^2$ pour tout élément z de $\mathbb{U}$. $\square$

2.32. Autre groupe isomorphe à son quotient par un sous-groupe non réduit à l'élément neutre.

Nous notons G le groupe additif quotient $\mathbb{Q}/\mathbb{Z}$ qui a été défini dans l'exemple 2.27 (page 27). Soit n un entier supérieur ou égal à 2 et a la classe d'équivalence de $1/n$. Clairement, a est d'ordre n dans le groupe G . Nous notons H_n le sous-groupe de G engendré par a et nous considérons l'application :

$$\left| \begin{array}{l} \varphi_n : G \longrightarrow G \\ x \longmapsto \varphi_n(x) = n \cdot x. \end{array} \right.$$

Comme dans tout groupe abélien additif, $n \cdot (x + y) = n \cdot x + n \cdot y$ quels que soient les éléments x et y de G , donc φ_n est un morphisme de groupe de G dans G . Soit x un élément de G , classe d'équivalence d'un rationnel r de représentant irréductible (p, q) tel que $0 \leq p < q$. Alors x appartient au noyau de φ_n si, et seulement si, $(np)/q$ appartient à $\mathbb{Z}$, ce qui équivaut au fait que q divise np , donc par le théorème de Gauss que q divise n , soit encore à l'existence d'un entier naturel k tel que $n = kq$, c'est-à-dire $r = (kp)/n$. Ainsi $\text{Ker } \varphi_n = H_n$. On en déduit que φ_n passe au quotient⁷ et définit un morphisme de groupe injectif ψ_n de G/H_n dans G . Comme par ailleurs φ_n est une surjection de G sur G — en effet, la classe d'équivalence de $p/(qn)$ est un antécédent par φ_n de celle de p/q —, ψ_n est une surjection de G/H_n sur G , donc G est isomorphe à G/H_n . $\square$

Il n'existe aucun morphisme non nul du groupe additif $\mathbb{Z}/3\mathbb{Z}$ vers le groupe additif $\mathbb{Z}/2\mathbb{Z}$, puisque le premier n'a que des éléments dont l'ordre divise 3 et le second des éléments dont l'ordre est un diviseur de 2, alors que l'ordre de l'image par un morphisme de groupe d'un élément d'ordre k divise k .

2.33. Deux groupes additifs infinis sans morphisme de groupe non nul de l'un vers l'autre.

Nous reprenons le groupe additif quotient $G = \mathbb{Q}/\mathbb{Z}$ des exemples 2.27 et 2.32. Nous montrons qu'il n'existe aucun morphisme de groupe différent de l'application nulle de G dans $\mathbb{Z}$ ni de $\mathbb{Z}$ dans G . Soit φ un morphisme de groupe de $(G, +)$ dans $(\mathbb{Z}, +)$.

7. Voir [SCHW], chapitre 1, §2, proposition 4.

Si x appartient à G , x est d'ordre fini donc aussi son image $\varphi(x)$, ce qui, 0 étant le seul élément d'ordre fini du groupe $(\mathbb{Z}, +)$, montre que $\varphi(x) = 0$. On prouve de même qu'il n'existe aucun morphisme de groupe non nul de $(\mathbb{Z}, +)$ dans $(G, +)$. $\square$

THÉORÈME 2.2. — L'ensemble des automorphismes d'un groupe G est un sous-groupe du groupe $\mathfrak{S}(G)$ des permutations de G .

On note $\text{Aut}(G)$ le groupe des automorphismes du groupe G ainsi obtenu.

2.34. Groupe isomorphe au groupe de ses automorphismes.

Nous associons à tout élément s du groupe $\mathfrak{S}_3$ des permutations de $\llbracket 1, 3 \rrbracket = \{1, 2, 3\}$ l'automorphisme intérieur de $\mathfrak{S}_3$:

$$\left| \begin{array}{l} f_s : \mathfrak{S}_3 \longrightarrow \mathfrak{S}_3 \\ r \longmapsto f_s(r) = s \circ r \circ s^{-1} \end{array} \right.$$

(voir page 24 ce qui précède l'exemple 2.19). Quels que soient les éléments s et s' de $\mathfrak{S}_3$, on a, pour tout élément r de $\mathfrak{S}_3$:

$$f_{s \circ s'}(r) = (s \circ s') \circ r \circ (s \circ s')^{-1} = s \circ (s' \circ r \circ s'^{-1}) \circ s^{-1} = f_s \circ f_{s'}(r),$$

donc l'application :

$$\left| \begin{array}{l} \Phi : \mathfrak{S}_3 \longrightarrow \text{Aut}(\mathfrak{S}_3) \\ s \longmapsto \Phi(s) = f_s \end{array} \right.$$

est un morphisme de groupe du groupe $\mathfrak{S}_3$ dans le groupe $G = \text{Aut}(\mathfrak{S}_3)$ des automorphismes de $\mathfrak{S}_3$. Si s appartient à $\mathfrak{S}_3$ et si f_s est l'identité de $\mathfrak{S}_3$ alors, pour tout élément r de $\mathfrak{S}_3$, $r = f_s(r) = s \circ r \circ s^{-1}$ donc $s \circ r = r \circ s$, d'où l'on déduit que s appartient au centre de $\mathfrak{S}_3$ qui — voir l'exemple 2.12, page 21 — est réduit à l'identité. Il en résulte que Φ est un morphisme de groupe injectif de $\mathfrak{S}_3$ dans le groupe G . Par conséquent $\text{Im}(\Phi) = \{f_s \mid s \in \mathfrak{S}_3\}$ est un sous-groupe de $G = \text{Aut}(\mathfrak{S}_3)$ isomorphe à $\mathfrak{S}_3$.

Nous démontrons que ces deux groupes coïncident. Pour un automorphisme g de $\mathfrak{S}_3$ et un élément s de $\mathfrak{S}_3$, on a, pour tout entier relatif k , $g(s^k) = g(s)^k$, donc $g(s^k) = \text{Id}$ si, et seulement si, $g(s)^k = \text{Id}$. Ainsi l'ordre d'un élément est conservé par Φ , donc l'image d'une transposition est une transposition. De plus l'ensemble des trois transpositions de $\{1, 2, 3\}$ engendre le groupe $\mathfrak{S}_3$, donc g est entièrement déterminé par la donnée de leurs trois images, trois images sont possibles pour la première et ce choix fait, il ne reste que deux possibilités pour la deuxième et la dernière est alors déterminée. Ceci montre que le groupe G possède au plus six éléments. En conclusion⁸ le groupe G est isomorphe au groupe $\mathfrak{S}_3$. $\square$

■ Groupes simples et groupes résolubles

DÉFINITION 2.11. — Un groupe simple est un groupe qui ne possède aucun sous-groupe distingué autre que lui-même et le singleton élément neutre.

Tous les groupes finis dont l'ordre est un nombre premier sont des groupes simples et, pour les groupes commutatifs finis, ce sont les seuls.

8. Ce résultat s'applique en fait au groupe $\mathfrak{S}_n$ pour tout entier $n \geq 3$ différent de 6 (voir [SCHW], chapitre 2, exercice II.14).

DÉFINITION 2.12. — Un groupe G est résoluble s'il existe une suite $H_0, H_1, \dots, H_n$ de sous-groupes de G telle que, pour tout $i \in \llbracket 1, n \rrbracket$, H_{i-1} est un sous-groupe distingué de H_i et le groupe quotient H_i/H_{i-1} est commutatif.

2.35. Groupe fini simple dont l'ordre n'est pas un nombre premier.

Nous allons démontrer que le groupe alterné $\mathfrak{A}_5$ — voir page 23 ce qui précède l'exemple 2.16 — répond à la question. Pour simplifier, nous supprimons le « rond » dans l'écriture des composés, que nous appelons des produits.

Soit H un sous-groupe distingué de $\mathfrak{A}_5$ différent de $\{\text{Id}\}$. Montrons d'abord que si H contient un 3-cycle, il est égal à $\mathfrak{S}_5$ tout entier. Comme 1, 2, 3, 4 et 5 jouent des rôles analogues, on peut supposer que ce trois-cycle est $(1/2/3)$. Pour $k = 4$ ou 5 on a $(1/k/2) = (3/2/k)(1/2/3)(3/2/k)^{-1}$ et $(1/k/2)^2 = (1/2/k)$, donc le 3-cycle $(1/2/k)$ appartient à H , de même que $(1/i/j) = (1/2/j)^{-1}(1/2/i)(1/2/j)$ lorsque $i, j \in \{3, 4, 5\}$ et $i \neq j$. Nous rappelons que $\mathfrak{S}_5$ est engendré par l'ensemble des transpositions $(1/i)$ pour $i \in \llbracket 2, 5 \rrbracket$, donc tout élément de $\mathfrak{A}_5$ est le produit d'un nombre pair de transpositions du type $(1/i)$. Pour des entiers distincts $i, j \in \llbracket 2, 5 \rrbracket$, $(1/i/j) = (1/j)(1/i)$, ce qui montre que l'ensemble des 3-cycles de cette forme engendre $\mathfrak{A}_5$; or ils appartiennent tous à H , donc $H = \mathfrak{A}_5$.

Il nous reste à prouver que H contient effectivement un 3-cycle.

Les éléments de $\mathfrak{A}_5$ différents de l'identité sont, soit des 3-cycles, soit des 5-cycles, soit les produits de deux transpositions disjointes. Supposons d'abord que H contienne $s = (1/2)(3/4)$. Si l'on pose $r = (1/2/3)$, $v = (r^{-1}sr)s^{-1} = (1/4)(2/3)$ appartient aussi à H donc, en posant $u = (1/4/5)$, $w = (uv)u^{-1} = (2/3)(4/5)$ appartient à H , ainsi que le 3-cycle $vw = (1/4)(4/5) = (1/4/5)$. Enfin, si s est un 5-cycle appartenant à H , par exemple $s = (1/2/3/4/5)$, alors, en posant $r = (1/2/3)$, $s^{-1}rsr^{-1} = (1/3/5)$ est un 3-cycle appartenant à H .

Tout ceci montre que H est égal à $\mathfrak{A}_5$. On a prouvé que $\mathfrak{A}_5$ est un groupe simple.

Enfin, l'ordre de $\mathfrak{A}_5$ est $5!/2 = 60$, qui n'est pas un nombre premier. $\square$

2.36. Groupe qui n'est pas résoluble.

Le groupe $\mathfrak{A}_5$ n'admet que $\{\text{Id}\}$ et lui-même comme sous-groupes distingués (voir l'exemple précédent 2.35); comme il n'est pas commutatif, le groupe $\mathfrak{A}_5$ n'est pas résoluble. $\square$

Lorsque p est un nombre premier, tout groupe fini d'ordre p est cyclique donc isomorphe au groupe abélien additif $\mathbb{Z}/p\mathbb{Z}$, mais ceci ne caractérise pas les nombres premiers. Pour le justifier, nous introduisons les sous-groupes de Sylow⁹.

DÉFINITION 2.13. — Si G est un groupe fini d'ordre $n = p^k m$ où p est un nombre premier, k un entier naturel non nul et m un entier naturel non divisible par p , un p -sous-groupe de Sylow est un sous-groupe d'ordre p^k de G .

9. Du nom du mathématicien norvégien Ludwig Sylow (1832-1918).

THÉORÈME 2.3. — Théorème de Sylow¹⁰.

Si p est un nombre premier et si p divise l'ordre n d'un groupe fini G , G possède au moins un p -sous-groupe de Sylow et le nombre des p -sous-groupes de Sylow de G est un diviseur de n congru à 1 modulo p .

2.37. Entier naturel $n \geq 2$ qui n'est pas un nombre premier tel que tout groupe fini d'ordre n est cyclique.

Soit G un groupe d'ordre $15 = 3 \times 5$. Le nombre des 5-sous-groupes de Sylow de G divise 15 et il est congru à 1 modulo 5. De même le nombre des 3-sous-groupes de Sylow de G est un diviseur de 15 congru à 1 modulo 3. Le groupe G n'admet donc qu'un seul 5-sous-groupe de Sylow et un seul 3-sous-groupe de Sylow. L'ordre d'un élément de G différent de l'élément neutre divise 15, c'est donc 3, 5 ou 15. Un élément d'ordre 5 de G engendre un groupe d'ordre 5, c'est donc l'unique 5-sous-groupe de Sylow. Il y a donc dans G au plus quatre éléments d'ordre 5, qui sont les quatre éléments différents de l'élément neutre. De même il y a dans G au plus deux éléments d'ordre 3. Or G possède quinze éléments, dont quatre sont d'ordre 5 et deux d'ordre 3, et l'élément neutre est d'ordre 1. Par conséquent il reste huit éléments d'ordre 15 et un tel élément engendre G , qui est donc cyclique. $\square$

On s'intéresse à deux diviseurs premiers p et q de l'ordre n d'un groupe fini G .

DÉFINITION 2.14. — Si G est un groupe fini d'ordre $n = p^k q^\ell m$ où p et q sont des nombres premiers, k et ℓ des entiers naturels non nuls et m un entier naturel qui n'est divisible ni par p ni par q , un (p, q) -sous-groupe de Hall¹¹ est un sous-groupe de G d'ordre $p^k q^\ell$.

Contrairement aux sous-groupes de Sylow, un groupe fini dont l'ordre possède les bonnes propriétés n'admet pas nécessairement un sous-groupe de Hall.

2.38. Groupe fini ne possédant pas de $(3, 5)$ -sous-groupe de Hall.

Nous considérons le groupe $\mathfrak{S}_5$ des permutations de $[1, 5] = \{1, 2, 3, 4, 5\}$, d'ordre $120 = 3^1 \times 5^1 \times 8$. Les entiers 3 et 5 sont des nombres premiers et 8 est premier avec 3 et avec 5, donc 3 et 5 ne divisent pas 8. Nous démontrons que $\mathfrak{S}_5$ ne contient aucun $(3, 5)$ -sous-groupe de Hall, c'est-à-dire aucun sous-groupe d'ordre 15.

Supposons que $\mathfrak{S}_5$ contienne un sous-groupe d'ordre 15. Ce sous-groupe est cyclique comme nous l'avons montré dans l'exemple précédent 2.37, donc il existe dans le groupe $\mathfrak{S}_5$ un élément d'ordre 15. Un cycle de $\mathfrak{S}_5$ est d'ordre 2, 3, 4 ou 5. D'après la décomposition d'une permutation en produit de cycles disjoints, un élément qui n'est ni l'identité ni un cycle est, soit le composé de deux transpositions disjointes, soit le composé d'une transposition et d'un 3-cycle disjoints. Dans le premier cas il est d'ordre 2 et dans le second d'ordre 6.

Le groupe $\mathfrak{S}_5$ ne possède donc aucun sous-groupe d'ordre 15, donc il ne contient aucun $(3, 5)$ -sous-groupe de Hall. $\square$

10. Voir [SCHW], chapitre 2, §3.

11. Du nom du mathématicien américain Philipp Hall (1904-1982)

Chapitre 3

Anneaux et corps

La notion d'anneau a été introduite, dans la deuxième moitié du XIX^e siècle, par les mathématiciens allemands Richard Dedekind (1831-1916) et David Hilbert (1862-1943) pour généraliser les ensembles de nombres $\mathbb{Z}$, $\mathbb{Q}$ et $\mathbb{R}$, munis de l'addition et de la multiplication. Son utilité pour étudier les matrices ou certains espaces de fonctions a justifié l'introduction d'anneaux non commutatifs. C'est à cette même époque que Richard Dedekind introduit la notion d'idéal, sur une idée d'Ernst Kummer (1810-1893), pour affiner la notion de divisibilité. Dans les années 1920, la théorie des anneaux s'étoffe avec l'introduction de différents types d'anneaux, comme les anneaux principaux, factoriels ou noethériens.

DÉFINITION 3.1. — Un anneau est un triplet $(A, +, \times)$ où $(A, +)$ est un groupe abélien additif et $\times$ une loi de composition interne sur A , associative et distributive par rapport à $+$ sur A et possédant un élément neutre.

Soit A — en fait $(A, +, \times)$ — un anneau. L'addition et la multiplication de A sont respectivement les lois de composition interne $+$ et $\times$, le zéro de l'anneau A , noté 0_A ou 0 , est l'élément neutre de $(A, +)$ et l'élément unité de l'anneau A , noté 1_A ou 1 , celui de $(A, \times)$. Pour tout élément x de A , $0_A \times x = x \times 0_A = 0_A$. L'anneau A est nul s'il est égal à $\{0\}$, ce qui équivaut à $1_A = 0_A$.

DÉFINITION 3.2. — Un anneau est commutatif si sa multiplication est commutative.

DÉFINITION 3.3. — Si a et b appartiennent à un anneau commutatif A , a divise b dans A s'il existe un élément q de A tel que $b = aq$.

DÉFINITION 3.4. — Un élément a d'un anneau A est simplifiable à gauche (resp. à droite) si, quels que soient les éléments x et y de A , $ax = ay$ (resp. $xa = ya$) entraîne $x = y$, et régulier s'il est simplifiable à droite et à gauche.

DÉFINITION 3.5. — Un anneau intègre est un anneau commutatif non nul dans lequel tout élément est régulier.

DÉFINITION 3.6. — Un corps est un anneau non nul K dans lequel tout élément non nul est inversible, c'est-à-dire symétrisable dans $(K, \times)$.

L'anneau $\mathbb{Z}$ des nombres entiers relatifs est un anneau intègre et tout corps commutatif est un anneau intègre.

DÉFINITION 3.7. — Si A et B sont des anneaux, un morphisme d'anneau de A dans B est un morphisme de groupe de $(A, +)$ dans $(B, +)$ tel que $f(1_A) = 1_B$ et, quels que soient les éléments x et y de A , $f(xy) = f(x)f(y)$.

Soit A et B des anneaux. Un isomorphisme d'anneau de A sur B est un morphisme d'anneau f de A dans B pour lequel il existe un morphisme d'anneau g de B dans A tel que $g \circ f = \text{Id}_A$ et $f \circ g = \text{Id}_B$. Les isomorphismes d'anneau de A sur B sont les morphismes d'anneau bijectifs de A sur B .

■ Propriétés générales

3.1. Anneau qui n'est pas un anneau commutatif.

L'anneau infini $\mathcal{M}_2(\mathbb{R})$ des matrices carrées d'ordre 2 à coefficients réels n'est pas commutatif. On a en effet :

$$\begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \neq \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}. \quad \square$$

THÉORÈME 3.1. — Théorème de Wedderburn¹. Tout corps fini est commutatif.

Ce résultat est faux pour un anneau quelconque.

3.2. Anneau fini qui n'est pas un anneau commutatif.

Les mêmes calculs que dans l'exemple précédent 3.1 montrent que l'anneau $\mathcal{M}_2(\mathbb{F}_2)$ des matrices carrées d'ordre 2 à coefficients dans le corps $\mathbb{F}_2 = \mathbb{Z}/2\mathbb{Z}$ des entiers modulo 2 n'est pas commutatif, alors qu'il est fini (de cardinal $2^4 = 16$). $\square$

Certains auteurs n'imposent pas l'existence de l'élément unité dans la définition d'un anneau. Nous appelons *pseudo-anneau* une telle structure.

3.3. Pseudo-anneau qui n'est pas un anneau.

Nous notons $\mathbb{P}$ l'ensemble des entiers relatifs pairs. C'est un sous-groupe de $(\mathbb{Z}, +)$, stable par multiplication $\times$ de $\mathbb{Z}$. L'addition et la multiplication de $\mathbb{Z}$ induisent donc sur $\mathbb{P}$ des lois de composition internes, notées encore $+$ et $\times$. Alors $(\mathbb{P}, +)$ est un groupe abélien additif et $\times$ est associative et distributive par rapport à $+$ sur $\mathbb{P}$, ce qui montre que $\mathbb{P}$ — c'est en fait $(\mathbb{P}, +, \times)$ — est un pseudo-anneau. Si $\mathbb{P}$ était un anneau, $(\mathbb{P}, \times)$ posséderait un élément neutre q , donc on aurait $2 \times q = 2$, d'où l'égalité $q = 1$, en contradiction avec l'appartenance de q à $\mathbb{P}$. En conclusion, le pseudo-anneau $\mathbb{P}$ n'est pas un anneau. $\square$

DÉFINITION 3.8. — Un sous-anneau d'un anneau A est un sous-groupe S du groupe $(A, +)$ contenant l'élément unité de A et stable par la multiplication de A .

Si S est un sous-anneau de A et si l'on note encore $+$ et $\times$ les lois de composition internes induites sur S par l'addition et la multiplication de A , $(S, +, \times)$ est un anneau ayant le même zéro et le même élément unité que A .

1. Ce théorème est démontré, en 1905, indépendamment par le mathématicien écossais MacLagan Wedderburn (1882-1948) et le mathématicien américain Leonard Dickson (1874-1954). Pour une démonstration, voir [PERR], chapitre III, théorème 4.9.

Pour *démontrer* qu'une partie S d'un anneau A est un sous-anneau de A , il suffit d'établir que l'élément unité de A appartient à S , que S est stable par l'addition et la multiplication de A et que, pour tout élément x de S , $-x$ appartient à S .

3.4. Partie d'un anneau qui n'est pas un sous-anneau, mais qui est un anneau pour l'addition et la multiplication induites.

Nous munissons $\mathbb{Z}^2 = \mathbb{Z} \times \mathbb{Z}$ de l'addition produit et de la multiplication produit :

$$(p_1, p_2) + (q_1, q_2) = (p_1 + q_1, p_2 + q_2) \text{ et } (p_1, p_2) \times (q_1, q_2) = (p_1 q_1, p_2 q_2).$$

Alors $\mathbb{Z}^2$ — c'est en fait $(\mathbb{Z}^2, +, \times)$ — est un anneau commutatif dont le zéro est $(0, 0)$ et l'élément unité $(1, 1)$. Nous posons $S = \mathbb{Z} \times \{0\} = \{(n, 0) \mid n \in \mathbb{Z}\}$. On a, quels que soient $p, q \in \mathbb{Z}$, $(p, 0) + (q, 0) = (p + q, 0)$ et $(p, 0) \times (q, 0) = (pq, 0)$, donc S est stable par $+$ et $\times$. De plus, pour tout élément $x = (n, 0)$ de S , $-x = (-n, 0)$ appartient à S et $x \times (1, 0) = (n \times 1, 0) = (n, 0) = x$, donc S , muni de l'addition et de la multiplication induites par celles de $\mathbb{Z}^2$, est un anneau commutatif dont l'élément unité est $(1, 0)$. Or $(1, 1)$ n'appartient pas à S , donc S n'est pas un sous-anneau de « l'anneau produit » $\mathbb{Z}^2$. $\square$

Remarquons que cependant, dans l'exemple précédent 3.4, S est un idéal² de l'anneau produit $\mathbb{Z}^2$.

■ Éléments inversibles et diviseurs de zéro

DÉFINITION 3.9. — Un élément a d'un anneau A est inversible à gauche (resp. à droite) s'il existe un élément a' de A tel que $a'a = 1_A$ (resp. $aa' = 1_A$), et a est inversible s'il existe un élément a' de A tel que $a'a = aa' = 1_A$.

Si un élément a d'un anneau A est inversible, il existe un élément a' de A et un seul tel que $a'a = aa' = 1_A$.

DÉFINITION 3.10. — Si un élément a d'un anneau A est inversible, l'inverse de a est l'unique élément a^{-1} de A tel que $a^{-1}a = aa^{-1} = 1_A$.

3.5. Élément inversible à droite qui n'est pas inversible à gauche.

Nous posons $E = \mathbb{R}[X]$, nous notons D la dérivation canonique de $\mathbb{R}[X]$, c'est-à-dire l'endomorphisme $D : P \mapsto D(P) = P'$ de l'espace vectoriel réel E , et nous considérons l'application :

$$\left| \begin{array}{l} J : E \longrightarrow E \\ P = \sum_{n=0}^{+\infty} a_n X^n \longmapsto J(P) = \sum_{n=0}^{+\infty} \frac{a_n}{n+1} X^{n+1}. \end{array} \right.$$

Les applications D et J sont des éléments de l'anneau $\mathcal{L}(E)$ des endomorphismes de E , dont l'élément unité est l'application identique Id_E de E . On a clairement $D \circ J = \text{Id}_E$, donc l'endomorphisme D de E est inversible à droite dans $\mathcal{L}(E)$. Si D était inversible à gauche dans l'anneau $\mathcal{L}(E)$, il existerait un endomorphisme

2. Voir dans la suite de ce chapitre, page 43 et suivantes.

T de E tel que $T \circ D = \text{Id}_E$, donc D serait injectif, en contradiction avec l'égalité $D(X^0) = D(0) (= 0)$. En conclusion, la dérivation canonique D de $E = \mathbb{R}[X]$ est inversible à droite mais pas à gauche dans l'anneau $\mathcal{L}(E)$. $\square$

THÉORÈME 3.2. — L'ensemble $U(A)$ des éléments inversibles d'un anneau A est stable par la multiplication de A et, muni de la multiplication induite, $U(A)$ est un groupe multiplicatif dont l'élément neutre est l'élément unité 1_A de l'anneau A .

DÉFINITION 3.11. — Si A est un anneau, $(A, +)$ est le groupe additif de l'anneau A et $U(A)$ — c'est en fait $(U(A), \times)$ — est le groupe multiplicatif de l'anneau A .

3.6. Anneau qui n'est pas commutatif mais dont le groupe multiplicatif est commutatif.

Nous travaillons sur le corps $\mathbb{F}_2 = \mathbb{Z}/2\mathbb{Z}$ des entiers modulo 2 et nous posons :

$$\mathcal{A} = \left\{ \begin{pmatrix} x & y \\ 0 & z \end{pmatrix} \mid x, y, z \in \mathbb{F}_2 \right\}.$$

L'ensemble $\mathcal{A}$ est clairement un sous-anneau de l'anneau $\mathcal{M}_2(\mathbb{F}_2)$ des matrices carrées d'ordre 2 à coefficients dans $\mathbb{F}_2$.

En introduisant les éléments $P = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$ et $Q = \begin{pmatrix} 0 & 1 \\ 0 & 1 \end{pmatrix}$ de $\mathcal{A}$, on vérifie que :

$$PQ = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \text{ et } QP = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}$$

donc $\mathcal{A}$ n'est pas un anneau commutatif. On prouve facilement que les matrices :

$$I_2 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \text{ (l'élément unité de } \mathcal{M}_2(\mathbb{F}_2)) \text{ et } M = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$$

sont les deux seuls éléments inversibles de l'anneau $\mathcal{A}$; par suite $U(\mathcal{A}) = \{I_2, M\}$. Il en résulte que le groupe multiplicatif de l'anneau $\mathcal{A}$ est commutatif. $\square$

3.7. Anneau infini n'ayant qu'un seul élément inversible.

Nous travaillons sur le corps $\mathbb{F}_2 = \mathbb{Z}/2\mathbb{Z}$ des entiers modulo 2, nous choisissons un ensemble infini E et nous munissons l'ensemble $\mathcal{A} = \mathcal{F}(E, \mathbb{F}_2)$ des applications de E dans $\mathbb{F}_2$ de l'addition et de la multiplication des applications à valeurs dans un anneau. Alors $\mathcal{A}$ est un anneau commutatif³ dont le zéro est l'application $\mathbf{0}$ de E dans $\mathbb{F}_2$ constante égale à 0 et l'élément unité l'application j de E dans $\mathbb{F}_2$ constante égale à 1. Si f est un élément inversible de l'anneau $\mathcal{A}$, il existe une application g de E dans $\mathbb{F}_2$ telle que $fg = j$, ce qui montre que, pour tout élément x de E , $f(x)g(x) = (fg)(x) = 1$ donc $f(x) = g(x) = 1$, d'où l'on déduit que $f = j$. Ainsi $U(\mathcal{A}) = \{j\}$, alors que l'anneau $\mathcal{A}$ est infini. $\square$

3. Remarquons que si l'on note Δ la différence symétrique sur l'ensemble $\mathcal{P}(E)$ des parties de E et si l'on associe à toute partie Z de E sa fonction caractéristique :

$$\left| \begin{array}{l} \chi_Z : E \longrightarrow \mathbb{F}_2 \\ x \longmapsto \chi_Z(x) = \begin{cases} 1 & \text{si } x \in Z, \\ 0 & \text{si } x \in \complement_E Z, \end{cases} \end{array} \right.$$

l'application $\varphi : Z \mapsto \varphi(Z) = \chi_Z$ est un isomorphisme d'anneau de l'anneau $(\mathcal{P}(E), \Delta, \cap)$ sur l'anneau $\mathcal{A}$.

DÉFINITION 3.12. — Un élément a d'un anneau A est un diviseur de zéro à droite (resp. à gauche) de l'anneau A s'il existe au moins un élément non nul b de A tel que $ba = 0_A$ (resp. $ab = 0_A$), et a est un diviseur de zéro de A si c'est un diviseur de zéro à droite ou à gauche.

L'existence dans un anneau d'un diviseur de zéro à droite entraîne celle d'un diviseur de zéro à gauche. Cependant un élément donné d'un anneau peut être un diviseur zéro « d'un seul côté ».

3.8. Diviseur de zéro à droite mais pas à gauche.

Nous posons $E = \mathbb{R}[X]$, nous introduisons l'anneau $\mathcal{L}(E)$ des endomorphismes de l'espace vectoriel réel E et nous considérons les endomorphismes de E :

$$\left| \begin{array}{l} \theta : E \longrightarrow E \\ P \longmapsto \theta(P) = XP \end{array} \right. \quad \text{et} \quad \left| \begin{array}{l} \varphi : E \longrightarrow E \\ P = \sum_{n=0}^{+\infty} a_n X^n \longmapsto \varphi(P) = a_0 X^0. \end{array} \right.$$

Si P est un polynôme, le « terme constant » du produit XP est nul, donc on a $\varphi(\theta(P)) = 0$. Par conséquent $\varphi \circ \theta = 0$; or $\varphi \neq 0$, donc θ est un diviseur de zéro à droite de l'anneau $\mathcal{L}(E)$. Soit ψ un endomorphisme de E tel que $\theta \circ \psi = 0$. Si P est un polynôme, on a $\theta(\psi(P)) = 0$ et, l'endomorphisme θ étant injectif, il en résulte que $\psi(P) = 0$. On a ainsi établi que $\psi = 0$. En conclusion, θ n'est pas un diviseur de zéro à gauche de l'anneau $\mathcal{L}(E)$. $\square$

Un anneau A n'admet pas de diviseur de zéro si, et seulement si, quels que soient les éléments non nuls x et y de A , le produit xy est différent du zéro de A . Un anneau intègre est donc un anneau commutatif non nul sans diviseur de zéro.

3.9. Anneau commutatif non nul qui n'est pas un anneau intègre.

Nous reprenons l'anneau produit $\mathbb{Z}^2$ défini dans l'exemple 3.4 (page 35). Les couples $(1, 0)$ et $(0, 1)$ sont des éléments non nuls de l'anneau $\mathbb{Z}^2$ et leur produit $(1, 0) \times (0, 1)$ est égal à $(0, 0)$, donc $(1, 0)$ et $(0, 1)$ sont des diviseurs de zéro de l'anneau $\mathbb{Z}^2$, qui n'est donc pas un anneau intègre. $\square$

Ainsi, dans l'anneau produit $\mathbb{Z}^2$ des exemples 3.4 et 3.9, on ne peut pas toujours simplifier ; par exemple, $(1, 0) \times (0, 1) = (0, 0) = (1, 0) \times (0, 2)$ alors que $(0, 1) \neq (0, 2)$.

■ Anneau des polynômes à coefficients dans un anneau commutatif

Si A est un anneau commutatif non nul, on pose $0 = 0_A$ et $1 = 1_A$, et on associe à A l'anneau commutatif non nul $A[X]$ des polynômes à coefficients dans A .

DÉFINITION 3.13. — Si A est un anneau commutatif non nul et si $P = \sum_{n=0}^{+\infty} a_n X^n$ est un polynôme à coefficients dans A , le degré $\deg P$ et la valuation $\text{val } P$ du polynôme P sont les éléments de $\mathbb{N} \cup \{-\infty\}$ et $\mathbb{N} \cup \{+\infty\}$ définis respectivement ainsi : si $P = 0$, $\deg P = -\infty$ et $\text{val } P = +\infty$; si $P \neq 0$, $\deg P$ est le plus grand et $\text{val } P$ le plus petit des entiers naturels n tels que $a_n \neq 0$.

Si P et Q sont des polynômes à coefficients dans un anneau commutatif non nul A , alors $\deg(PQ) \leq \deg P + \deg Q$ et $\text{val}(PQ) \geq \text{val} P + \text{val} Q$, et si de plus A est un anneau intègre, ce sont des égalités.

3.10. Polynômes P et Q tels que $\deg(PQ) < \deg P + \deg Q$.

Nous travaillons dans l'anneau $A[X]$ des polynômes à coefficients dans l'anneau $A = \mathbb{Z}/4\mathbb{Z}$ des entiers modulo 4. Nous considérons le polynôme $P = 2X + 1$. Comme $P \times P = X^0$, $\deg(P \times P) = 0 < 2 = 1 + 1 = \deg P + \deg P$. $\square$

Remarquons que le polynôme P de l'exemple précédent 3.10 est inversible dans l'anneau $A[X]$ bien que ce ne soit pas un polynôme constant.

3.11. Polynômes P et Q tels que $\text{val}(PQ) > \text{val} P + \text{val} Q$.

Nous travaillons de nouveau dans l'anneau $A[X]$ des polynômes à coefficients dans l'anneau $A = \mathbb{Z}/4\mathbb{Z}$ des entiers modulo 4. Nous introduisons le polynôme $P = X + 2$. Comme $P \times P = X^2$, $\text{val}(P \times P) = 2 > 0 = 0 + 0 = \text{val} P + \text{val} P$. $\square$

Soit A un anneau commutatif non nul. Si $P = \sum_{n=0}^{+\infty} a_n X^n$ est un polynôme à coefficients dans A , on définit, pour tout élément x de l'anneau A , l'élément $P(x)$ de A par :

$$P(x) = \sum_{n=0}^{+\infty} a_n x^n \left(= \sum_{n=0}^p a_n x^n \text{ si } p \in \mathbb{N} \text{ et si } \deg P \leq p \right).$$

Quels que soient les polynômes P et Q à coefficients dans A , on a, pour tout $x \in A$:
 $(-P)(x) = -P(x)$, $(P + Q)(x) = P(x) + Q(x)$ et $(PQ)(x) = P(x)Q(x)$.

DÉFINITION 3.14. — Si P est un polynôme à coefficients dans A , une racine de P dans A est un élément x de A tel que $P(x) = 0$.

DÉFINITION 3.15. — Si P est un polynôme à coefficients dans A , l'application :

$$\left| \begin{array}{l} \tilde{P} : A \longrightarrow A \\ x \longmapsto \tilde{P}(x) = P(x) \end{array} \right.$$

s'appelle l'application polynomiale sur A associée au polynôme P .

Certains résultats qui nous sont familiers si l'anneau de base est intègre, *a fortiori* si c'est un corps commutatif, sont mis en défaut dans le cas général. Montrons par exemple que dans le cas où A n'est pas un anneau intègre, un polynôme de degré 1 n'a pas forcément de racine dans A et un polynôme de degré n peut admettre plus de n racines dans A .

3.12. Polynômes de degré 1 ayant zéro racine ou deux racines.

Nous introduisons l'anneau $A[X]$ des polynômes à coefficients dans l'anneau $A = \mathbb{Z}/4\mathbb{Z}$ des entiers modulo 4. Le polynôme $2X + 1$ n'admet pas de racine dans A , alors que le polynôme $2X$ admet 0 et 2 pour racines dans A . $\square$

3.13. Polynômes de degré 2 qui admet 4 racines.

Nous considérons l'anneau $A[X]$ des polynômes à coefficients dans l'anneau $A = \mathbb{Z}/12\mathbb{Z}$ des entiers modulo 12. Dans A , l'équation $x^2 = 4$ possède quatre

solutions : $x_1 = 2$, $x_2 = 4$, $x_3 = 8$ et $x_4 = 10$, et l'équation $x^2 + 3x + 2 = 0$ également quatre : $y_1 = 2$, $y_2 = 7$, $y_3 = 10$ et $y_4 = 11$. Ainsi les polynômes $X^2 - 4$ et $X^2 + 3X + 2$ sont de degré 2 et admettent tous deux quatre racines dans A . $\square$

Remarquons que le théorème affirmant qu'un élément a de A est une racine de P dans A si, et seulement si, $X - a$ divise⁴ P dans l'anneau $A[X]$, est valable pour l'anneau $A = \mathbb{Z}/12\mathbb{Z}$ de l'exemple précédent 3.13, mais que dans l'anneau $A[X]$, un polynôme peut avoir plusieurs factorisations différentes ; on a par exemple :

$$X^2 - 4 = (X - 2)(X - 10) = (X - 8)(X - 10)$$

et :

$$X^2 + 3X + 2 = (X - 2)(X - 7) = (X - 10)(X - 11).$$

3.14. Polynôme de degré 2 admettant une infinité de racines.

Nous choisissons un ensemble infini E et nous notons $+$ la différence symétrique et $\times$ l'intersection sur l'ensemble $\mathcal{P}(E)$ des parties de E ; on a donc, quelles que soient les parties X et Y de E , $X + Y = (X \setminus Y) \cup (Y \setminus X) = (X \cup Y) \setminus (X \cap Y)$ et $X \times Y = X \cap Y$. Alors $(\mathcal{P}(E), +, \times)$ est anneau commutatif non nul, dont le zéro est $\emptyset$ et dont l'élément unité est E et, pour toute partie X de E , $X + X = \emptyset$. Si nous choisissons un élément a de E et si nous posons $A = \{a\}$ et $B = \mathbb{C}_E A$, A et B sont des éléments non nuls de l'anneau $\mathcal{A} = (\mathcal{P}(E), +, \times)$ et $A \times B = \emptyset$, donc $\mathcal{A}$ n'est pas un anneau intègre. Pour toute partie X de E , $X \cap X = X$ donc $(X \cap X) + X = \emptyset$, ce qui montre que toute partie de E est une racine dans $\mathcal{A}$ du polynôme $X^2 + X$, qui admet donc une infinité de racines dans $\mathcal{A}$. $\square$

THÉORÈME 3.3. — Si K est un corps commutatif de caractéristique nulle⁵, P un polynôme non nul à coefficients dans K , a un élément de K et m un entier naturel, il y a équivalence entre les assertions :

- (i) $P(a) = P'(a) = \dots = P^{(m-1)}(a) = 0$ et $P^{(m)}(a) \neq 0$.
- (ii) $(X - a)^m$ divise P et $(X - a)^{m+1}$ ne divise pas P dans l'anneau $K[X]$.

Pour un corps K quelconque, $P \in K[X]$ et $a \in K$, l'entier naturel m défini dans l'assertion (ii) du théorème 3.3 est unique, on l'appelle l'ordre de multiplicité de a comme racine de P dans K et on dit que a est une racine d'ordre m de P dans K .

Le théorème 3.3 devient faux si K n'est pas de caractéristique nulle.

3.15. Polynôme dont 1 est une racine d'ordre p mais dont tous les polynômes dérivés admettent 1 pour racine.

Soit p un nombre premier p . Le corps $\mathbb{F}_p = \mathbb{Z}/p\mathbb{Z}$ des entiers modulo p est un corps commutatif de caractéristique p . Nous travaillons dans l'anneau $\mathbb{F}_p[X]$ des polynômes à coefficients dans $\mathbb{F}_p$. Nous posons $P = (X - 1)^p$. L'assertion (ii) du théorème 3.3 est vraie pour $a = 1$ et $m = p$. On a $P' = p(X - 1)^{p-1} = 0$ donc le polynôme $P^{(n)}$ est nul pour tout entier $n \geq 1$; de plus $P(1) = 0$, donc $P^{(n)}(1) = 0$ pour tout entier naturel n : l'assertion (i) du théorème 3.3 est fausse pour le polynôme P , l'élément $a = 1$ de $\mathbb{F}_p$ et n'importe quel entier naturel m . $\square$

4. Voir la définition 3.3 page 33.

5. Voir dans la suite de ce chapitre, page 42.

Soit K un corps commutatif. Soit P et Q des polynômes à coefficients dans K . Nous rappelons que P est associé à Q si P divise Q et Q divise P , ce qui équivaut à l'existence de $\lambda \in K \setminus \{0\}$ tel que $P = \lambda Q$. Si P divise Q , P est associé à Q si, et seulement si, P et Q ont le même degré.

DÉFINITION 3.16. — Un polynôme P à coefficients dans un corps commutatif K est irréductible sur K — ou dans l'anneau $K[X]$ — si $\deg P \geq 1$ et si tout diviseur de P dans $K[X]$ est un polynôme constant ou un polynôme de même degré que P — c'est-à-dire un polynôme constant ou associé à P .

Sur un corps commutatif, tout polynôme de degré 1 est irréductible.

DÉFINITION 3.17. — Un polynôme P à coefficients dans un corps commutatif K est réductible sur K si $\deg P \geq 1$ et si P n'est pas irréductible sur K .

Soit K un corps commutatif et $P \in K[X]$. Le polynôme P est réductible sur K si, et seulement si, $\deg P \geq 2$ et P admet au moins un diviseur D dans $K[X]$ tel que $1 \leq \deg D < \deg P$. Si $\deg P \geq 2$ et si P admet une racine a dans K , $D = X - a$ divise P dans $K[X]$ et $1 = \deg D < \deg P$, donc P est réductible sur K .

THÉORÈME 3.4. — Sur un corps commutatif K , un polynôme de degré 2 ou 3 est irréductible si, et seulement si, il n'admet pas de racine dans K .

Démontrons-le. Soit P un polynôme de degré 2 (resp. de degré 3) à coefficients dans K . Si P admet une racine dans K , P est réductible sur K — voir ce qui précède. Si P est réductible sur K , il admet dans $K[X]$ un diviseur D tel que $1 \leq \deg D < 2$ (resp. $1 \leq \deg D < 3$), donc un diviseur de degré 1, et ce diviseur admet une racine dans K , qui est aussi une racine de P .

3.16. Polynôme de degré 4 réductible sans racine.

La décomposition $X^4 + 1 = (X^2 - X\sqrt{2} + 1)(X^2 + X\sqrt{2} + 1)$ nous montre que le polynôme $X^4 + 1$ est réductible sur $\mathbb{R}$. Or $x^4 + 1 > 0$ pour tout réel x , donc $X^4 + 1$ n'admet pas de racine réelle. $\square$

Plus généralement, le produit de deux polynômes de degré 2 à coefficients réels de discriminants strictement négatifs est un polynôme de degré 4 réductible sur $\mathbb{R}$ et sans racine réelle.

THÉORÈME 3.5. — Si A est un anneau intègre infini, l'application :

$$\begin{array}{l} \Phi : A[X] \longrightarrow \mathcal{F}(A, A) \\ P \longmapsto \Phi(P) = \tilde{P} \text{ (application polynomiale associée à } P\text{).} \end{array}$$

est un morphisme d'anneau injectif.

3.17. Des polynômes différents dont les applications polynomiales associées sont égales.

Nous travaillons dans l'anneau $\mathbb{F}_3[X]$ des polynômes à coefficients dans le corps $\mathbb{F}_3 = \mathbb{Z}/3\mathbb{Z}$ des entiers modulo 3.

Nous posons $P = X$ et $Q = X^3$. On a $P(0) = 0 = Q(0)$, $P(1) = 1 = Q(1)$ et $P(2) = 2 = Q(2)$, donc $\tilde{P} = \tilde{Q}$, alors que $P \neq Q$. $\square$

Grâce au petit théorème de Fermat⁶, le résultat de l'exemple précédent 3.17 se généralise, pour tout nombre premier p , au corps $\mathbb{F}_p = \mathbb{Z}/p\mathbb{Z}$ des entiers modulo p et aux polynômes $P = X$ et $Q = X^p$.

Les polynômes irréductibles sur $\mathbb{C}$ sont les polynômes de degré 1, et les polynômes irréductibles sur $\mathbb{R}$ sont les polynômes de degré 1 et les polynômes de degré 2 de discriminant strictement négatif.

3.18. Polynôme irréductible de degré 3.

Nous travaillons dans l'anneau $\mathbb{Q}[X]$ des polynômes à coefficients dans le corps $\mathbb{Q}$ des nombres rationnels. Nous considérons le polynôme $P = X^3 - 2$. Supposons qu'il existe un nombre rationnel a tel que $a^3 = 2$. Notons (n, d) le représentant irréductible de a de dénominateur positif. On a, dans $\mathbb{Z}$, $n^3 = 2d^3$, donc 2 divise n^3 ; or 2 est un nombre premier, donc 2 divise n , ce qui prouve l'existence d'un entier relatif k tel que $n = 2k$. On a alors $8k^3 = 2d^3$ donc $4k^3 = d^3$, d'où l'on déduit que 2 divise d^3 et que par conséquent 2 divise d , en contradiction avec le fait que n et d sont premiers entre eux dans $\mathbb{Z}$. Il en résulte qu'il n'existe aucun nombre rationnel a tel que $a^3 = 2$, donc que le polynôme P n'admet pas de racine rationnelle. Le théorème 3.4 montre alors que $P = X^3 - 2$ est irréductible sur $\mathbb{Q}$. $\square$

3.19. Autre polynôme irréductible de degré 3.

Nous travaillons dans l'anneau $\mathbb{F}_5[X]$ des polynômes à coefficients dans le corps $\mathbb{F}_5 = \mathbb{Z}/5\mathbb{Z}$ des entiers modulo 5. Nous posons $P = X^3 + 4X^2 + X + 1$. En substituant successivement chacun des cinq éléments de $\mathbb{F}_5$ à X , on obtient : $P(0) = 1$, $P(1) = 2$, $P(2) = 2$, $P(3) = 2$ et $P(4) = 3$, et 2 et 3 sont différents de zéro dans $\mathbb{F}_5$, donc P n'admet pas de racine dans $\mathbb{F}_5$. Une nouvelle application du théorème 3.4 montre que le polynôme P est irréductible sur $\mathbb{F}_5$. $\square$

Si a et b sont des éléments d'un anneau commutatif non nul, a est associé à b si a divise b et b divise a , ce qui équivaut à l'existence d'un élément inversible u tel que $a = ub$. On retrouve ainsi, si K est un corps commutatif et si l'anneau est $K[X]$, la notion de polynôme associé à un autre rappelée en haut de la page 40.

DÉFINITION 3.18. — Un élément a d'un anneau commutatif non nul est irréductible s'il n'est ni nul ni inversible et si tout diviseur de a est inversible ou associé à a , et l'élément a est réductible si a n'est ni nul ni inversible et si a admet au moins un diviseur qui n'est ni inversible ni associé à a .

Si K est un corps commutatif, les éléments inversibles de l'anneau $K[X]$ sont les polynômes constants non nuls, c'est-à-dire les λX^0 pour $\lambda \in K \setminus \{0\}$, donc la définition 3.18 appliquée dans l'anneau $K[X]$ redonne les mêmes polynômes irréductibles et réductibles que les définitions 3.16 et 3.17 (page 40).

Si A est un anneau commutatif non nul, il en est de même de l'anneau $\mathcal{A} = A[X]$ donc la définition 3.18 s'applique dans $A[X]$, et on dit d'un polynôme à coefficients dans A qu'il est irréductible sur A s'il est irréductible dans l'anneau $A[X]$ et réductible sur A s'il est réductible dans l'anneau $A[X]$. Bien entendu, tout ceci s'applique à $A = \mathbb{Z}$.

6. C'est le théorème 1.4 du chapitre 1, page 16.

THÉORÈME 3.6. — Si p est un nombre premier, tout polynôme réductible sur $\mathbb{Z}$ est réductible⁷ sur le corps $\mathbb{F}_p = \mathbb{Z}/p\mathbb{Z}$ des entiers modulo p .

La réciproque est fautive, comme le montre l'exemple suivant.

3.20. Polynôme irréductible sur $\mathbb{Z}$ mais réductible sur $\mathbb{F}_p = \mathbb{Z}/p\mathbb{Z}$.

Les éléments inversibles de l'anneau $\mathbb{Z}$ étant 1 et -1 , les éléments inversibles de l'anneau $\mathbb{Z}[X]$ sont $1 = X^0$ et $-1 = -X^0$. Il en résulte que si Q est un polynôme à coefficients dans $\mathbb{Z}$, les éléments de l'anneau $\mathbb{Z}[X]$ associés à Q sont Q et $-Q$. Nous posons $P = X^4 + 1$, polynôme à coefficients dans $\mathbb{Z}$. Il se décompose :

$$P = (X^2 - X\sqrt{2} + 1)(X^2 + X\sqrt{2} + 1)$$

en produit de facteurs premiers dans $\mathbb{R}[X]$. Les diviseurs de P dans $\mathbb{R}[X]$ sont donc les $\lambda(X^2 - X\sqrt{2} + 1)^k(X^2 + X\sqrt{2} + 1)^\ell$ où $\lambda \in \mathbb{R} \setminus \{0\}$ et $k, \ell \in \{0, 1\}$. Ainsi un diviseur D de P dans $\mathbb{R}[X]$ est égal à λX^0 , $\lambda(X^2 - X\sqrt{2} + 1)$, $\lambda(X^2 + X\sqrt{2} + 1)$ ou λP avec $\lambda \in \mathbb{R} \setminus \{0\}$. Soit D un diviseur de P dans $\mathbb{Z}[X]$. Comme D divise P dans $\mathbb{R}[X]$ et que ses coefficients sont des entiers relatifs, D est égal à nX^0 , $n(X^2 - X\sqrt{2} + 1)$, $n(X^2 + X\sqrt{2} + 1)$ ou nP avec $n \in \mathbb{Z} \setminus \{0\}$ et, $\sqrt{2}$ étant irrationnel⁸, $n\sqrt{2}$ n'appartient pas à $\mathbb{Z}$ donc $D = nX^0$ ou $D = nP$. Nous notons Q le polynôme à coefficients dans $\mathbb{Z}$ tel que $P = DQ$ et q son coefficient dominant. Alors $q \in \mathbb{Z}$ et $1 = nq$, donc $n = \pm 1$, ce qui montre que $D = \pm X^0$ ou $D = \pm P$. En conclusion, tout diviseur de P dans $\mathbb{Z}[X]$ est inversible ou associé à P , donc le polynôme $P = X^4 + 1$ est irréductible⁹ sur $\mathbb{Z}$.

De plus, $X^4 + 1 = (X + 1)^4$ dans $\mathbb{F}_2[X]$ et $X^4 + 1 = (X^2 + X + 2)(X^2 + 2X + 2)$ dans $\mathbb{F}_3[X]$, donc, sur $\mathbb{F}_2$ et sur $\mathbb{F}_3$, le polynôme $P = X^4 + 1$ est réductible¹⁰. $\square$

■ Caractéristique d'un anneau non nul

Rappelons que si $(G, +)$ est un groupe abélien additif, alors, pour tout entier relatif n et tout élément x de G , $n \cdot x$ (lire : n fois x) est l'élément de G défini par :

$$0 \cdot x = 0, \quad n \cdot x = \underbrace{x + \cdots + x}_{n \text{ termes}} \text{ si } n \geq 1 \text{ et } n \cdot x = \underbrace{(-x) + \cdots + (-x)}_{(-n) \text{ termes}} \text{ si } n \leq -1.$$

On a, quels que soient les entiers relatifs n, p et q et les éléments x et y de G :

$$(p + q) \cdot x = p \cdot x + q \cdot x, \quad p \cdot (q \cdot x) = (pq) \cdot x \text{ et } n \cdot (x + y) = n \cdot x + n \cdot y.$$

DÉFINITION 3.19. — Soit A un anneau non nul. S'il existe au moins un entier relatif non nul n tel que $n \cdot 1 = 0$, la caractéristique de A est le plus petit des entiers k strictement positifs tels que $k \cdot 1 = 0$. Si $n \cdot 1 \neq 0$ pour tout entier relatif non nul n , on dit de l'anneau A qu'il est de caractéristique nulle.

7. En fait, $P = a_0 + a_1X + a_2X^2 + \cdots$ étant un polynôme à coefficients dans $\mathbb{Z}$, on dit de P qu'il est irréductible (resp. réductible) sur $\mathbb{F}_p$ si le polynôme $\bar{P} = \alpha_0 + \alpha_1X + \alpha_2X^2 + \cdots$, où α_i est, pour tout i , la classe de a_i modulo p , est irréductible (resp. réductible) sur $\mathbb{F}_p$.

8. Voir l'exemple 5.3, page 84.

9. On peut le prouver plus rapidement à l'aide du critère d'Eisenstein et du « changement d'indéterminée » $Y = X - 1$; voir [PERR], chapitre III, théorème 3.2 et proposition 3.11.

10. On peut en fait démontrer que, pour tout nombre premier p , $X^4 + 1$ est réductible sur le corps commutatif $\mathbb{F}_p = \mathbb{Z}/p\mathbb{Z}$; voir [PERR], chapitre III, proposition 3.11.

Par exemple, pour tout entier $p \geq 2$, l'anneau $\mathbb{Z}/p\mathbb{Z}$ des entiers modulo p est un anneau commutatif non nul de caractéristique p .

Si un anneau non nul A est de caractéristique p , l'ensemble des entiers relatifs n tels que $n \cdot 1 = 0$ est l'ensemble $p\mathbb{Z} = \{kp \mid k \in \mathbb{Z}\}$ des multiples de p et, comme $1 \neq 0$, nous voyons que si $p \neq 0$, p est supérieur ou égal à 2.

Soit A un anneau intègre de caractéristique non nulle p . Alors p est un entier supérieur ou égal à 2. Soit n un diviseur positif de p . Il existe $m \in \mathbb{N}^*$ tel que $nm = p$, donc on a, dans l'anneau A , $0 = p \cdot 1 = (n \cdot 1)(m \cdot 1)$, ce qui montre que $n \cdot 1 = 0$ ou $m \cdot 1 = 0$. Par suite, p divise n ou m , donc $n = p$ ou $m = p$, d'où l'on déduit que $n = p$ ou $n = 1$. L'entier p est donc un nombre premier.

Ainsi la caractéristique d'un anneau intègre est nulle ou un nombre premier.

3.21. Anneau qui n'est pas intègre mais dont la caractéristique est un nombre premier.

Le corps commutatif $\mathbb{F}_5 = \mathbb{Z}/5\mathbb{Z}$ des entiers modulo 5 est un anneau non nul de caractéristique 5. Nous considérons l'anneau produit ¹¹ $A = \mathbb{F}_5 \times \mathbb{F}_5$ dont le zéro est $(0, 0)$ et dont l'élément unité est $(1, 1)$. L'anneau A est commutatif et n'est pas nul, mais ce n'est pas un anneau intègre ; en effet, $(1, 0)(0, 1) = (0, 0)$ alors que $(1, 0) \neq (0, 0)$. Pour tout entier relatif n , on a, en notant $\bar{n}$ la classe de n modulo 5, $n \cdot (1, 1) = (n \cdot 1, n \cdot 1) = (\bar{n}, \bar{n})$, donc $n \cdot (1, 1)$ est nul dans A si, et seulement si, 5 divise n . Il en résulte que la caractéristique de A est égale à 5. $\square$

Si un anneau non nul A est de caractéristique nulle, l'application $n \mapsto n \cdot 1$ est un morphisme d'anneau injectif de $\mathbb{Z}$ dans A , donc l'ensemble A est infini. Cependant un anneau infini n'est pas forcément de caractéristique nulle.

3.22. Anneau infini qui n'est pas de caractéristique nulle.

Nous choisissons un ensemble infini E et nous reprenons l'anneau commutatif non nul $\mathcal{A} = (\mathcal{P}(E), +, \times)$ défini dans l'exemple 3.14 (page 39) : $+$ est la différence symétrique et $\times$ l'intersection. Alors $\mathcal{A}$ est infini, son zéro est $\emptyset$ et son élément unité est E . Or $E \neq \emptyset$ et $2 \cdot E = E + E = \emptyset$, donc $\mathcal{A}$ est de caractéristique 2. $\square$

■ Idéaux d'un anneau non nul

DÉFINITION 3.20. — Soit A un anneau non nul. Un idéal de A à droite (resp. à gauche) est une partie I de A contenant 0, stable par l'addition et telle que, quels que soient les éléments x de A et y de I , le produit yx (resp. xy) appartient encore à I . Un idéal bilatère de A est un idéal de A à droite et à gauche.

Si I est un idéal à droite, à gauche ou bilatère d'un anneau non nul A , alors I est un sous-groupe de $(A, +)$ et $I = A$ si, et seulement si, l'élément unité de A appartient à I . Si A est un anneau commutatif non nul, les trois notions : « idéal de A à droite », « idéal de A à gauche » et « idéal bilatère de A » se confondent en la notion d'idéal de l'anneau A .

11. Voir l'exemple 3.4, page 35, dans lequel on remplace $\mathbb{Z}$ par $\mathbb{F}_5$.

DÉFINITION 3.21. — Un anneau non nul A est simple si ses seuls idéaux bilatères sont $\{0\}$ et A .

Un anneau commutatif non nul est un corps si, et seulement si, c'est un anneau simple.

3.23. Anneau simple qui n'est pas un corps.

Nous considérons l'anneau $\mathcal{A} = \mathcal{M}_2(\mathbb{R})$ des matrices carrées d'ordre 2 à coefficients réels. C'est un anneau non nul, son élément unité est la matrice unité I_2 et ce n'est pas un anneau commutatif — voir l'exemple 3.1 (page 34). Nous posons :

$$P = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}, \quad Q = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix} \quad \text{et} \quad R = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}.$$

La matrice P n'étant pas inversible, l'anneau $\mathcal{A}$ n'est pas un corps.

Soit $\mathcal{J}$ un idéal bilatère de $\mathcal{A}$, différent de $\{0\}$. Nous choisissons un élément :

$$M = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

non nul de l'idéal $\mathcal{J}$. Nous supposons d'abord que $a \neq 0$. Le calcul donne :

$$\left(\frac{1}{a}P\right)MP = P$$

donc $P \in \mathcal{J}$. Nous obtenons aussi l'égalité $RPR = Q$, donc $Q \in \mathcal{J}$. Comme $I_2 = P + Q$, I_2 appartient à $\mathcal{J}$. Supposons enfin que $a = 0$. Alors b, c ou d est différent de zéro et on obtient par le calcul que :

$$M_1 = MR = \begin{pmatrix} b & a \\ d & c \end{pmatrix}, \quad M_2 = RM = \begin{pmatrix} c & d \\ a & b \end{pmatrix} \quad \text{et} \quad M_3 = RM_1 = \begin{pmatrix} d & c \\ b & a \end{pmatrix}.$$

Les matrices M_1, M_2 et M_3 appartiennent donc à $\mathcal{J}$ et, en remplaçant M par M_1, M_2 ou M_3 , le même raisonnement que pour $a \neq 0$ montre que $I_2 \in \mathcal{J}$. Ainsi, dans tous les cas, $I_2 \in \mathcal{J}$ donc $\mathcal{J} = \mathcal{A}$. En conclusion, $\mathcal{A}$ est un anneau simple¹². $\square$

Soit A un anneau commutatif non nul. Pour l'étude de la divisibilité dans A , on introduit, pour tout $a \in A$, l'ensemble (a) des multiples de a dans A , c'est-à-dire :

$$(a) = \{xa \mid x \in A\}.$$

Pour tout élément a de A , (a) est l'idéal de A engendré par a , c'est-à-dire le plus petit pour l'inclusion des idéaux de A contenant a . De plus, quels que soient les éléments a et b de A , a divise b dans A si, et seulement si, $(b) \subset (a)$.

DÉFINITION 3.22. — Si A est un anneau commutatif non nul, un idéal principal de A est un idéal de l'anneau A engendré par un seul élément, c'est-à-dire un idéal $I = (a) = \{xa \mid x \in A\}$ où a est un élément de A .

DÉFINITION 3.23. — Un anneau principal est un anneau intègre dont tout idéal est un idéal principal.

12. On peut abrégé cette démonstration si l'on sait que deux matrices sont équivalentes si, et seulement si, elles ont le même rang.

Les sous-groupes de $(\mathbb{Z}, +)$ sont les $n\mathbb{Z} = \{kn \mid k \in \mathbb{Z}\}$ pour n parcourant $\mathbb{Z}$ (en fait $\mathbb{N}$ suffit) et, pour tout entier relatif n , on a $n\mathbb{Z} = (n)$, donc les idéaux de l'anneau $\mathbb{Z}$ sont les (n) pour n parcourant $\mathbb{Z}$ ($\mathbb{N}$ suffit) : l'anneau intègre $\mathbb{Z}$ est un anneau principal. Pour tout corps commutatif K , l'anneau $K[X]$ des polynômes à coefficients dans K est un anneau principal.

Nous avons vu dans ce qui précède ce que sont les éléments irréductibles d'un anneau commutatif non nul (définition 3.18, page 41).

Nous travaillons dans l'anneau intègre $\mathbb{Z}$. Si p est un entier naturel non nul, il y a équivalence entre les assertions :

- (i) p est un nombre premier,
- (ii) p est irréductible dans $\mathbb{Z}$,
- (iii) p n'est pas inversible dans $\mathbb{Z}$ et, pour tout couple (m, n) d'entiers relatifs, si p divise le produit mn , alors p divise m ou p divise n .

En s'exprimant en termes d'idéaux, nous obtenons, pour tout entier naturel non nul p , l'équivalence entre les assertions :

- (1) p est un nombre premier,
- (2) $(p) = p\mathbb{Z}$ est un idéal strict de l'anneau $\mathbb{Z}$ et il n'existe aucun idéal I de $\mathbb{Z}$ tel que $(p) \subsetneq I \subsetneq \mathbb{Z}$,
- (3) $(p) = p\mathbb{Z}$ est un idéal strict de $\mathbb{Z}$ et, pour tout couple (m, n) d'entiers relatifs, si mn appartient à (p) , alors $m \in (p)$ ou $n \in (p)$.

Compte tenu de l'équivalence entre (1) et (3), on introduit la définition suivante.

DÉFINITION 3.24. — Un idéal premier d'un anneau commutatif non nul A est un idéal strict I de A tel que, quels que soient les éléments x et y de l'anneau A , $xy \in I$ entraîne $x \in I$ ou $y \in I$.

DÉFINITION 3.25. — Un idéal maximal d'un anneau commutatif non nul A est un élément maximal de l'ensemble ordonné par l'inclusion des idéaux stricts de A .

Un idéal maximal d'un anneau commutatif non nul A est donc un idéal strict M de A tel que, pour tout idéal I de A , $M \subset I$ entraîne $I = M$ ou $I = A$, ce qui signifie qu'il n'existe aucun idéal I de A tel que $M \subsetneq I \subsetneq A$. On déduit des équivalences entre (1), (2) et (3) que si p est un entier naturel non nul, p est un nombre premier si, et seulement si, (p) est un idéal maximal de l'anneau $\mathbb{Z}$, et p est un nombre premier si, et seulement si, (p) est un idéal premier de l'anneau $\mathbb{Z}$.

Si A est un anneau commutatif, tout idéal maximal de A est un idéal premier de A et, a étant un élément de A , si (a) est un idéal premier de A , alors a est irréductible. Cependant, les réciproques sont en général fausses.

3.24. Idéal premier qui n'est pas un idéal maximal.

Nous considérons l'anneau produit $\mathcal{A} = \mathbb{Z} \times \mathbb{Z}$ — voir l'exemple 3.4 (page 35). C'est un anneau commutatif non nul de zéro $(0, 0)$ et d'élément unité $(1, 1)$.

Nous posons $\mathcal{J} = \{0\} \times \mathbb{Z}$. Clairement, $(0, 0) \in \mathcal{J}$ et $\mathcal{J}$ est stable par l'addition ; si $x = (p, q)$ est un élément de $\mathcal{A}$ et $y = (0, n)$ un élément de $\mathcal{J}$, alors $xy = (0, qn)$ donc xy appartient à $\mathcal{J}$. Par conséquent, $\mathcal{J}$ est un idéal de $\mathcal{A}$.

Nous posons $\mathcal{J} = (2\mathbb{Z}) \times \mathbb{Z} = \{(2p, q) \mid p, q \in \mathbb{Z}\}$. On prouve comme pour $\mathcal{I}$ que $\mathcal{J}$ est un idéal de $\mathcal{A}$ et, comme $(1, 1)$ n'appartient pas à $\mathcal{J}$, c'est un idéal strict. De plus $\mathcal{J}$ est inclus dans $\mathcal{J}$ et $(2, 0)$ appartient à $\mathcal{J}$ mais pas à $\mathcal{I}$, donc $\mathcal{I} \subsetneq \mathcal{J}$, d'où finalement les inclusions strictes $\mathcal{I} \subsetneq \mathcal{J} \subsetneq \mathcal{A}$, qui montrent que $\mathcal{I}$ n'est pas un idéal maximal de $\mathcal{A}$. Soit $x = (m, n)$ et $y = (p, q)$ des éléments de $\mathcal{A}$ tels que $xy \in \mathcal{I}$. Alors (mp, nq) appartient à $\mathcal{I}$, donc $mp = 0$. Il en résulte que $m = 0$ ou $p = 0$; par conséquent x appartient à $\mathcal{I}$ ou y appartient à $\mathcal{I}$. Ainsi $\mathcal{I}$ est un idéal premier de $\mathcal{A}$. $\square$

Remarquons que l'anneau $\mathcal{A}$ de l'exemple précédent 3.24 n'est pas intègre.

3.25. Idéal premier d'un anneau intègre qui n'est pas maximal.

Nous considérons l'anneau $\mathcal{A} = \mathbb{Z}[X]$ des polynômes à coefficients entiers et les idéaux :

$$\mathcal{I} = (X) = \{XR \mid R \in \mathbb{Z}[X]\} \text{ et } \mathcal{J} = \{2Q + XR \mid Q, R \in \mathbb{Z}[X]\}$$

de l'anneau $\mathcal{A}$; $\mathcal{I}$ est l'idéal engendré par X et $\mathcal{J}$ l'idéal engendré par $\{2X^0, X\}$.

Clairement, $\mathcal{I}$ est l'ensemble des polynômes $\sum_{n=0}^{+\infty} a_n X^n$ à coefficients entiers dont le « terme constant » a_0 est nul et $\mathcal{J}$ l'ensemble des polynômes à coefficients entiers dont le terme constant est pair. Par suite $\mathcal{I}$ est inclus dans $\mathcal{J}$ et $\mathcal{J}$ est un idéal strict de $\mathcal{A}$. De plus, $2X^0$ appartient à $\mathcal{J}$ mais n'appartient pas à $\mathcal{I}$, donc $\mathcal{I} \subsetneq \mathcal{J}$. Les inclusions strictes $\mathcal{I} \subsetneq \mathcal{J} \subsetneq \mathcal{A}$ montrent que $\mathcal{I}$ n'est pas un idéal maximal de $\mathcal{A}$.

Soit $P = \sum_{n=0}^{+\infty} a_n X^n$ et $Q = \sum_{n=0}^{+\infty} b_n X^n$ des éléments de $\mathcal{A}$ tels que $PQ \in \mathcal{I}$. Comme $PQ = \sum_{n=0}^{+\infty} c_n X^n$ où $c_0 = a_0 b_0$, que c_0 est nul et que $\mathbb{Z}$ est un anneau intègre, $a_0 = 0$ ou $b_0 = 0$, donc P appartient à $\mathcal{I}$ ou Q appartient à $\mathcal{I}$. Par suite $\mathcal{I}$ est un idéal premier. De plus $\mathcal{A}$ est un anneau intègre. $\square$

3.26. Élément irréductible p d'un anneau intègre tel que (p) n'est pas un idéal premier.

Nous considérons la partie $\mathcal{A} = \{a + ib\sqrt{5} \mid a, b \in \mathbb{Z}\}$ de $\mathbb{C}$. Comme $1 = 1 + i0\sqrt{5}$, 1 appartient à $\mathcal{A}$. Si $z = a + ib\sqrt{5}$ et $z' = c + id\sqrt{5}$ sont des éléments de $\mathcal{A}$, on a $z + z' = (a + c) + i(b + d)\sqrt{5}$ et $zz' = (ac - 5bd) + i(ad + bc)\sqrt{5}$, donc $z + z'$ et zz' appartiennent à $\mathcal{A}$. Enfin, si $z = a + ib\sqrt{5}$ est un élément de $\mathcal{A}$, son opposé $-z = (-a) + i(-b)\sqrt{5}$ appartient à $\mathcal{A}$. Par suite $\mathcal{A}$ est un sous-anneau du corps commutatif $\mathbb{C}$, donc, muni de l'addition et de la multiplication induites par celles de $\mathbb{C}$, $\mathcal{A}$ est un anneau intègre.

Nous posons, pour tout élément $z = a + ib\sqrt{5}$ de $\mathcal{A}$, $N(z) = |z|^2 = a^2 + 5b^2$. Les propriétés du module sur $\mathbb{C}$ montrent que, quels que soient $z \in \mathcal{A}$ et $z' \in \mathcal{A}$:

$$N(zz') = N(z)N(z').$$

De plus, pour tout élément z de $\mathcal{A}$, $N(z)$ est un entier naturel, donc, si z et z' sont des éléments de $\mathcal{A}$ et si z divise z' dans l'anneau $\mathcal{A}$, $N(z)$ divise $N(z')$ dans $\mathbb{N}$.

Si z est inversible dans l'anneau $\mathcal{A}$, z divise 1 dans $\mathcal{A}$, donc $N(z)$ divise $N(1) = 1$ dans $\mathbb{N}$, ce qui prouve que $N(z) = 1$. Réciproquement, si $z \in \mathcal{A}$ et si $N(z) = 1$, on a $1 = z\bar{z}$, donc, $\bar{z}$ appartenant clairement à $\mathcal{A}$, z est inversible dans $\mathcal{A}$. Les éléments inversibles de l'anneau $\mathcal{A}$ sont donc les éléments z de $\mathcal{A}$ tels que $N(z) = 1$. De plus, si $z = a + ib\sqrt{5}$ appartient à $\mathcal{A}$ et si $N(z) = 1$, alors $1 = a^2 + 5b^2$, $a^2 \in \mathbb{N}$ et $b^2 \in \mathbb{N}$, donc $b^2 = 0$ et $a^2 = 1$, ce qui montre que $z = \pm 1$; la réciproque étant immédiate, les éléments inversibles de $\mathcal{A}$ sont 1 et -1 .

Nous posons $p = 2 + i\sqrt{5}$. Alors p appartient à $\mathcal{A}$ et $N(p) = 2^2 + 5 \times 1^2 = 9$. On a $9 = |p|^2 = p\bar{p}$ et $\bar{p} = 2 - i\sqrt{5}$ appartient à $\mathcal{A}$, donc 9 appartient à (p) . Supposons que (p) soit un idéal premier de $\mathcal{A}$. Comme $3 \times 3 = 9$ appartient à (p) , $3 \in (p)$ donc il existe un élément z de $\mathcal{A}$ tel que $3 = zp$. On a $9 = N(3) = N(z)N(p) = N(z) \times 9$, donc $N(z) = 1$. Par conséquent $z = \pm 1$, donc $3 = \pm p = (\pm 2) + i(\pm 1)\sqrt{5}$, ce qui est bien sûr faux. Il en résulte que (p) n'est pas un idéal premier de l'anneau $\mathcal{A}$. Comme $N(p) = 9 \neq 1$, p n'est pas inversible dans $\mathcal{A}$. Soit $z = a + ib\sqrt{5}$ un diviseur non inversible de p dans $\mathcal{A}$. Alors $N(z) \neq 1$ et $N(z)$ divise $N(p) = 9$ dans $\mathbb{N}$, donc $N(z) = 3$ ou $N(z) = 9$. Si $N(z) = 3$, alors $a^2 + 5b^2 = 3$ donc $b^2 = 0$ et $a^2 = 3$, en contradiction avec le fait que 3 n'est pas un carré dans $\mathbb{N}$. Il en résulte que $N(z) = 9$. Nous notons u l'élément de $\mathcal{A}$ tel que $p = uz$. Alors $9 = N(p) = N(u)N(z) = 9N(u)$, donc $N(u) = 1$; par conséquent u est inversible dans $\mathcal{A}$, donc z est associé à p . En conclusion, p est un élément irréductible de $\mathcal{A}$. $\square$

■ Divisibilité dans un anneau intègre

Les propriétés les plus usitées de la divisibilité dans $\mathbb{Z}$ sont les suivantes :

- (1) l'existence des pgcd et des ppcm,
- (2) la décomposition unique en produit de facteurs premiers,
- (3) le théorème de Bézout,
- (4) la division euclidienne.

La divisibilité s'étudie dans un anneau intègre, pour ne pas avoir à considérer les notions de diviseurs d'un élément à gauche ou à droite — grâce à la commutativité de la multiplication — et pour éviter les diviseurs de zéro. Pour généraliser les propriétés (1), (2), (3) et (4) de la divisibilité dans $\mathbb{Z}$, différents types d'anneau ont été introduits, en particulier les anneaux *factoriels*, les anneaux *euclidiens* et bien sûr les anneaux principaux.

DÉFINITION 3.26. — Un anneau factoriel est un anneau intègre A dans lequel tout élément non nul et non inversible admet une décomposition en produit de facteurs irréductibles «unique à l'ordre près des facteurs» — ce qui signifie que tout élément non nul et non inversible de A est un produit fini d'éléments irréductibles de l'anneau A et que, si $n, m \in \mathbb{N}^*$, si $p_1, \dots, p_n, q_1, \dots, q_m$ sont des éléments irréductibles de A et si $p_1 \times \dots \times p_n = q_1 \times \dots \times q_m$, alors $m = n$ et il existe une permutation σ de $\llbracket 1, n \rrbracket$ telle que q_i est associé à $p_{\sigma(i)}$ pour tout $i \in \llbracket 1, n \rrbracket$.

DÉFINITION 3.27. — Un anneau euclidien est un anneau intègre A muni d'une application φ de $A \setminus \{0\}$ dans $\mathbb{N}$ telle que, quels que soient $a \in A$ et $b \in A$:

- (1) Si $a \neq 0$, si $b \neq 0$ et si a divise b dans A , alors $\varphi(a) \leq \varphi(b)$.
- (2) Si $b \neq 0$, il existe $q, r \in A$ tels que $a = bq + r$, et $r = 0$ ou $\varphi(r) < \varphi(b)$.

Nous disposons des implications suivantes :

- si A est un corps commutatif, A est un anneau euclidien,
- si A est un anneau euclidien, A est un anneau principal,

- si A est un anneau principal, A est un anneau factoriel,
- si A est un anneau factoriel, alors, quels que soient les éléments a et b de A , le couple (a, b) admet un pgcd et un ppcm,

mais les implications réciproques sont en général fausses.

3.27. Anneau intègre qui n'est pas un anneau factoriel.

Nous reprenons l'anneau intègre $\mathcal{A} = \{a + ib\sqrt{5} \mid a, b \in \mathbb{Z}\}$ de l'exemple 3.26 (pages 46 et 47). L'ensemble des éléments inversibles de $\mathcal{A}$ est :

$$U(\mathcal{A}) = \{z \in \mathcal{A} \mid N(z) = 1\} = \{1, -1\},$$

$p = 2 + i\sqrt{5}$ est un élément irréductible de $\mathcal{A}$ et $9 = p\bar{p}$. La même démonstration que pour p montre que $q = \bar{p} = 2 - i\sqrt{5}$ et 3 sont des éléments irréductibles de $\mathcal{A}$; en effet, cette preuve repose simplement sur le fait que $N(p) = 9$, et on a $N(q) = N(3) = 9$. Les éléments associés à un élément z de $\mathcal{A}$ sont z et $-z$, donc 3 n'est associé ni à p , ni à q . On a $3 \times 3 = p \times q (= 9)$ donc, si $\mathcal{A}$ était un anneau factoriel, 3 serait associé à p ou à q , en contradiction avec les résultats précédents. Ainsi l'anneau intègre $\mathcal{A}$ n'est pas un anneau factoriel. $\square$

3.28. Anneau intègre dans lequel deux éléments n'ont pas toujours un pgcd.

Nous utilisons de nouveau l'anneau $\mathcal{A}$ des deux exemples précédents 3.26 et 3.27. Nous rappelons que si $z \in \mathcal{A}$, $N(z) = 1$ si, et seulement si, $z = \pm 1$. Nous posons :

$$p = 2 + i\sqrt{5}, \quad q = \bar{p} = 2 - i\sqrt{5}, \quad a = pq = 9 \quad \text{et} \quad b = 6 + i3\sqrt{5} = 3p.$$

Supposons que a et b possèdent un pgcd d dans l'anneau $\mathcal{A}$. Ceci signifie que, dans $\mathcal{A}$, d est un diviseur commun à a et à b et tout diviseur commun à a et b divise d . Comme 3 et p sont des diviseurs communs à a et à b , 3 et p divisent d dans $\mathcal{A}$. Par suite, $9 = N(3) = N(p) = \text{divise } N(d)$ et $N(d)$ divise $81 = N(a) = N(b)$ dans $\mathbb{N}$, donc $N(d)$ est égal à 9 , 27 ou 81 . Notons u et v les éléments de $\mathcal{A}$ tels que $d = 3u$ et $d = pv$. Si $N(d) = 9$, alors $9 = N(3u) = N(3)N(u) = 9N(u)$ et $9 = N(pv) = N(p)N(v) = 9N(v)$, donc $N(u) = N(v) = 1$, ce qui montre que $u = \pm 1$ et $v = \pm 1$, d'où l'on déduit que $3 = \pm p$, égalité manifestement fautive ; par suite $N(d) \neq 9$. En introduisant les éléments x et y de $\mathcal{A}$ tels que $a = dx$ et $b = dy$, on prouve que si $N(d) = 81$, alors $a = \pm b$, ce qui est faux. Par suite $N(d) \neq 81$ donc $N(d) = 27$. On voit facilement que l'équation $m^2 + 5n^2 = 27$ n'admet pas de solution (m, n) dans $\mathbb{N}^2$, en contradiction avec l'égalité $N(d) = 27$.

Ainsi a et b n'ont pas de pgcd dans l'anneau $\mathcal{A}$. $\square$

Remarquons que, dans l'anneau $\mathcal{A}$ des trois exemples précédents, le théorème de Gauss¹³ est faux. En effet, dans l'anneau $\mathcal{A}$, 3 divise $9 = pq$, 3 est premier avec p et 3 ne divise pas q — les deux dernières affirmations se prouvent par les mêmes méthodes que dans l'exemple 3.28.

THÉORÈME 3.7. — Théorème de Gauss¹⁴.

Si A est un anneau factoriel, alors $A[X]$ est un anneau factoriel.

13. Si a divise bc et si a est premier avec b , alors a divise c .

14. Pour une démonstration, voir [PERR], chapitre II, théorème 4.2.

3.29. Anneau factoriel qui n'est pas un anneau principal.

Nous reprenons l'anneau $\mathcal{A} = \mathbb{Z}[X]$ des polynômes à coefficients entiers et l'idéal $\mathcal{J} = \{2Q + XR \mid Q, R \in \mathcal{A}\}$ de $\mathcal{A}$ engendré par $\{2X^0, X\}$, qui ont été introduits dans l'exemple 3.25 (page 46). Rappelons que $\mathcal{J}$ est l'ensemble des polynômes $\sum_{n=0}^{+\infty} a_n X^n$ à coefficients entiers dont le « terme constant » a_0 est pair, d'où l'on déduit que $\mathcal{J}$ est un idéal strict de $\mathcal{A}$.

Supposons que $\mathcal{J}$ soit un idéal principal de l'anneau $\mathcal{A}$. Il existe un polynôme P à coefficients dans $\mathbb{Z}$ tel que $\mathcal{J} = (P) = \{PQ \mid Q \in \mathcal{A}\}$. Comme $2X^0$ et X appartiennent à $\mathcal{J}$, il existe des éléments Q et R de $\mathcal{A}$ tel que $2X^0 = PQ$ et $X = PR$. L'anneau $\mathbb{Z}$ étant un anneau intègre, $0 = \deg(PQ) = \deg P + \deg Q$, donc $\deg P = 0$, d'où l'existence d'un entier relatif non nul p tel que $P = pX^0$. De plus $1 = X(1) = P(1)R(1) = pR(1)$, donc p divise 1 dans $\mathbb{Z}$, ce qui montre que $p = \pm 1$. Par suite $P = \pm X^0$, donc $X^0 = \pm P$ appartient à $\mathcal{J}$. Or X^0 est l'élément unité de l'anneau $\mathcal{A}$, donc $\mathcal{J} = \mathcal{A}$, en contradiction avec $\mathcal{J} \subsetneq \mathcal{A}$. Il en résulte que $\mathcal{J}$ n'est pas un idéal principal de $\mathcal{A}$.

En conclusion, $\mathcal{A}$ n'est pas un anneau principal.

Comme $\mathbb{Z}$ est un anneau factoriel — car principal —, on déduit du théorème 3.7 (page 48) que $\mathcal{A}$ est un anneau factoriel. $\square$

Remarquons que le théorème de Bézout n'est pas vérifié dans l'anneau $\mathcal{A} = \mathbb{Z}[X]$. En effet l'exemple précédent 3.29 montre que, dans $\mathcal{A}$, tout diviseur commun à $2X^0$ et à X dans $\mathcal{A}$ est égal à $\pm X^0$, donc X^0 est un pgcd $(2X^0, X)$; or $X^0 \notin \mathcal{J}$, donc il n'existe aucun couple (Q, R) d'éléments de $\mathcal{A}$ tels que $(2X^0)Q + XR = X^0 (= 1)$.

En fait le théorème de Bézout est vérifié dans une classe plus large d'anneaux intègres que celle des anneaux principaux, ceux dans lesquels l'intersection de deux idéaux principaux est un idéal principal; de tels anneaux sont appelés les anneaux de Gauss et nous avons prouvé que $\mathbb{Z}[X]$ n'est pas un anneau de Gauss.

3.30. Anneau principal qui n'est pas un anneau euclidien.

Posons $w = \frac{1}{2} + i\frac{\sqrt{19}}{2}$ et $\mathcal{A} = \{a + bw \mid a, b \in \mathbb{Z}\}$. On a $w^2 = -\frac{9}{2} + i\frac{\sqrt{19}}{2} = -5 + w$.

On en déduit facilement que $\mathcal{A}$ est un sous-anneau du corps commutatif $\mathbb{C}$. On a $\bar{w} = 1 - w$ donc, pour tout élément z de $\mathcal{A}$, $\bar{z} \in \mathcal{A}$. Nous munissons $\mathcal{A}$ de l'addition et de la multiplication induites par celles de $\mathbb{C}$. Alors $\mathcal{A}$ est un anneau intègre.

Nous posons, pour tout élément z de $\mathcal{A}$, $N(z) = z\bar{z} = |z|^2$. Le calcul donne : $w + \bar{w} = 1$ et $N(w) = w\bar{w} = |w|^2 = 5$, donc, pour tout élément $z = a + bw$ de $\mathcal{A}$, $N(z) = (a + bw)(a + b\bar{w}) = a^2 + ab + 5b^2$, ce qui montre que $N(z)$ est un nombre entier, donc un entier naturel, et on obtient, en passant par $\mathbb{Q}$, l'égalité :

$$(E) \quad N(z) = \left(a + \frac{b}{2}\right)^2 + \frac{19b^2}{4}.$$

Compte tenu des propriétés du module sur $\mathbb{C}$, on a, quels que soient les éléments z et z' de $\mathcal{A}$, $N(zz') = N(z)N(z')$, d'où l'on déduit que si z divise z' dans l'anneau $\mathcal{A}$, alors $N(z)$ divise $N(z')$ dans $\mathbb{N}$. Si $z = a + bw$ est un élément inversible de l'anneau $\mathcal{A}$, z divise 1 dans $\mathcal{A}$, donc $N(z)$ divise $N(1) = 1$ dans $\mathbb{N}$, ce qui montre que $N(z) = 1$. Réciproquement, si $N(z) = 1$, $z\bar{z} = 1$ donc, puisque $\bar{z} \in \mathcal{A}$, z est

inversible dans $\mathcal{A}$ et son inverse dans $\mathcal{A}$ est $\bar{z}$. Si $z = a + bw$ est un élément de $\mathcal{A}$ et si $N(z) = 1$, on déduit de (E) que $b = 0$ et de l'égalité $1 = a^2 + ab + 5b^2$ que $a^2 = 1$, ce qui montre que $z = \pm 1$. De plus $N(\pm 1) = 1$, donc l'ensemble des éléments inversibles de l'anneau $\mathcal{A}$ est $U(\mathcal{A}) = \{z \in \mathcal{A} \mid N(z) = 1\} = \{1, -1\}$.

Supposons que $\mathcal{A}$ soit un anneau euclidien. Nous disposons d'une application φ de $\mathcal{A} \setminus \{0\}$ dans $\mathbb{N}$ qui vérifie les propriétés (1) et (2) de la définition 3.27 (page 47).

Nous posons $X = \{\varphi(z) \mid z \in \mathcal{A} \setminus \{0, 1, -1\}\}$. L'entier naturel $\varphi(w)$ appartenant à X , X est une partie non vide de $\mathbb{N}$, donc X possède un plus petit élément m .

Nous choisissons un élément u de $\mathcal{A} \setminus \{0, 1, -1\}$ tel que $m = \varphi(u)$. Comme $u \neq 0$, $N(u) > 0$, et comme u est différent de 1 et de -1 , u n'est pas inversible dans $\mathcal{A}$ donc $N(u) \neq 1$, ce qui prouve que $N(u) \geq 2$.

Soit z un élément de $\mathcal{A} \setminus \{0, 1, -1\}$. Il existe $q, r \in \mathcal{A}$ tels que $z = uq + r$ avec $r = 0$ ou $\varphi(r) < \varphi(u)$. Compte tenu de la définition de u , r n'appartient pas à X , donc $r = 0, 1$ ou -1 , d'où les trois possibilités : $z = uq$, $z = uq + 1$ et $z = uq - 1$. Par suite, u divise z , $z - 1$ ou $z + 1$ dans $\mathcal{A}$; nous appelons $P(z)$ cette propriété.

Nous posons $u = a + bw$ où $a, b \in \mathbb{Z}$.

Appliquons $P(z)$ à $z = 2$: u divise 2, 1 ou 3 dans $\mathcal{A}$. Si u divisait 1 dans $\mathcal{A}$, u serait inversible dans $\mathcal{A}$, ce qui est faux; par suite u divise 2 ou 3 dans $\mathcal{A}$.

Supposons que u divise 2. Alors $N(u)$ divise $N(2) = 4$ dans $\mathbb{N}$ et $N(u) \neq 1$, donc $N(u) = 2$ ou 4. Si $b \neq 0$, (E) donne : $N(u) \geq 19/4 > 4$, ce qui est faux. Par suite $b = 0$ donc $u = a$, ce qui montre que $N(u) = a^2$, et comme 2 n'est pas un carré dans $\mathbb{N}$, $N(u) \neq 2$ donc $N(u) = 4$.

Supposons que u divise 3. Alors $N(u)$ divise $N(3) = 9$ dans $\mathbb{N}$ et $N(u) \neq 1$, donc $N(u) = 3$ ou 9. L'égalité (E) donne : $(19b^2)/4 \leq 9$, donc $b^2 \leq 36/19 < 2$. Par suite $b = 0$ ou $b = \pm 1$. Si $b = \pm 1$, $N(u) \geq 19/4 > 4$, donc $N(u) = 9$; en multipliant par 4, il vient : $(2a \pm 1)^2 + 19 = 36$, donc $17 = (2a \pm 1)^2$, en contradiction avec le fait que 17 n'est pas un carré dans $\mathbb{N}$. Ainsi $b = 0$, donc $u = a$, ce qui montre que $N(u) = a^2$, et comme 3 n'est pas un carré dans $\mathbb{N}$, $N(u) \neq 3$ donc $N(u) = 9$.

Il résulte de cette étude que $N(u) = 4$ ou 9. Nous savons que $w \in \mathcal{A} \setminus \{0, 1, -1\}$, donc, en appliquant $P(z)$ à $z = w$, on obtient que u divise w , $w - 1$ ou $w + 1$ dans $\mathcal{A}$, donc que $N(u)$ divise $N(w) = 5$, $N(w - 1) = 5$ ou $N(w + 1) = 7$ dans $\mathbb{N}$, en contradiction avec $N(u) = 4$ ou 9.

En conclusion, $\mathcal{A}$ n'est pas un anneau euclidien.

Il reste à prouver que $\mathcal{A}$ est un anneau principal.

Soit $\mathcal{I}$ un idéal de $\mathcal{A}$ différent de $\{0\}$. L'ensemble $T = \{N(t) \mid t \in \mathcal{I} \setminus \{0\}\}$ est une partie non vide de $\mathbb{N}^*$, donc T admet un plus petit élément k , et $k \in \mathbb{N}^*$. Nous choisissons un élément y de $\mathcal{I} \setminus \{0\}$ tel que $k = N(y)$.

Soit x un élément non nul de $\mathcal{I}$. Comme x et $\bar{y}$ appartiennent à $\mathcal{A}$, leur produit appartient à $\mathcal{A}$; de plus $N(y) \in \mathbb{N}^*$, donc, dans le corps commutatif $\mathbb{C}$:

$$z = \frac{x}{y} = \frac{x\bar{y}}{N(y)} = \frac{\alpha + \beta w}{\gamma}$$

où $\alpha, \beta, \gamma \in \mathbb{Z}$ et $\gamma > 0$. Quitte à les diviser par leur pgcd, nous supposons que α , β et γ sont premiers entre eux dans leur ensemble.

Supposons que $\gamma = 2$. Alors α ou β est impair. Posons $u = \alpha + \beta\bar{w}$; on a $z = \bar{u}/2$. On prouve, en examinant les parités possibles de α et β , que l'entier naturel

$N(u) = \alpha^2 + \alpha\beta + 5\beta^2$ est impair, d'où l'existence de $v \in \mathbb{N}$ tel que $N(u) = 2v + 1$. Alors :

$$zu - v = \frac{1}{2}N(u) - v = \frac{1}{2}.$$

Nous supposons que $\gamma \geq 3$. Le théorème de Bézout fournit $m, n, q \in \mathbb{Z}$ tels que $n\alpha + (m+n)\beta - q\gamma = 1$. Notons p l'entier le plus proche — ou l'un des entiers les plus proches — de $(m\alpha - 5n\beta)/\gamma$ et posons $u = m + nw$ et $v = p + qw$. On a :

$$\begin{aligned} zu - v &= \frac{\alpha + \beta w}{\gamma}(m + nw) - (p + qw) \\ &= \frac{m\alpha - 5n\beta}{\gamma} - p + \frac{n\alpha + (m+n)\beta - q\gamma}{\gamma}w = \frac{m\alpha - 5n\beta}{\gamma} - p + \frac{w}{\gamma} = r + \frac{w}{\gamma} \end{aligned}$$

où r est un nombre rationnel tel que $|r| \leq 1/2$, donc :

$$\begin{aligned} N(zu - v) &= \left| r + \frac{w}{\gamma} \right|^2 = \left(r + \frac{w}{\gamma} \right) \left(r + \frac{\bar{w}}{\gamma} \right) = r^2 + r \frac{w + \bar{w}}{\gamma} + \frac{w\bar{w}}{\gamma^2} \\ &= r^2 + \frac{r}{\gamma} + \frac{5}{\gamma^2} \leq \frac{1}{4} + \frac{1}{6} + \frac{5}{9} = \frac{35}{36}. \end{aligned}$$

De plus $zu - v \neq 0$, car sinon $w = -r\gamma$ donc w est réel, ce qui est évidemment faux. Dans les deux cas ($\gamma = 2$ et $\gamma \geq 3$) on a trouvé des éléments u et v de $\mathcal{A}$ tels que $zu - v \neq 0$ et $N(zu - v) < 1$. Or $ux - vy = y(zu - v)$, donc $ux - vy \neq 0$ et $N(ux - vy) = N(y)N(zu - v) < 1$. Comme x et y appartiennent à l'idéal $\mathcal{I}$ et u et v à $\mathcal{A}$, $ux - vy$ appartient à $\mathcal{I} \setminus \{0\}$, donc l'inégalité $N(ux - vy) < N(y)$ contredit la définition de y . Par conséquent il est impossible que γ soit supérieur ou égal à 2, donc $\gamma = 1$, d'où l'on déduit que $z = \alpha + \beta w$, ce qui montre que $z \in \mathcal{A}$, et comme $x = yz$, x appartient à l'idéal (y) de $\mathcal{I}$ engendré par y ; de plus $y \in \mathcal{I}$, donc $\mathcal{I} = (y)$. En conclusion, $\mathcal{A}$ est un anneau principal. $\square$

■ Autres types d'anneaux

DÉFINITION 3.28. — Un anneau commutatif non nul A est noethérien¹⁵ si toute suite croissante pour l'inclusion d'idéaux de A est stationnaire — ou encore, ce qui est équivalent, si tout idéal de A est engendré par une partie finie de A .

En particulier, un anneau principal est un anneau noethérien.

THÉORÈME 3.8. — **Théorème de Hilbert**¹⁶.

Si A un anneau commutatif non nul et si A est noethérien, l'anneau $A[X]$ des polynômes à coefficients dans A est noethérien.

3.31. Anneau noethérien qui n'est pas un anneau principal.

Nous avons démontré dans l'exemple 3.29 que l'anneau $\mathbb{Z}[X]$ n'est pas un anneau principal. Par ailleurs, $\mathbb{Z}$ est un anneau principal donc un anneau noethérien. Le théorème de Hilbert ci-dessus montre que $\mathbb{Z}[X]$ est un anneau noethérien. $\square$

15. Du nom de la mathématicienne allemande Emmy Nöther (1882-1935), à qui l'on doit des travaux sur les anneaux et les idéaux.

16. Pour une démonstration, voir [PERR], chapitre II, théorème 2.3.

3.32. Anneau commutatif qui n'est pas un anneau noethérien.

Nous munissons l'ensemble $\mathcal{A}$ des suites de réels de l'addition et de la multiplication terme à terme des coefficients ; alors $\mathcal{A}$ est un anneau commutatif non nul dont le zéro est la suite constante égale à 0 et l'élément unité la suite constante égale à 1. Nous notons, pour tout entier naturel n , $\mathcal{I}_n$ l'ensemble des suites nulles à partir du rang n , ce qui signifie qu'un élément $x = (x_k)_{k \in \mathbb{N}}$ de $\mathcal{A}$ appartient à $\mathcal{I}_n$ si, et seulement si, $x_k = 0$ pour tout entier $k \geq n$. Clairement $\mathcal{I}_n$ est, pour tout entier naturel n , un idéal de l'anneau $\mathcal{A}$ et la suite $(\mathcal{I}_n)_{n \in \mathbb{N}}$ est strictement croissante pour l'inclusion. Il en résulte que $\mathcal{A}$ n'est pas un anneau noethérien. $\square$

Le théorème de Krull¹⁷ affirme que dans un anneau, tout idéal à gauche (resp. à droite) strict est inclus dans au moins un idéal à gauche (resp. à droite) maximal. Si un idéal d'un anneau A contient un élément inversible, il contient 1 donc il est égal à A . Il en résulte que si A est un anneau commutatif non nul et si l'ensemble N des éléments non inversibles de A est un idéal, N est un idéal maximal de A et c'est le seul. De tels anneaux existent et s'appellent des anneaux locaux.

3.33. Anneau commutatif non nul n'admettant qu'un seul idéal maximal.

Soit p un nombre premier et k un entier ≥ 1 . Posons $q = p^k$ et considérons l'anneau $\mathcal{A} = \mathbb{Z}/q\mathbb{Z}$ des entiers modulo q . Si $n \in \mathbb{Z}$, la classe $\bar{n}$ de n modulo q est inversible dans $\mathcal{A}$ si, et seulement si, n est premier avec q , donc, p étant un nombre premier, $\bar{n}$ est inversible dans l'anneau $\mathcal{A}$ si, et seulement si, p ne divise pas n dans $\mathbb{Z}$.

Nous notons $\mathcal{N}$ l'ensemble des éléments de $\mathcal{A}$ qui ne sont pas inversibles dans $\mathcal{A}$. On a donc $\mathcal{N} = \{\bar{n} \mid n \in \mathbb{Z} \text{ et } p \text{ divise } n\}$. Comme p divise 0, le zéro $\bar{0}$ de $\mathcal{A}$ appartient à $\mathcal{N}$. Soit x et y des éléments de $\mathcal{N}$ et z un élément de $\mathcal{A}$. Alors $x = \bar{\ell}$, $y = \bar{m}$ et $z = \bar{n}$ où ℓ , m et n sont des entiers relatifs tels que p divise ℓ et m . L'entier p divise $\ell + m$ et $n\ell$, donc $x + y = \bar{\ell} + \bar{m}$ et $zx = \bar{n}\bar{\ell}$ appartiennent à $\mathcal{N}$. En conclusion $\mathcal{N}$ est un idéal de $\mathcal{A}$, donc — voir ci-dessus — $\mathcal{N}$ est un idéal maximal de $\mathcal{A}$ et c'est le seul. $\square$

L'anneau de l'exemple précédent 3.33 n'est pas un anneau intègre — sauf si $k = 1$, auquel cas c'est un corps.

3.34. Anneau intègre ne comportant qu'un seul idéal maximal.

Nous considérons l'anneau $\mathcal{A} = \mathbb{R}[[X]]$ des séries formelles¹⁸ à une indéterminée à coefficients réels ; c'est un anneau intègre. Les éléments inversibles de $\mathcal{A}$ sont les séries formelles $\sum_{n \in \mathbb{N}} a_n X^n$ de « terme constant » a_0 différent de zéro. Les éléments de $\mathcal{A}$ qui ne sont pas inversibles dans $\mathcal{A}$ sont donc les séries formelles $\sum_{n \in \mathbb{N}} a_n X^n$ dont le terme constant a_0 est nul. On en déduit que l'ensemble $\mathcal{N}$ des éléments de $\mathcal{A}$ qui ne sont pas inversibles dans $\mathcal{A}$ est l'idéal de $\mathcal{A}$ engendré par X . En particulier $\mathcal{N}$ est un idéal de $\mathcal{A}$, et on conclut comme dans l'exemple précédent 3.33 que $\mathcal{N}$ est un idéal maximal de $\mathcal{A}$ et que c'est le seul. $\square$

17. Du nom du mathématicien allemand Wolfgang Krull (1899-1970). Ce théorème découle de l'axiome du choix par l'intermédiaire du lemme de Zorn — voir le chapitre 1, pages 14 et 15.

18. Voir [ARN1], §VIII.4 ou [SCHW], chapitre 3, §5.

3.35. Autre anneau commutatif non nul n'ayant qu'un seul idéal maximal.

Nous introduisons l'ensemble $\mathcal{H}$ des couples (f, a) où a est un réel strictement positif et f une fonction de $\mathbb{R}$ dans $\mathbb{R}$ définie et continue sur l'intervalle $[0, a[$. En notant $\mathbf{0}$ l'application de $\mathbb{R}_+$ dans $\mathbb{R}$ constante égale à 0 et $\mathbf{1}$ l'application de $\mathbb{R}_+$ dans $\mathbb{R}$ constante égale à 1, $\omega = (\mathbf{0}, 1)$ et $\varepsilon = (\mathbf{1}, 1)$ sont des éléments de $\mathcal{H}$. Nous posons, si (f, a) et (g, b) appartiennent à $\mathcal{H}$:

$$(f, a) \oplus (g, b) = (f + g, \text{Min}(a, b)) \text{ et } (f, a) \otimes (g, b) = (f \times g, \text{Min}(a, b)).$$

Il est clair que nous obtenons ainsi sur $\mathcal{H}$ des lois de composition interne $\oplus$ et $\otimes$ commutatives et associatives. La relation binaire $\mathcal{R}$, définie sur l'ensemble $\mathcal{H}$ par : $(f, a) \mathcal{R} (g, b)$ s'il existe un réel c tel que $0 < c \leq \text{Min}(a, b)$ et $f(t) = g(t)$ pour tout $t \in [0, c[$, est clairement une relation d'équivalence sur $\mathcal{H}$ compatible avec $\oplus$ et $\otimes$.

Nous notons $\mathcal{A} = \mathcal{H}/\mathcal{R}$ l'ensemble quotient¹⁹ de $\mathcal{H}$ par $\mathcal{R}$ et, pour tout élément φ de $\mathcal{A}$, $\overline{\varphi}$ la classe d'équivalence de φ modulo $\mathcal{R}$. Si $x \in \mathcal{A}$, un représentant de x est un élément $\varphi = (f, a)$ de $\mathcal{H}$ tel que $x = \overline{\varphi}$. La compatibilité de $\mathcal{R}$ avec $\oplus$ et $\otimes$ justifie l'introduction des lois de composition interne $+$ et $\times$ définies sur $\mathcal{A}$ de la manière suivante : si x et y appartiennent à $\mathcal{A}$, $x + y = \overline{\varphi \oplus \psi}$ et $x \times y = \overline{\varphi \otimes \psi}$ où φ et ψ sont des représentants respectifs (quelconques) de x et de y . On vérifie sans difficulté que l'ensemble $\mathcal{A}$, muni de l'addition $+$ et de la multiplication $\times$ ainsi définies, est un anneau commutatif non nul dont le zéro est $0_{\mathcal{A}} = \overline{\omega}$ et dont l'élément unité est $1_{\mathcal{A}} = \overline{\varepsilon}$.

Soit u un élément inversible de l'anneau $\mathcal{A}$. Soit $\varphi = (f, a)$ un représentant de u . Il existe $v \in \mathcal{A}$ tel que $u \times v = 1_{\mathcal{A}}$. Choisissons un représentant $\psi = (g, b)$ de v . Comme $(f \times g, \text{Min}(a, b)) \mathcal{R} \varepsilon$, il existe un réel c tel que $0 < c \leq \text{Min}(a, b, 1)$ et $f(t)g(t) = \mathbf{1}(t) = 1$ pour tout $t \in [0, c[$; en particulier $f(0)g(0) = 1$ donc $f(0) \neq 0$. Réciproquement, soit u un élément de $\mathcal{A}$ admettant un représentant $\varphi = (f, a)$ tel que $f(0) \neq 0$. La continuité de f à droite en 0 justifie l'existence de $b \in]0, a]$ tel que $f(t) \neq 0$ pour tout point t de $[0, b[$. La fonction inverse $1/f$ est alors définie et continue sur $[0, b[$, donc $\psi = (1/f, b)$ appartient à $\mathcal{H}$; posons $v = \overline{\psi}$. Clairement, $(\varphi \otimes \psi) \mathcal{R} \varepsilon$ donc $u \times v = 1_{\mathcal{A}}$, ce qui montre que u est inversible dans $\mathcal{A}$.

Il résulte de cette étude qu'un élément u de $\mathcal{A}$ est inversible dans l'anneau $\mathcal{A}$ si, et seulement si, il admet un représentant (f, a) tel que $f(0) \neq 0$.

Si $\varphi = (f, a)$ et $\psi = (g, b)$ appartiennent à $\mathcal{H}$ et si $\varphi \mathcal{R} \psi$, alors $f(0) = g(0)$, ce qui justifie l'existence de l'application $T : \mathcal{A} \rightarrow \mathbb{R}$ définie ainsi : si $u \in \mathcal{A}$, $T(u) = f(0)$ où (f, a) est un représentant (quelconque) de u . Des calculs immédiats montrent que T est un morphisme d'anneau de $\mathcal{A}$ dans $\mathbb{R}$, donc $\mathcal{N} = \text{Ker}(T)$ est un idéal de $\mathcal{A}$. On déduit de la description ci-dessus des éléments inversibles de $\mathcal{A}$ que $\mathcal{N}$ est l'ensemble des éléments de $\mathcal{A}$ qui ne sont pas inversibles dans $\mathcal{A}$, donc, comme dans les exemples précédents 3.33 et 3.34, $\mathcal{N}$ est l'unique idéal maximal de $\mathcal{A}$. $\square$

La divisibilité sur un anneau intègre est une relation de préordre, mais ce n'est pas en général un préordre total²⁰; c'est le cas de $\mathbb{Z}$, puisque 2 ne divise pas 3 et 3

19. Un élément de $\mathcal{A}$ s'appelle un germe de fonctions à droite de 0.

20. Une relation de préordre sur un ensemble E est une relation binaire $\mathcal{R}$ réflexive et transitive sur E , et $\mathcal{R}$ est un préordre total si, quels que soient $a, b \in E$, $a \mathcal{R} b$ ou $b \mathcal{R} a$.

ne divise pas 2. Dans un corps commutatif, tout élément non nul divise n'importe quel autre : on dit dans ce cas que la relation de divisibilité est triviale. Donnons un exemple dans lequel la relation de divisibilité est totale sans être triviale.

3.36. Anneau intègre dans lequel, pour tout couple (a, b) , a divise b ou b divise a .

Nous reprenons l'anneau intègre $\mathcal{A} = \mathbb{R}[[X]]$ des séries formelles à une indéterminée à coefficients réels de l'exemple 3.34 (page 52). Les éléments inversibles de $\mathcal{A}$ sont les séries formelles $\sum_{n \in \mathbb{N}} a_n X^n$ de « terme constant » a_0 différent de 0. Tout élément de $\mathcal{A}$ divise 0. Soit donc $S = \sum_{n \in \mathbb{N}} a_n X^n$ et $T = \sum_{n \in \mathbb{N}} b_n X^n$ des éléments non nuls de $\mathcal{A}$. Il existe $p, q \in \mathbb{N}$ tels que $a_p \neq 0$ et $b_q \neq 0$. Notons k le plus petit des entiers naturels p tel que $a_p \neq 0$ et ℓ le plus petit des entiers naturels q tel que $b_q \neq 0$. On a $S = X^k P$ où $P = \sum_{n \in \mathbb{N}} a_{k+n} X^n$; or $a_{k+0} = a_k \neq 0$, donc P est un élément inversible de $\mathcal{A}$. De même, $T = X^\ell Q$ où Q est un élément inversible de $\mathcal{A}$. Si $k \leq \ell$, $T = X^{\ell-k} X^k Q = X^{\ell-k} (P^{-1} S) Q = (X^{\ell-k} P^{-1} Q) S$, donc S divise T ; un raisonnement identique montre que si $\ell \leq k$, alors T divise S . Or $k \leq \ell$ ou $\ell \leq k$, donc S divise T ou T divise S . $\square$

3.37. Anneau intègre qui n'est pas un anneau noethérien.

Nous notons $\mathcal{A}_0$ le sous-ensemble de l'anneau $\mathcal{A}$ de l'exemple 3.35 constitué des classes d'équivalence modulo $\mathcal{R}$ des éléments $\varphi = (f, a)$ de $\mathcal{H}$ pour lesquels il existe un nombre rationnel $r > 0$ et une fonction f_0 , somme d'une série entière de rayon de convergence strictement positif, tels que $f(t) = f_0(t^r)$ pour tout $t \in [0, a]$.

Clairement, $1_{\mathcal{A}} \in \mathcal{A}_0$ et, pour tout élément x de $\mathcal{A}_0$, $-x \in \mathcal{A}_0$. Soit $x, y \in \mathcal{A}_0$. On a $x = \overline{\varphi}$ et $y = \overline{\psi}$ avec $\varphi = (f, a)$ et $\psi = (g, b)$ où $f : t \mapsto f(t) = f_0(t^r)$ et $g : t \mapsto g(t) = g_0(t^s)$ pour des rationnels $r > 0$ et $s > 0$ et des sommes f_0 et g_0 de séries entières de rayon de convergence strictement positif. En réduisant r et s au même dénominateur, on obtient des représentants (p, d) de r et (q, d) de s où $p, q, d \in \mathbb{N}^*$. Les fonctions $f_1 : t \mapsto f_1(t) = f_0(t^p)$ et $g_1 : t \mapsto g_1(t) = g_0(t^q)$ sont des sommes de séries entières de rayon de convergence strictement positif, $1/d$ est un rationnel strictement positif, $f(t) = f_1(t^{1/d})$ pour tout $t \in [0, a]$ et $g(t) = g_1(t^{1/d})$ pour tout $t \in [0, b]$. Alors $k_1 = f_1 + g_1$ et $\ell_1 = f_1 g_1$ sont des sommes de séries entières de rayon de convergence strictement positif, et $(f + g)(t) = k_1(t^{1/d})$ et $(fg)(t) = \ell_1(t^{1/d})$ pour tout $t \in [0, \min(a, b)]$, donc $x + y$ et xy appartiennent à $\mathcal{A}_0$. En conclusion, $\mathcal{A}_0$ est un sous-anneau de l'anneau $\mathcal{A}$. Nous munissons $\mathcal{A}_0$ de l'addition et de la multiplication induites par celle de $\mathcal{A}$; alors $\mathcal{A}_0$ est un anneau commutatif non nul de zéro $0_{\mathcal{A}}$ et d'élément unité $1_{\mathcal{A}}$.

Soit $x, y \in \mathcal{A}_0$ tels que $x \neq 0_{\mathcal{A}}$ et $x \times y = 0_{\mathcal{A}}$. Nous utilisons les mêmes notations $\varphi = (f, a)$, $\psi = (g, b)$, r, s, f_0 et g_0 que ci-dessus, avec les mêmes hypothèses. Comme $(f \times g, \min(a, b)) \mathcal{R} 0$, il existe un réel c tel que $0 < c \leq \min(a, b, 1)$ et $f(t)g(t) = 0(t) = 0$ pour tout $t \in [0, c]$. Si $f(0) \neq 0$, alors, par continuité de f à droite en 0, il existe $d \in]0, c]$ tel que $f(t) \neq 0$ pour tout $t \in [0, d]$, donc $g(t) = 0$ pour tout $t \in [0, c]$. Si $f(0) = 0$, alors $f_0(0) = 0$ et f_0 n'est pas la fonction nulle donc, f_0 étant analytique, 0 est un zéro isolé de f_0 , et comme $f_0(t^r)g_0(t^s) = 0$ pour tout $t \in [0, c]$, il existe un réel $d_0 > 0$ tel que $g_0(t) = 0$ pour tout $t \in]0, d_0]$, donc pour tout $t \in [0, d_0]$, donc un réel $d > 0$ tel que $g(t) = 0$ pour tout $t \in [0, d]$. Ainsi, dans les deux cas, $y = 0_{\mathcal{A}}$. En conclusion, $\mathcal{A}_0$ est anneau intègre.

Nous introduisons, pour tout nombre rationnel $r > 0$, l'application :

$$\left| \begin{array}{l} h_r : \mathbb{R}_+ \longrightarrow \mathbb{R} \\ t \longmapsto h_r(t) = t^r \end{array} \right.$$

et la classe d'équivalence z_r de $(h_r, 1)$ modulo $\mathcal{R}$; en définissant l'application $\ell_0 : t \mapsto \ell_0(t) = t$ de $\mathbb{R}$ dans $\mathbb{R}$, alors $h_r(t) = \ell_0(t^r)$ pour tout $t \in [0, 1]$, donc $z_r \in \mathcal{A}_0$. Nous notons, pour tout rationnel $r > 0$, $\mathcal{I}_r$ l'idéal de $\mathcal{A}_0$ engendré par z_r . Soit $r, s \in \mathbb{Q}$ tels que $0 < r < s$. Pour tout $t \in [0, 1]$, $h_s(t) = h_{s-r}(t)h_r(t)$, donc $z_s = z_{s-r} \times z_r$, ce qui montre que z_r divise z_s dans l'anneau $\mathcal{A}_0$. Par suite $\mathcal{I}_s \subset \mathcal{I}_r$, et cette inclusion est stricte, puisque si f est une fonction définie et continue sur un intervalle $[0, a]$ où $a > 0$ et si $t^r = t^s f(t)$ pour tout $t \in]0, a[$, alors $f(t)$ tend vers $+\infty$ quand t tend vers 0 à droite, en contradiction avec la continuité de f à droite en 0. Par conséquent $(\mathcal{I}_{1/n})_{n \in \mathbb{N}^*}$ est une suite strictement croissante pour l'inclusion d'idéaux de $\mathcal{A}_0$, donc $\mathcal{A}_0$ n'est pas un anneau noethérien. $\square$

3.38. Anneau intègre sans élément irréductible.

Nous reprenons l'anneau intègre $\mathcal{A}_0$ de l'exemple précédent 3.37.

Soit z un élément non nul et non inversible de $\mathcal{A}_0$. On a $z = \overline{\varphi}$ où $\varphi = (f, a)$ est un élément de $\mathcal{H}$ pour lequel il existe un rationnel $r > 0$ et une fonction f_0 , somme d'une série entière $\sum_{n \geq 0} \alpha_n t^n$ de rayon de convergence $R > 0$, tels que $f(t) = f_0(t^r)$ pour tout $t \in [0, a]$.

Si $\alpha_0 \neq 0$, la série formelle $\sum_{n \in \mathbb{N}} \alpha_n X^n$ est inversible dans l'anneau $\mathbb{R}[[X]]$ — voir les exemples 3.34 et 3.36 — et si l'on note $\sum_{n \in \mathbb{N}} \beta_n X^n$ son inverse, la série entière $\sum_{n \geq 0} \beta_n t^n$ admet un rayon de convergence strictement positif et il existe un réel $b > 0$ tel que $1/f_0(t) = \sum_{n=0}^{+\infty} \beta_n t^n$ pour tout $t \in [0, b]$, d'où l'on déduit, comme dans l'exemple 3.35 pour $\mathcal{A}$, que z est inversible dans l'anneau $\mathcal{A}_0$, en contradiction avec l'hypothèse. Par conséquent, $\alpha_0 = 0$. De plus z n'est pas nul, donc il existe au moins un entier naturel k tel que $\alpha_k \neq 0$. Notons m le plus petit des entiers naturels k tels que $\alpha_k \neq 0$. Alors $m \geq 1$, $\alpha_m \neq 0$, $\alpha_k = 0$ pour $k < m$ et :

$$f_0(t) \underset{t \rightarrow 0^+}{\sim} \alpha_m t^m.$$

La fonction $f_1 : t \mapsto f_1(t) = f_0(t^2)$ est la somme de la série entière $\sum_{n \geq 0} \alpha_n t^{2n}$, de rayon de convergence $R_1 \geq \sqrt{R} > 0$, et $f(t) = f_0(t^r) = f_1(t^{r/2})$ pour tout $t \in [0, a]$. De plus :

$$f_1(t) \underset{t \rightarrow 0^+}{\sim} \alpha_m t^{2m} \text{ donc } \frac{f_1(t)}{t} \underset{t \rightarrow 0^+}{\sim} \alpha_m t^{2m-1}.$$

Or $2m - 1 \geq 1$, donc la fonction :

$$g_1 : t \mapsto g_1(t) = \frac{f_1(t)}{t},$$

prolongée en 0 par $g_1(0) = 0$, est la somme de la série entière $\sum_{n \geq 0} \alpha_n t^{2n-1}$, de rayon de convergence $R_1 > 0$, et $f_1(t) = t g_1(t)$ pour tout point t de $[0, \sqrt{a}]$.

Nous introduisons enfin les fonctions :

$$g : t \mapsto g(t) = g_1(t^{r/2}), \quad \ell_1 : t \mapsto \ell_1(t) = t \text{ et } \ell : t \mapsto \ell(t) = \ell_1(t^{r/2}) = t^{r/2}.$$

Notons x et y les classes d'équivalence respectives de (ℓ, a) et (g, a) modulo $\mathcal{R}$. Alors x et y appartiennent à $\mathcal{A}_0$. Or $f(t) = f_1(t^{r/2}) = t^{r/2} g_1(t^{r/2}) = \ell(t) g(t)$ pour tout $t \in [0, a]$, donc $z = x \times y$. Si x ou y était inversible dans $\mathcal{A}_0$, il serait inversible

dans $\mathcal{A}$, donc on aurait $\ell(0) \neq 0$ ou $g(0) \neq 0$, en contradiction avec $\ell(0) = g(0) = 0$. Par suite ni x ni y n'est inversible dans $\mathcal{A}_0$, donc z est réductible dans $\mathcal{A}_0$. Ainsi $\mathcal{A}_0$ est un anneau intègre qui ne possède aucun élément irréductible. $\square$

■ Corps

Tous les corps finis sont commutatifs d'après le théorème de Wedderburn (théorème 3.1, page 34) ; ceci devient faux pour les corps infinis.

3.39. Corps non commutatif²¹.

Nous posons $\mathbb{H} = \left\{ \begin{pmatrix} u & -\bar{v} \\ v & \bar{u} \end{pmatrix} \mid u, v \in \mathbb{C} \right\}$.

C'est un sous-ensemble de l'anneau $\mathcal{M}_2(\mathbb{C})$ des matrices carrées d'ordre 2 à coefficients complexes. Clairement, $\mathbb{H}$ est un sous-groupe du groupe additif de $\mathcal{M}_2(\mathbb{C})$ et la matrice unité I_2 appartient à $\mathbb{H}$.

Si $A = \begin{pmatrix} a & -\bar{b} \\ b & \bar{a} \end{pmatrix}$ et $B = \begin{pmatrix} c & -\bar{d} \\ d & \bar{c} \end{pmatrix}$ sont des éléments de $\mathbb{H}$, on a :

$$AB = \begin{pmatrix} ac - \bar{b}d & -a\bar{d} + \bar{c}\bar{b} \\ cb + \bar{a}d & -b\bar{d} + \bar{a}\bar{c} \end{pmatrix} = \begin{pmatrix} u & -\bar{v} \\ v & \bar{u} \end{pmatrix}$$

où $u = ac - \bar{b}d$ et $v = cb + \bar{a}d$, donc AB appartient à $\mathbb{H}$.

Il en résulte que $\mathbb{H}$ est un sous-anneau de l'anneau $\mathcal{M}_2(\mathbb{C})$. Si :

$$A = \begin{pmatrix} a & -\bar{b} \\ b & \bar{a} \end{pmatrix}$$

est un élément non nul de $\mathbb{H}$, alors $a \neq 0$ ou $b \neq 0$, donc $\delta = \det A = |a|^2 + |b|^2 > 0$, ce qui montre que A est inversible dans $\mathcal{M}_2(\mathbb{C})$ et que, puisque $1/\delta$ est un réel, son inverse :

$$A^{-1} = \frac{1}{\delta} \begin{pmatrix} \bar{a} & \bar{b} \\ -b & a \end{pmatrix}$$

appartient à $\mathbb{H}$, d'où l'on déduit que A est inversible dans l'anneau $\mathbb{H}$. Il en résulte que l'anneau $\mathbb{H}$ est un corps. Or les matrices :

$$P = \begin{pmatrix} i & 0 \\ 0 & -i \end{pmatrix} \text{ et } Q = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$$

appartiennent à $\mathbb{H}$ et :

$$PQ = \begin{pmatrix} 0 & -i \\ -i & 0 \end{pmatrix} \neq \begin{pmatrix} 0 & i \\ i & 0 \end{pmatrix} = QP,$$

donc $\mathbb{H}$ n'est pas un corps commutatif. $\square$

Le corps non commutatif $\mathbb{H}$ construit dans l'exemple précédent 3.39 est appelé le corps des quaternions d'Hamilton ; sa multiplication permet de décrire les rotations d'un espace vectoriel euclidien de dimension 3, de même que la multiplication de $\mathbb{C}$ décrit celles d'un espace vectoriel euclidien de dimension 2.

21. On en trouvera un autre exemple dans [BLAN] (chapitre I, §4), le corps non commutatif obtenu étant un espace vectoriel de dimension finie 9 sur le corps $\mathbb{Q}$ des nombres rationnels.

DÉFINITION 3.29. — Un corps premier est un corps commutatif n'ayant aucun sous-corps autre que lui-même.

THÉORÈME 3.9. — Les corps premiers sont, à un isomorphisme de corps près, $\mathbb{Q}$ et les $\mathbb{F}_p = \mathbb{Z}/p\mathbb{Z}$ où p est un nombre premier.

THÉORÈME 3.10. — Le seul automorphisme de corps d'un corps premier est l'identité.

3.40. Corps commutatif qui n'a pas d'autre automorphisme que l'identité, mais qui n'est pas un corps premier.

Il suffit de considérer le corps $\mathbb{R}$ des nombres réels. Tout d'abord, $\mathbb{Q}$ est un sous-corps de $\mathbb{R}$ différent de $\mathbb{R}$, donc $\mathbb{R}$ n'est pas un corps premier.

Soit f un automorphisme du corps $\mathbb{R}$. En particulier, f est un endomorphisme du groupe additif de $\mathbb{R}$, $f(1) = 1$ et, quels que soient les réels x et y , $f(xy) = f(x)f(y)$. Si t est un réel et si $t \geq 0$, alors $f(t) = f(\sqrt{t}\sqrt{t}) = (f(\sqrt{t}))^2 \geq 0$. Si x et y sont des nombres réels et si $x \leq y$, alors $y - x \geq 0$ donc $f(y) - f(x) = f(y - x) \geq 0$. L'application f est donc croissante. En appliquant le théorème 5.3 — le texte et la preuve de ce théorème figurent dans le chapitre 5, à la page 95 —, nous concluons qu'il existe un nombre réel α tel que $f = \alpha \text{Id}_{\mathbb{R}}$; or $1 = f(1) = \alpha$, donc $f = \text{Id}_{\mathbb{R}}$. $\square$

■ Corps ordonnés

DÉFINITION 3.30. — Un corps commutatif totalement ordonné est un corps commutatif K dans lequel il existe une partie P , dont les éléments sont appelés les éléments positifs de K , vérifiant les quatre axiomes suivants :

- (i) $P + P \subset P$. (ii) $P \times P \subset P$.
- (iii) $P \cup (-P) = K$. (iv) $P \cap (-P) = \{0\}$.

Pour alléger le vocabulaire, nous appelons simplement dans ce qui suit «corps ordonné» tout corps commutatif totalement ordonné, et nous disons que le corps commutatif K est ordonné par P si P est une partie de K vérifiant les axiomes (i), (ii), (iii) et (iv) de la définition 3.30.

THÉORÈME 3.11. — Si K est un corps ordonné par P , la relation binaire $\leq$, définie sur K par : $x \leq y$ si $y - x$ appartient à P , est une relation d'ordre total sur K et, quels que soient les éléments x , y et z de K , $x \leq y$ entraîne $x + z \leq y + z$, et $(x \leq y \text{ et } z \geq 0)$ entraîne $xz \leq yz$.

Remarquons que $\mathbb{R}$ est ordonné par $P = \mathbb{R}_+$. Un corps commutatif K , ordonné par une partie P de K , est systématiquement muni de la relation d'ordre total définie dans le texte du théorème 3.11, qui montre de plus que les rapports entre l'addition, la multiplication et la relation d'ordre total $\leq$ sur un corps ordonné sont les mêmes que dans $\mathbb{R}$.

Si K est un corps ordonné, la valeur absolue et les notions de partie majorée, minorée et bornée, de suite convergente, de limite d'une suite convergente et de suite de Cauchy se définissent comme dans $\mathbb{R}$, en remplaçant $\mathbb{R}$ par K .

DÉFINITION 3.31. — Un corps ordonné K est archimédien si, quels que soient les éléments $a > 0$ et $b \geq 0$ de K , il existe un entier $n \geq 1$ tel que $b < n \cdot a$.

Clairement, un corps ordonné K est archimédien si, et seulement si, la suite $(1/n)$ — c'est-à-dire la suite $(1/(n \cdot 1))$ — converge vers 0 dans K .

DÉFINITION 3.32. — Un corps ordonné K est complet si toute suite de Cauchy de K converge dans K .

Le corps $\mathbb{R}$ des nombres réels est caractérisé par le fait d'être un corps ordonné archimédien et complet, ce qui signifie que tout corps ordonné archimédien et complet est isomorphe à $\mathbb{R}$ — pour les structures de corps et d'ensemble ordonné. Le corps ordonné $\mathbb{Q}$ étant archimédien mais n'étant pas complet, on ne peut pas supprimer *complet* ; on ne peut pas non plus supprimer *archimédien*.

3.41. Corps ordonné qui n'est pas archimédien.

Nous considérons le corps commutatif $\mathcal{F} = \mathbb{R}(X)$ des fractions rationnelles à coefficients réels. Si P et Q sont des polynômes non nuls, nous notons $\gamma(P, Q)$ le quotient du coefficient dominant de P par celui de Q et, si Q est un polynôme non nul, nous posons $\gamma(0, Q) = 0$. Si F est un élément de $\mathcal{F}$, $\gamma(P, Q)$ a clairement la même valeur pour tous les représentants (P, Q) de F ; nous notons $\psi(F)$ cette valeur commune. Par exemple :

$$\psi\left(\frac{-4X+3}{-6X^2+5}\right) = \frac{2}{3}.$$

Nous posons alors $\mathcal{P} = \{F \in \mathcal{F} \mid \psi(F) \geq 0\}$. Clairement, une fraction rationnelle non nulle F appartient à $\mathcal{P}$ si, et seulement si, pour tout représentant (P, Q) de F , les coefficients dominants des polynômes P et Q sont de même signe, ce qui prouve l'existence d'au moins un représentant *positif* de F , c'est-à-dire un représentant (P, Q) de F tel que les coefficients dominants des polynômes P et Q sont strictement positifs — si ce n'est pas le cas, il suffit de remplacer le représentant (P, Q) par le représentant $(-P, -Q)$. Nous démontrons que la partie $\mathcal{P}$ de $\mathcal{F}$ vérifie les quatre axiomes de la définition 3.30 des corps ordonnés.

Soit F_1 et F_2 des éléments non nuls de $\mathcal{P}$. Nous notons (P_1, Q_1) et (P_2, Q_2) des représentants positifs respectifs de F_1 et de F_2 , et a_1, a_2, b_1 et b_2 les coefficients dominants respectifs des polynômes P_1, P_2, Q_1 et Q_2 ; on a donc $a_1 > 0, a_2 > 0, b_1 > 0$ et $b_2 > 0$. Le couple $(A, B) = (P_1Q_2 + P_2Q_1, Q_1Q_2)$ est un représentant de $F_1 + F_2$, le coefficient dominant de A est $a_1b_2 > 0, a_2b_1 > 0$ ou $a_1b_2 + a_2b_1 > 0$, et celui de B est $b_1b_2 > 0$, donc $F_1 + F_2$ appartient à $\mathcal{P}$. Le couple (P_1P_2, Q_1Q_2) est un représentant de F_1F_2 , le coefficient dominant de P_1P_2 est $a_1a_2 > 0$ et celui de Q_1Q_2 est $b_1b_2 > 0$, donc F_1F_2 appartient à $\mathcal{P}$. Si F_1 et F_2 sont des éléments de $\mathcal{P}$ et si $F_1 = 0$ ou $F_2 = 0$, alors $F_1 + F_2 = F_2$ ou F_1 appartient à $\mathcal{P}$ et $F_1F_2 = 0$ également. Par conséquent :

$$\mathcal{P} + \mathcal{P} \subset \mathcal{P} \text{ et } \mathcal{P} \times \mathcal{P} \subset \mathcal{P}.$$

Remarquons que $-\mathcal{P} = \{F \in \mathcal{F} \mid \psi(F) \leq 0\}$. Pour toute fraction rationnelle F , $\psi(F) \geq 0$ ou $\psi(F) \leq 0$, et on a à la fois $\psi(F) \geq 0$ et $\psi(F) \leq 0$ si, et seulement si, $F = 0$, donc $\mathcal{P} \cup (-\mathcal{P}) = \mathcal{F}$ et $\mathcal{P} \cap (-\mathcal{P}) = \{0\}$.

Ainsi $\mathcal{P}$ fait de $\mathcal{F}$ un corps ordonné. Nous prouvons qu'il n'est pas archimédien.

Les fractions rationnelles $F = 1/X$ et $G = 1/1 = 1$ sont des éléments non nuls de $\mathcal{P}$ donc $F > 0$ et $G > 0$ dans le corps ordonné $\mathcal{F}$. On a, pour tout entier naturel n :

$$G - n \cdot F = 1 - \frac{n}{X} = \frac{X - n}{X},$$

donc $\psi(G - n \cdot F) = 1 > 0$, ce qui montre que $G - n \cdot F$ est un élément non nul de $\mathcal{P}$, d'où l'on déduit que $n \cdot F < G$ dans le corps ordonné $\mathcal{F}$. En conclusion, le corps ordonné $\mathcal{F}$ n'est pas archimédien. $\square$

Le corps ordonné $\mathcal{F}$ de l'exemple précédent 3.41 n'est pas complet. On démontre en effet que la suite $(F_n)_{n \geq 1}$ d'éléments de $\mathcal{F}$, de terme général :

$$F_n = \sum_{p=1}^n \frac{1}{pX^p} = \frac{1}{X} + \frac{1}{2X^2} + \cdots + \frac{1}{nX^n},$$

est de Cauchy dans $\mathcal{F}$ mais ne converge pas dans $\mathcal{F}$. On peut, comme pour construire $\mathbb{R}$ à l'aide des suites de Cauchy de $\mathbb{Q}$, plonger $\mathcal{F}$ dans un corps ordonné complet, qui n'est pas archimédien puisqu'il contient $\mathcal{F}$ à un isomorphisme près.

3.42. Corps commutatif que l'on ne peut pas ordonner.

Il découle du théorème 3.11 (page 57) que, dans un corps ordonné, les carrés sont positifs, d'où l'on déduit que $1 > 0$, donc $-1 < 0$. Nous examinons le cas du corps $\mathbb{C}$ des nombres complexes. On a $-1 = i^2$, donc -1 est un carré dans $\mathbb{C}$; si l'on pouvait ordonner $\mathbb{C}$, -1 serait positif dans ce corps ordonné, en contradiction avec $-1 < 0$. Il est donc impossible de faire de $\mathbb{C}$ un corps ordonné. $\square$

3.43. Corps que l'on peut ordonner de plusieurs manières.

Le nombre réel $\sqrt{2}$ est irrationnel²² donc, si x est un nombre réel, si a et b sont des rationnels et si $x = a + b\sqrt{2}$, un tel couple (a, b) est unique, ce qui permet de prouver que $K = \mathbb{Q}(\sqrt{2}) = \{a + b\sqrt{2} \mid a, b \in \mathbb{Q}\}$ est un sous-corps de $\mathbb{R}$ — donc K , muni de l'addition et de la multiplication induites par celles de $\mathbb{R}$, est un corps commutatif —, de justifier l'existence de l'application :

$$\left| \begin{array}{l} \varphi : K \longrightarrow K \\ x = a + b\sqrt{2} \longmapsto \varphi(x) = a - b\sqrt{2} \end{array} \right.$$

et d'établir que φ est un automorphisme du corps commutatif K . L'ordre induit sur K par l'ordre canonique de $\mathbb{R}$ fait de K un corps ordonné par la partie $P_0 = K \cap \mathbb{R}_+$ de K . Nous posons :

$$P = \{a + b\sqrt{2} \mid a, b \in \mathbb{Q} \text{ et } a - b\sqrt{2} \in \mathbb{R}_+\} = \{x \in K \mid \varphi(x) \in P_0\} = \varphi^{-1}(P_0).$$

On a :

- (i) $P + P = \varphi(P_0) + \varphi(P_0) = \varphi(P_0 + P_0) \subset \varphi(P_0) = P$,
- (ii) $P \times P = \varphi(P_0) \times \varphi(P_0) = \varphi(P_0 \times P_0) \subset \varphi(P_0) = P$,
- (iii) $P \cup (-P) = \varphi(P_0) \cup \varphi(-P_0) = \varphi(P_0 \cup (-P_0)) = \varphi(K) = K$,
- (iv) $P \cap (-P) = \varphi(P_0) \cap \varphi(-P_0) = \varphi(P_0 \cap (-P_0)) = \varphi(\{0\}) = \{0\}$,

donc P définit sur K une structure de corps ordonné. Dans K ordonné par P_0 , $1 - \sqrt{2} < 0$, alors que dans K ordonné par P , $1 - \sqrt{2} > 0$, donc la structure de corps ordonné définie sur K par P est différente de celle définie par P_0 . $\square$

²² Voir l'exemple 5.3, page 84.

Chapitre 4

Espaces vectoriels

La notion d'espace vectoriel a été introduite dans les années 1840 par Arthur Cayley (1821-1895) et Hermann Grassmann (1809-1877). Le premier a considéré des n -uplets de réels comme étant chacun un élément et a défini des opérations sur ces objets. Le second a fourni une théorie confuse, mais ne dépendant d'aucune base. Les espaces vectoriels ont été formalisés en 1888 par Giuseppe Peano et sont devenus le cadre naturel de la géométrie, mais aussi celui de nombreuses théories comme l'analyse fonctionnelle. L'algèbre linéaire permet en effet d'utiliser l'intuition géométrique dans les théories mathématiques dépourvues de support intuitif apparent.

DÉFINITION 4.1. — Un espace vectoriel sur un corps commutatif K , dont l'élément unité est noté 1, est un groupe abélien additif $(E, +)$ muni d'une loi de composition externe de domaine K :

$$\begin{array}{l|l} K \times E & \longrightarrow E \\ (\lambda, x) & \longmapsto \lambda x \end{array}$$

vérifiant les quatre axiomes suivants, valables quels que soient les éléments λ et μ de K et les éléments x et y de E :

$$\begin{array}{ll} \text{(I)} & \lambda(x + y) = \lambda x + \lambda y. \\ \text{(II)} & (\lambda + \mu)x = \lambda x + \mu x. \\ \text{(III)} & \lambda(\mu x) = (\lambda\mu)x. \\ \text{(IV)} & 1x = x. \end{array}$$

Si K est un corps gauche — c'est-à-dire non commutatif —, on obtient ainsi un espace vectoriel à gauche sur K , et on définit un espace vectoriel à droite sur K en remplaçant λx par $x\lambda$ et les axiomes (I), (II), (III) et (IV) respectivement par : $(x + y)\lambda = x\lambda + y\lambda$, $x(\lambda + \mu) = x\lambda + x\mu$, $(x\lambda)\mu = x(\lambda\mu)$ et $x1 = x$.

Dans tout le chapitre, K est un corps commutatif, les espaces vectoriels sont des espaces vectoriels sur K et si l'on travaille dans un espace vectoriel E , K est le *corps de base*, les éléments de K sont les *scalaires* et ceux de E les *vecteurs*, et l'élément neutre 0 de $(E, +)$ est le *vecteur zéro* de l'espace vectoriel E .

On utilise toutes les notions classiques de la structure d'espace vectoriel sur un corps commutatif et leurs propriétés fondamentales, ainsi que celles des matrices à coefficients dans un corps commutatif et des déterminants¹.

1. Voir par exemple [ARN1], chapitres VI, IX et de XI à XV, ou [RAM1], chapitres 9 à 12.

■ Nécessité des axiomes

Nous allons montrer grâce à quatre exemples que chacun des axiomes de la définition 4.1 est nécessaire pour définir la structure d'espace vectoriel telle que nous la connaissons, c'est-à-dire qu'aucun des quatre n'est une conséquence des trois autres. Ces exemples sont tous obtenus à partir d'un espace vectoriel connu dont nous conservons l'addition, et sur lequel nous définissons une nouvelle loi de composition externe ayant pour domaine le même corps, que nous notons $*$.

4.1. Axiome (I) non vérifié.

Nous considérons un espace vectoriel E de dimension finie $n \geq 2$ sur le corps des nombres complexes, nous choisissons un sous-espace vectoriel D de dimension 1 de E et nous introduisons la loi de composition externe $*$ de domaine $\mathbb{C}$ sur E définie ainsi :

$$\left\{ \begin{array}{l} \mathbb{C} \times E \longrightarrow E \\ (\lambda, x) \longmapsto \lambda * x = \begin{cases} \lambda x & \text{si } x \in D, \\ \bar{\lambda} x & \text{si } x \notin D. \end{cases} \end{array} \right.$$

Nous choisissons un vecteur y_0 non nul de D et un vecteur z_0 n'appartenant pas à D — il en existe puisque $n \geq 2$. Alors le vecteur $y_0 + z_0$ n'appartient pas à D donc, en choisissant un complexe non réel λ (par exemple $\lambda = i$), on a :

$$\lambda * (y_0 + z_0) = \bar{\lambda}(y_0 + z_0) = \bar{\lambda}y_0 + \bar{\lambda}z_0 \quad \text{et} \quad \lambda * y_0 + \lambda * z_0 = \lambda y_0 + \bar{\lambda}z_0$$

ce qui montre que $\lambda * (y_0 + z_0) \neq \lambda * y_0 + \lambda * z_0$; l'axiome (I) n'est donc pas vérifié.

En revanche $\bar{1} = 1$ et, quels que soient les scalaires λ et μ , $\overline{\lambda + \mu} = \bar{\lambda} + \bar{\mu}$ et $\overline{\lambda\mu} = \bar{\lambda}\bar{\mu}$, donc les axiomes (IV), (II) et (III) sont clairement vérifiés. $\square$

Nous pouvons cependant remarquer que si les axiomes (II), (III) et (IV) sont vérifiés, l'axiome (I) s'applique à tout élément λ du sous-corps premier² du corps de base.

4.2. Axiome (II) non vérifié.

Nous considérons un espace vectoriel réel³ E différent de $\{0\}$ et nous définissons la loi de composition externe $*$ de domaine $\mathbb{R}$ sur E de la manière suivante : pour tout vecteur x de E et tout nombre réel λ , $\lambda * x = x$.

Nous choisissons un vecteur non nul x_0 de E et nous posons $\mu = \lambda = 1$. Alors $(\lambda + \mu) * x_0 = x_0$ et $\lambda * x_0 + \mu * x_0 = x_0 + x_0 = 2x_0$. Par suite l'axiome (II) n'est pas vérifié. Quels que soient les vecteurs x et y de E et les nombres réels λ et μ , on a $\lambda * (x + y) = x + y = \lambda * x + \lambda * y$, $\lambda * (\mu * x) = \mu * x = x = (\lambda\mu) * x$ et $1 * x = x$, donc les axiomes (I), (III) et (IV) sont vérifiés. $\square$

4.3. Axiome (III) non vérifié.

Nous considérons un espace vectoriel complexe⁴ $E \neq \{0\}$ et nous définissons la loi de composition externe $*$ de domaine $\mathbb{C}$ sur E ainsi : pour tout vecteur x de E et tout nombre complexe λ , $\lambda * x = \operatorname{Re}(\lambda)x$ — où $\operatorname{Re}(\lambda)$ désigne la partie réelle de λ .

2. Le sous-corps premier d'un corps K est le sous-corps de K engendré par l'élément unité de K ou encore le plus petit modulo l'inclusion des sous-corps de K . Par exemple $\mathbb{Q}$ est le sous-corps premier des corps $\mathbb{R}$ et $\mathbb{C}$.

3. Ceci signifie qu'il s'agit d'un espace vectoriel sur le corps $\mathbb{R}$ des nombres réels.

4. Ceci signifie qu'il s'agit d'un espace vectoriel sur le corps $\mathbb{C}$ des nombres complexes.

En choisissant un vecteur non nul x_0 de E , on a $i^2 * x_0 = (-1) * x_0 = -x_0$ et $i * (i * x) = i * 0 = 0$, donc l'axiome (iii) n'est pas vérifié. L'axiome (i) est clairement vérifié et, comme $\operatorname{Re}(1) = 1$ et que, quels que soient les nombres complexes λ et μ , $\operatorname{Re}(\lambda + \mu) = \operatorname{Re}(\lambda) + \operatorname{Re}(\mu)$, il en est de même des axiomes (iv) et (ii). $\square$

Plus généralement, en considérant un espace vectoriel $E \neq \{0\}$ sur un corps K et un endomorphisme f du groupe additif $(K, +)$ tel que $f(1) = 1$ mais qui n'est pas un morphisme pour la multiplication, et en définissant la loi de composition externe $*$ de domaine K sur E par $\lambda * x = f(\lambda)x$, les axiomes (i), (ii) et (iv) sont vérifiés, mais pas l'axiome (iii).

4.4. Axiome (iv) non vérifié.

Nous munissons $E = \mathbb{R}^2$ de sa structure canonique d'espace vectoriel sur $\mathbb{R}$ et nous définissons la loi de composition externe $*$ de domaine $\mathbb{R}$ sur E de la manière suivante : pour tout vecteur $z = (x, y)$ de E et tout nombre réel λ , $\lambda * z = (\lambda x, 0)$. Comme $1 * (0, 1) = (0, 0) \neq (0, 1)$, l'axiome (iv) n'est pas vérifié, alors que des calculs faciles permettent de vérifier les trois premiers axiomes. $\square$

Plus généralement, en considérant un espace vectoriel E et des sous-espaces vectoriels F et G supplémentaires dans E et différents de $\{0\}$, en notant p la projection sur F parallèlement à G et en définissant la loi de composition externe $*$ de domaine K sur E par $\lambda * x = \lambda p(x)$, les trois premiers axiomes sont vérifiés mais, en choisissant un vecteur non nul x_0 de G , on a $1 * x_0 = 1 p(x_0) = p(x_0) = 0 \neq x_0$.

■ Sous-espaces vectoriels

Parmi les sous-espaces vectoriels d'un espace vectoriel E , figurent en particulier les droites vectorielles de E (sous-espaces vectoriels de dimension finie 1 de E), les plans vectoriels de E (sous-espaces vectoriels de dimension finie 2 de E) et les hyperplans vectoriels de E , éléments maximaux de l'ensemble ordonné par l'inclusion des sous-espaces vectoriels stricts de E , ou encore sous-espaces vectoriels admettant une droite vectorielle pour supplémentaire dans E . Les hyperplans vectoriels sont les noyaux des formes linéaires non nulles sur E et, si E est de dimension finie $n \geq 1$, ce sont les sous-espaces vectoriels de dimension $n - 1$ de E .

La réunion de deux sous-groupes stricts d'un groupe ne peut pas être égale au groupe tout entier donc, *a fortiori*, la réunion de deux sous-espaces vectoriels stricts n'est jamais égale à l'espace tout entier. Cependant, un espace vectoriel peut être la réunion de trois sous-espaces vectoriels stricts.

4.5. Plan vectoriel réunion de trois droites vectorielles.

Nous considérons le corps $\mathbb{F}_2 = \mathbb{Z}/2\mathbb{Z} (= \{0, 1\})$ des entiers modulo 2 et l'espace vectoriel $E = (\mathbb{F}_2)^2$ sur $\mathbb{F}_2$. Nous posons $D_1 = \{(0, 0), (1, 0)\}$, $D_2 = \{(0, 0), (0, 1)\}$ et $D_3 = \{(0, 0), (1, 1)\}$; D_1 est la droite vectorielle de E engendrée par $d_1 = (1, 0)$, D_2 la droite vectorielle engendrée par $d_2 = (0, 1)$ et D_3 la droite vectorielle engendrée par $d_3 = (1, 1)$. Or $D_1 \cup D_2 \cup D_3 = \{(0, 0), (1, 0), (0, 1), (1, 1)\} = E$, donc la réunion de D_1 , D_2 et D_3 est le plan vectoriel E tout entier. $\square$

Remarquons que si l'on munit l'espace vectoriel $E = (\mathbb{F}_2)^2$ sur $\mathbb{F}_2$ de sa structure canonique d'espace affine, les quatre points $(0, 0)$, $(0, 1)$, $(1, 1)$ et $(1, 0)$ du plan affine E forment un parallélogramme dont les diagonales sont parallèles⁵.

4.6. Espace vectoriel de dimension infinie, réunion de trois hyperplans vectoriels.

Nous considérons l'espace vectoriel $E = \mathbb{F}_2[X]$ des polynômes à coefficients dans le corps $\mathbb{F}_2 = \mathbb{Z}/2\mathbb{Z}$ des entiers modulo 2 et nous posons :

$$H_0 = \left\{ \sum_{n=0}^{+\infty} a_n X^n \mid a_0 = 0 \right\}, \quad H_1 = \left\{ \sum_{n=0}^{+\infty} a_n X^n \mid a_1 = 0 \right\}$$

et :

$$H_2 = \left\{ \sum_{n=0}^{+\infty} a_n X^n \mid a_0 = a_1 \right\}.$$

Les applications φ_0 , φ_1 et φ_2 de E dans $\mathbb{F}_2$ qui, au polynôme $P = \sum_{n=0}^{+\infty} a_n X^n$, associent respectivement $\varphi_0(P) = a_0$, $\varphi_1(P) = a_1$ et $\varphi_2(P) = a_0 - a_1$, sont des formes linéaires non nulles sur E , $H_0 = \text{Ker}(\varphi_0)$, $H_1 = \text{Ker}(\varphi_1)$ et $H_2 = \text{Ker}(\varphi_2)$, donc H_0 , H_1 et H_2 sont des hyperplans vectoriels de E .

Soit $P = \sum_{n=0}^{+\infty} a_n X^n$ un vecteur de E . Si $a_0 \neq 0$ et $a_1 \neq 0$, alors $a_0 = a_1 = 1$ — en effet, $\mathbb{F}_2 = \{0, 1\}$ — donc P appartient à H_2 . En conclusion, $E = H_0 \cup H_1 \cup H_2$. $\square$

Nous utilisons ici la notion de valuation d'un polynôme (définition 3.13, page 37). Si $P = \sum_{k=0}^{+\infty} a_k X^k$ est un polynôme et n un entier naturel, la valuation $\text{val } P$ de P est supérieure ou égale à n si, et seulement si, $a_k = 0$ pour tout entier naturel $k < n$, donc $\text{val } P \geq n$ si, et seulement si, X^n divise P dans l'anneau $K[X]$. Il en résulte que, pour tout entier naturel n , l'ensemble V_n des polynômes de valuation supérieure ou égale à n est l'idéal de $K[X]$ engendré par le polynôme X^n ; en particulier, V_n est un sous-espace vectoriel de l'espace vectoriel $K[X]$ sur K .

En dimension finie, un sous-espace vectoriel strict n'est jamais isomorphe à l'espace tout entier. Ceci est faux en dimension infinie.

4.7. Sous-espaces vectoriels stricts d'un espace vectoriel E isomorphes à E .

Nous considérons l'espace vectoriel réel $E = \mathbb{R}[X]$ et nous notons, pour tout entier naturel n , V_n le sous-espace vectoriel des polynômes de valuation supérieure ou égale à n ; V_n est aussi l'idéal de l'anneau $\mathbb{R}[X]$ engendré par X^n .

Soit n un entier naturel. Pour tout polynôme P , X^n divise $X^n P$ donc le polynôme $X^n P$ appartient à V_n , ce qui justifie l'existence de l'application :

$$\left| \begin{array}{ll} f_n : E & \longrightarrow V_n \\ P & \longmapsto f_n(P) = X^n P \end{array} \right.$$

qui est clairement linéaire. Si Q appartient à V_n , X^n divise Q donc il existe un polynôme P et un seul tel que $Q = X^n P$, c'est-à-dire $Q = f_n(P)$. Par suite, f_n est

5. Le milieu d'un segment n'existe pas puisque, dans le corps $\mathbb{F}_2$, $2 = 1 + 1 = 0$.

une bijection de E sur V_n , donc un isomorphisme d'espace vectoriel de E sur V_n . Pour tout entier naturel n , $V_{n+1} \subset V_n$, $X^n \in V_n$, $X^n \notin V_{n+1}$, donc $V_{n+1} \subsetneq V_n$; enfin, si $n \geq 1$, $X^{n-1} \notin V_n$, donc V_n est un sous-espace vectoriel strict de E .

En conclusion, $(V_n)_{n \in \mathbb{N}^*}$ est une suite strictement décroissante de sous-espaces vectoriels stricts de E tous isomorphes à E . $\square$

THÉORÈME 4.1. — Si F , G et H sont des sous-espaces vectoriels d'un espace vectoriel E , $(F \cap G) + H \subset (F + H) \cap (G + H)$.

En général l'égalité est fausse, comme le montre l'exemple suivant.

4.8. Sous-espaces vectoriels F , G et H d'un espace vectoriel E tels que $(F \cap G) + H \subsetneq (F + H) \cap (G + H)$.

Nous considérons trois droites vectorielles distinctes F , G et H d'un plan vectoriel E — par exemple $E = K^2$ et F (resp. G , resp. H) est la droite vectorielle de E engendrée par $a = (1, 0)$ (resp. $b = (0, 1)$, resp. $c = (1, 1)$). La réunion de deux droites vectorielles distinctes de E engendre le plan E , donc $F + H = G + H = E$, d'où l'égalité $(F + H) \cap (G + H) = E$. Par ailleurs, les droites vectorielles F et G étant distinctes, $F \cap G = \{0\}$ donc $(F \cap G) + H = H$. Par suite⁶ :

$$(F \cap G) + H \subsetneq (F + H) \cap (G + H). \quad \square$$

THÉORÈME 4.2. — Si F , G et H sont des sous-espaces vectoriels d'un espace vectoriel E , $(F \cap H) + (G \cap H) \subset (F + G) \cap H$.

L'égalité est en général fausse, comme dans l'exemple suivant.

4.9. Sous-espaces vectoriels F , G et H d'un espace vectoriel E tels que $(F \cap H) + (G \cap H) \subsetneq (F + G) \cap H$.

Nous considérons, comme dans l'exemple précédent 4.8, trois droites vectorielles distinctes F , G et H d'un plan vectoriel E . Alors $F \cap H = G \cap H = \{0\}$, donc $(F \cap H) + (G \cap H) = \{0\}$. Comme les droites vectorielles F et G sont distinctes, leur réunion engendre le plan E , donc $(F + G) \cap H = E \cap H = H$, ce qui montre que⁷ $(F \cap H) + (G \cap H) \subsetneq (F + G) \cap H$. $\square$

■ Endomorphismes d'un espace vectoriel

Si E et F sont des espaces vectoriels, l'ensemble $\mathcal{L}(E, F)$ des applications linéaires de E dans F est un espace vectoriel sur K — sous-espace vectoriel de l'espace vectoriel sur K des applications de E dans F . Si E est un espace vectoriel, on dispose donc de l'espace vectoriel $\mathcal{L}(E) = \mathcal{L}(E, E)$ des endomorphismes de E et

6. Considérons un espace vectoriel E , notons $\mathcal{G}(E)$ l'ensemble des sous-espaces vectoriels de E et munissons $\mathcal{G}(E)$ de la relation d'inclusion. Pour des sous-espaces vectoriels F et G de E , $F + G$ et $F \cap G$ sont respectivement la borne supérieure et la borne inférieure de $\{F, G\}$ dans l'ensemble ordonné $(\mathcal{G}(E), \subset)$, donc l'inclusion fait de $\mathcal{G}(E)$ un treillis (définition 1.11, page 11). On prouve aisément que c'est un treillis modulaire (définition 1.13, page 13), mais l'inclusion stricte de l'exemple 4.8 montre qu'en général ce n'est pas un treillis distributif (définition 1.12, page 12).

7. C'est un nouvel exemple de treillis non distributif, obtenu en munissant $\mathcal{G}(E)$ (voir la note précédente) de la relation d'ordre $\supset$; alors les bornes supérieures et inférieures s'intervertissent.

$\mathcal{L}(E)$, muni de l'addition de cette structure d'espace vectoriel et de $\circ$, loi induite sur $\mathcal{L}(E)$ par la composition des applications, est un anneau dont l'élément unité est l'application identique Id_E de E .

THÉORÈME 4.3. — Théorème du rang.

Si E et F sont des espaces vectoriels et f une application linéaire de E dans F , et si E est de dimension finie, l'espace vectoriel $\text{Im}(f)$ est de dimension finie et :

$$\dim(E) = \dim(\text{Ker } f) + \dim(\text{Im } f) = \dim(\text{Ker } f) + \text{rg}(f).$$

En particulier, si f est un endomorphisme d'un espace vectoriel E de dimension finie, le noyau et l'image de f sont des sous-espaces vectoriels de E et on a l'égalité $\dim(E) = \dim(\text{Ker } f) + \dim(\text{Im } f)$, ce qui ne signifie nullement que le noyau et l'image de f sont en somme directe, et il est même possible, si E est de dimension finie paire, que $\text{Im}(f) = \text{Ker}(f)$.

4.10. Endomorphisme f tel que $\text{Im}(f) = \text{Ker}(f)$.

Nous introduisons l'espace vectoriel $E = \mathbb{K}^2$ sur $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$ et l'endomorphisme $f : z = (x, y) \mapsto f(z) = (y, 0)$ de l'espace vectoriel E . Si $z = (x, y)$ est un vecteur de E , $f(z) = 0$ si, et seulement si, $(y, 0) = (0, 0)$, ce qui équivaut à $y = 0$. Par suite $\text{Ker}(f) = \{(x, 0) \mid x \in \mathbb{K}\} = \{\lambda(1, 0) \mid \lambda \in \mathbb{K}\}$, donc $\text{Ker}(f)$ est la droite vectorielle D de E engendrée par le vecteur $d = (1, 0)$. Or la définition de f montre que $\text{Im}(f) = \{(y, 0) \mid y \in \mathbb{K}\} = D$, donc $\text{Im}(f) = \text{Ker}(f)$. $\square$

THÉORÈME 4.4. — Si E est un espace vectoriel de dimension finie et f un endomorphisme de E , il y a équivalence entre les assertions :

- (I) f est injectif.
- (II) f est une surjection de E sur E .
- (III) f est une bijection de E sur E .

L'assertion (III) s'écrit aussi : f est un automorphisme de l'espace vectoriel E .

Ces équivalences deviennent fausses en dimension quelconque.

4.11. Endomorphisme injectif qui n'est pas surjectif.

Nous considérons l'espace vectoriel $E = \mathbb{R}[X]$ sur $\mathbb{R}$ et l'application, clairement linéaire :

$$\left| \begin{array}{l} f : E \longrightarrow E \\ P \longmapsto f(P) = XP \end{array} \right.$$

L'anneau $\mathbb{R}[X]$ étant un anneau intègre, $\text{Ker}(f) = \{0\}$, donc l'endomorphisme f de E est injectif. L'image de f est l'ensemble des polynômes dont 0 est une racine, donc par exemple $X^0 \notin \text{Im}(f)$, ce qui montre que $\text{Im}(f) \subsetneq E$. Il en résulte que l'endomorphisme f n'est pas une surjection de E sur E . $\square$

Si E est un espace vectoriel et f un endomorphisme de E , la puissance k -ième f^k de f dans l'anneau $\mathcal{L}(E)$ est, pour tout entier naturel k , l'endomorphisme de E défini par $f^0 = \text{Id}_E$ et, pour $k \geq 1$, $f^k = \underbrace{f \circ \dots \circ f}_{k \text{ facteurs}}$.

Nous remarquons que dans l'exemple 4.11, la suite $(I_k)_{k \in \mathbb{N}}$ des images itérées de f , de terme général $I_k = \text{Im}(f^k)$, est une suite strictement décroissante de sous-espaces vectoriels de E .

4.12. Endomorphisme surjectif qui n'est pas injectif.

Nous considérons l'espace vectoriel $E = \mathbb{R}[X]$ sur $\mathbb{R}$ et sa dérivation canonique :

$$\begin{cases} D : E \longrightarrow E \\ P \longmapsto D(P) = P'. \end{cases}$$

Le polynôme $X^0 = 1$ appartient à $\text{Ker}(D)$ et il n'est pas nul, donc l'endomorphisme D n'est pas injectif. Si $Q = a_0 + a_1X + \dots + a_qX^q$ est un polynôme, $Q = D(P)$ pour le polynôme :

$$P = \sum_{k=0}^q \frac{a_k}{k+1} X^{k+1},$$

donc l'endomorphisme D est une surjection de E sur E . $\square$

Comme dans l'exemple 4.11 pour les images, nous remarquons que dans l'exemple 4.12, la suite $(N_k)_{k \in \mathbb{N}}$ des noyaux itérés de $f = D$, de terme général $N_k = \text{Ker}(f^k)$, est une suite strictement croissante de sous-espaces vectoriels de E .

THÉORÈME 4.5. — Si f et g sont des endomorphismes d'un même espace vectoriel E , alors $\text{Im}(f + g) \subset \text{Im}(f) + \text{Im}(g)$ et $\text{Ker}(f) \cap \text{Ker}(g) \subset \text{Ker}(f + g)$.

Cependant les inclusions du théorème 4.5 n'ont aucune raison d'être des égalités.

4.13. Endomorphismes f et g tels que $\text{Im}(f + g) \subsetneq \text{Im}(f) + \text{Im}(g)$.

Nous choisissons un espace vectoriel réel E différent de $\{0\}$ et un automorphisme f de E , c'est-à-dire un isomorphisme d'espace vectoriel de E sur E — par exemple $f = \text{Id}_E$ —, et nous posons $g = -f$. Comme $f + g = 0$, $\text{Im}(f + g) = \{0\}$, alors que $\text{Im}(f) = \text{Im}(g) = E$ donc $\text{Im}(f) + \text{Im}(g) = E$. $\square$

4.14. Endomorphismes f et g tels que $\text{Ker}(f) \cap \text{Ker}(g) \subsetneq \text{Ker}(f + g)$.

Nous choisissons un espace vectoriel réel E différent de $\{0\}$ et un automorphisme f de E , et nous posons $g = -f$. Comme $f + g = 0$, $\text{Ker}(f + g) = E$, alors que $\text{Ker}(f) = \text{Ker}(g) = \{0\}$ donc $\text{Ker}(f) \cap \text{Ker}(g) = \{0\}$. $\square$

DÉFINITION 4.2. — Un projecteur d'un espace vectoriel E est un endomorphisme f de E tel que $f^2 = f$.

THÉORÈME 4.6. — Théorème des projecteurs.

Si f est un projecteur d'un espace vectoriel E , alors $E = \text{Im}(f) \oplus \text{Ker}(f)$ et f est la projection sur $\text{Im}(f)$ parallèlement à $\text{Ker}(f)$.

Si F et G sont des sous-espaces vectoriels supplémentaires d'un espace vectoriel E , la projection f sur F parallèlement à G est un projecteur de E , $\text{Im}(f) = F$ et $\text{Ker}(f) = G$, donc $E = \text{Im}(f) \oplus \text{Ker}(f)$. Cependant un endomorphisme f d'un espace vectoriel E tel que $E = \text{Im}(f) \oplus \text{Ker}(f)$ n'est pas forcément un projecteur.

4.15. Endomorphisme f de E tel que $E = \text{Im}(f) \oplus \text{Ker}(f)$ mais qui n'est pas un projecteur.

Nous considérons l'espace vectoriel $E = \mathbb{R}^2$ sur $\mathbb{R}$ et l'application :

$$\left| \begin{array}{l} f : E \longrightarrow E \\ z = (x, y) \longmapsto f(z) = (x - 2y, 6y - 3x), \end{array} \right.$$

qui est clairement linéaire. Un vecteur $z = (x, y)$ de E appartient à $\text{Ker}(f)$ si, et seulement si, $x = 2y$, donc $\text{Ker}(f) = \{(2y, y) \mid y \in \mathbb{R}\} = \{\lambda(2, 1) \mid \lambda \in \mathbb{R}\}$; par suite, $\text{Ker}(f)$ est la droite vectorielle de E engendrée par le vecteur $e = (2, 1)$. Si $z = (x, y)$ est un vecteur de E , on a $f(z) = (\lambda, -3\lambda) = \lambda(1, -3)$ où $\lambda = x - 2y$; de plus, pour tout réel λ , $\lambda(1, -3) = (\lambda, -3\lambda) = f((\lambda, 0))$. Par conséquent, $\text{Im}(f)$ est la droite vectorielle de E engendrée par le vecteur $d = (1, -3)$. La famille (d, e) est clairement une base de E , donc $E = \text{Im}(f) \oplus \text{Ker}(f)$. Or, en posant $a = (1, 0)$, on a $f(a) = (1, -3)$ donc $f^2(a) = f(f(a)) = (7, -21) \neq f(a)$, ce qui montre que $f^2 \neq f$. Par conséquent, f n'est pas un projecteur de E . $\square$

En fait, si f est un endomorphisme de E , $\text{Im}(f) \cap \text{Ker}(f) = \{0\}$ si, et seulement si, $\text{Ker}(f^2) = \text{Ker}(f)$, et $E = \text{Im}(f) + \text{Ker}(f)$ si, et seulement si, $\text{Im}(f^2) = \text{Im}(f)$, donc $E = \text{Im}(f) \oplus \text{Ker}(f)$ si, et seulement si, $\text{Ker}(f^2) = \text{Ker}(f)$ et $\text{Im}(f^2) = \text{Im}(f)$.

DÉFINITION 4.3. — Un endomorphisme f d'un espace vectoriel E est nilpotent s'il existe un entier naturel p tel que $f^p = 0$ — ce qui signifie que $f^p(x) = 0$ pour tout vecteur x de E .

Si $E \neq \{0\}$, alors $\text{Id}_E \neq 0$, donc, si f est un endomorphisme nilpotent de E , tout entier naturel p tel que $f^p = 0$ est supérieur ou égal à 1. Si E est de dimension finie $n \geq 1$, alors $f^n = 0$ pour tout endomorphisme nilpotent f de E .

4.16. Endomorphisme f de E tel que, pour tout vecteur x de E , il existe $p \in \mathbb{N}$ tel que $f^p(x) = 0$, mais qui n'est pas nilpotent.

Nous considérons l'espace vectoriel $E = \mathbb{R}[X]$ sur $\mathbb{R}$ et sa dérivation canonique f , c'est-à-dire l'endomorphisme $f : P \mapsto f(P) = P'$ de E . On a, pour tout entier naturel p , $f^p(X^p) = p! \neq 0$ donc $f^p \neq 0$. Par suite l'endomorphisme f de E n'est pas nilpotent. Cependant, pour tout polynôme P , on obtient, en choisissant un entier p strictement plus grand que le degré de P , l'égalité $f^p(P) = 0$. $\square$

■ Valeurs propres et vecteurs propres, polynôme caractéristique et polynôme minimal

Dans tout ce paragraphe, n est un entier naturel non nul, $\mathcal{M}_n(K)$ est l'espace vectoriel et l'anneau des matrices carrées d'ordre n à coefficients dans K et I_n la matrice unité d'ordre n .

Soit f un endomorphisme d'un espace vectoriel E . Une *valeur propre* de f est un scalaire λ pour lequel il existe au moins un vecteur non nul a de E tel que $f(a) = \lambda a$. Si λ est une valeur propre de f , le *sous-espace propre* de f associé à λ est $V(f, \lambda) = \text{Ker}(f - \lambda \text{Id}_E) = \{x \in E \mid f(x) = \lambda x\}$. Un *vecteur propre* de f est un vecteur $a \neq 0$ pour lequel il existe un scalaire λ tel que $f(a) = \lambda a$.

Soit $A \in \mathcal{M}_n(K)$. Les valeurs propres, sous-espaces propres et vecteurs propres de A sont ceux de l'endomorphisme Φ_A de K^n canoniquement associé à A , c'est-à-dire de l'application linéaire $\Phi_A : X \mapsto \Phi_A(X) = AX$ de K^n dans K^n . Si λ est une valeur propre de A , le sous-espace propre de A associé à λ est le sous-espace vectoriel $V(A, \lambda) = V(\Phi_A, \lambda) = \text{Ker}(\Phi_A - \lambda \text{Id}_{K^n}) = \{X \in K^n \mid AX = \lambda X\}$ de K^n des solutions du système linéaire homogène $AX = 0$. En travaillant sur le corps $K(X)$ des fractions rationnelles, dont K est un sous-corps, la matrice $A - \lambda \text{Id}_n$ est à coefficients dans le sous-anneau $K[X]$ de $K(X)$, donc son déterminant $\Gamma_A(X) = \det(A - X \text{Id}_n)$ est un polynôme, appelé le *polynôme caractéristique* de A , polynôme de degré n tel que, pour tout scalaire λ , $\Gamma_A(\lambda) = \det(A - \lambda \text{Id}_n)$.

Si des matrices A et B carrées d'ordre n sont semblables, elles ont le même polynôme caractéristique donc, si f est un endomorphisme d'un espace vectoriel E de dimension finie n , les matrices représentatives de f dans une base de E ont toutes le même polynôme caractéristique, ce qui justifie la définition suivante : le polynôme caractéristique de f est $\Gamma_f(X) = \Gamma_A(X)$ où A est la matrice représentative de f dans une base (quelconque) de E ; alors, $\Gamma_f(\lambda) = \det(f - \lambda \text{Id}_E)$ pour tout élément λ de K , d'où le théorème suivant.

THÉORÈME 4.7. — Si E est un espace vectoriel de dimension finie et si $f \in \mathcal{L}(E)$, les valeurs propres de f sont les racines dans K du polynôme caractéristique de f .

Si f est un endomorphisme d'un espace vectoriel E de dimension finie n et A la matrice représentative de f dans une base de E , f et A ont les mêmes valeurs propres et, si λ est l'une d'elles, les sous-espaces propres $V(f, \lambda)$ et $V(A, \lambda)$ sont isomorphes. On en déduit que si des matrices A et B carrées d'ordre n sont semblables, elles ont les mêmes valeurs propres et que, si λ est l'une d'elles, les sous-espaces propres $V(A, \lambda)$ et $V(B, \lambda)$ sont isomorphes.

Si E est un espace vectoriel de dimension finie n et f un endomorphisme de E , le polynôme caractéristique de f est de degré n , donc f admet au plus n valeurs propres ; par conséquent le spectre de f , c'est-à-dire l'ensemble de ses valeurs propres, est fini. Ce résultat tombe en défaut en dimension quelconque.

4.17. Endomorphisme admettant tout élément du corps pour valeur propre.

Nous considérons l'espace vectoriel $\mathcal{D}$ sur le corps $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$ des applications indéfiniment dérivables de $\mathbb{R}$ dans $\mathbb{K}$ et l'endomorphisme de dérivation :

$$\left| \begin{array}{l} D : \mathcal{D} \longrightarrow \mathcal{D} \\ f \longmapsto D(f) = f' \end{array} \right.$$

de l'espace vectoriel $\mathcal{D}$ sur $\mathbb{K}$. Nous introduisons, pour tout élément λ de $\mathbb{K}$, l'application :

$$\left| \begin{array}{l} h_\lambda : \mathbb{R} \longrightarrow \mathbb{K} \\ x \longmapsto h_\lambda(x) = e^{\lambda x}. \end{array} \right.$$

Si λ appartient à $\mathbb{K}$, $h_\lambda'(x) = \lambda e^{\lambda x}$ pour tout nombre réel x , donc $D(h_\lambda) = \lambda h_\lambda$, ce qui montre que λ est une valeur propre de la dérivation D et h_λ un vecteur propre associé à λ . Ainsi, tout élément du corps $\mathbb{K}$ est une valeur propre de D . $\square$

Tout endomorphisme f d'un espace vectoriel de dimension finie n sur $\mathbb{C}$ possède au moins une valeur propre — en effet, le polynôme caractéristique de f admet au moins une racine dans $\mathbb{C}$ (théorème de d'Alembert-Gauss). Ceci devient faux en dimension infinie.

4.18. Endomorphisme d'un espace vectoriel complexe sans valeur propre.

Nous introduisons l'espace vectoriel $E = \mathbb{C}[X]$ sur le corps $\mathbb{C}$ des nombres complexes et l'endomorphisme $f : P \mapsto f(P) = XP$ de E . Soit P un polynôme différent de zéro. Son degré n est un entier naturel et le degré de $f(P) = XP$ est $n + 1$. Or, pour tout nombre complexe λ , λP est nul ou de degré n , donc l'égalité $f(P) = \lambda P$ est impossible. Ainsi, f n'admet aucune valeur propre. $\square$

Un endomorphisme d'un espace vectoriel réel de dimension 3 admet au moins une valeur propre ; en effet, son polynôme caractéristique est de degré 3 et 3 est impair, donc il admet au moins une racine réelle.

4.19. Endomorphisme sans valeur propre d'un espace vectoriel de dimension 3 sur $\mathbb{Q}$.

Nous notons f l'endomorphisme de l'espace vectoriel $E = \mathbb{Q}^3$ sur $\mathbb{Q}$ dont la matrice représentative dans la base canonique est :

$$\begin{pmatrix} 1 & -1 & 0 \\ 1 & -1 & -2 \\ 1 & 0 & 0 \end{pmatrix}.$$

Le calcul montre que le polynôme caractéristique de f est $\Gamma_f(X) = -X^3 + 2$. Il n'existe aucun rationnel r tel que $r^3 = 2$ — voir l'exemple 3.18 (page 41) — donc $\Gamma_f(X)$ n'admet pas de racine rationnelle : f n'admet pas de valeur propre. $\square$

Le polynôme $-X^3 + 2$ est de degré 3 et n'admet pas de racine dans $\mathbb{Q}$, donc il est irréductible sur $\mathbb{Q}$ (théorème 3.4, page 40). Par suite, f étant l'endomorphisme de l'exemple 4.19, E n'admet aucun sous-espace vectoriel stable par f autre que $\{0\}$ et E ; en effet, si F était un sous-espace vectoriel de dimension 1 ou 2 stable par f , le polynôme caractéristique de la restriction de f à F serait un diviseur de degré 1 ou 2 de $-X^3 + 2$ dans $\mathbb{Q}[X]$, ce qui est impossible puisque $-X^3 + 2$ est irréductible.

4.20. Endomorphisme sans valeur propre d'un espace vectoriel de dimension 3 sur $\mathbb{F}_5 = \mathbb{Z}/5\mathbb{Z}$.

Nous introduisons l'espace vectoriel $E = (\mathbb{F}_5)^3$ sur le corps $\mathbb{F}_5 = \mathbb{Z}/5\mathbb{Z}$ des entiers modulo 5 et l'endomorphisme g de E de matrice représentative :

$$\begin{pmatrix} 0 & 1 & 1 \\ 0 & 0 & 1 \\ 4 & 0 & 0 \end{pmatrix}$$

dans la base canonique. Un calcul facile montre que le polynôme caractéristique de g est $\Gamma_g(X) = -X^3 - X - 1 = 4X^3 + 4X + 4$. En substituant à X les cinq éléments de $\mathbb{F}_5$, on vérifie que le polynôme $\Gamma_g(X)$ n'admet pas de racine dans $\mathbb{F}_5$. Par conséquent l'endomorphisme g n'admet aucune valeur propre. $\square$

Comme dans l'exemple 4.19, le polynôme caractéristique de l'endomorphisme g de l'exemple 4.20 est irréductible, ici sur $\mathbb{F}_5$, donc il n'existe aucun sous-espace vectoriel de E stable par g autre que $\{0\}$ et E .

DÉFINITION 4.4. — Un endomorphisme f d'un espace vectoriel E de dimension finie est trigonalisable (resp. diagonalisable) s'il existe une base de E dans laquelle la matrice représentative de f est triangulaire (resp. diagonale).

Si E est un espace vectoriel de dimension finie n et f un endomorphisme trigonalisable de E , il existe une base de E dans laquelle la matrice représentative A de f est triangulaire et, en notant $(\alpha_1, \dots, \alpha_n)$ la diagonale de A , le polynôme caractéristique de f est $\Gamma_f(X) = (\alpha_1 - X) \cdots (\alpha_n - X)$, donc ses valeurs propres sont $\alpha_1, \dots, \alpha_n$; en particulier, f admet au moins une valeur propre.

Si le corps K est algébriquement clos, tout endomorphisme d'un espace vectoriel de dimension finie est trigonalisable⁸, mais il n'est pas forcément diagonalisable.

4.21. Endomorphisme non trigonalisable.

Nous considérons l'endomorphisme f de l'espace vectoriel $\mathbb{R}^2$ sur $\mathbb{R}$ dont la matrice représentative dans la base canonique de $\mathbb{R}^2$ est :

$$\begin{pmatrix} 1 & -1 \\ 1 & 1 \end{pmatrix}.$$

Le polynôme caractéristique de f est $(1 - X)(1 - X) - (-1) = X^2 - 2X + 2$, c'est-à-dire $(X - 1)^2 + 1$. Ce polynôme n'ayant pas de racine réelle, f n'admet aucune valeur propre, donc f n'est pas trigonalisable. $\square$

4.22. Endomorphisme non diagonalisable.

Nous posons $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$, et nous introduisons l'endomorphisme f de l'espace vectoriel $\mathbb{K}^2$ sur $\mathbb{K}$ dont la matrice représentative dans la base canonique est :

$$\begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}.$$

Le polynôme caractéristique de f est $(1 - X)(1 - X) - 0 = (X - 1)^2$, donc f admet 1 pour unique valeur propre. Par suite, si f était diagonalisable, il existerait une base de $\mathbb{K}^2$ dans laquelle sa matrice représentative serait :

$$\begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$$

et par conséquent f serait l'application identique de $\mathbb{K}^2$, ce qui n'est pas le cas. Il en résulte que l'endomorphisme f n'est pas diagonalisable. $\square$

Nous rappelons maintenant la définition d'un *polynôme minimal*.

Soit $\mathcal{A}$ une algèbre sur K ; c'est un espace vectoriel muni d'une multiplication qui en fait un anneau telle que, quels que soient $x, y \in \mathcal{A}$ et $\lambda, \mu \in K$, $(\lambda x)(\mu y) = (\lambda \mu)xy$.

8. En fait, un endomorphisme d'un espace vectoriel de dimension finie est trigonalisable si, et seulement si, son polynôme caractéristique est scindé sur le corps de base.

Si x est un vecteur de $\mathcal{A}$, les *polynômes annulateurs* de x sont les polynômes P à coefficients dans K tels que $P(x) = 0$ et, l'application $P \mapsto P(x)$ étant un morphisme d'anneau de $K[X]$ dans $\mathcal{A}$, l'ensemble des polynômes annulateurs de x est un idéal de $K[X]$, appelé l'*idéal annulateur* de x . Si un vecteur x de $\mathcal{A}$ admet au moins un polynôme annulateur différent de zéro, le *polynôme minimal* de x est le générateur unitaire de son idéal annulateur, ou encore le polynôme annulateur unitaire de x de plus bas degré. Si $\mathcal{A} \neq \{0\}$, alors, pour tout $x \in \mathcal{A}$, $X^0(x) = 1 \neq 0$, donc $1 = X^0$ n'est le polynôme minimal d'aucun vecteur de $\mathcal{A}$.

Si E est un espace vectoriel différent de $\{0\}$, $\mathcal{L}(E)$ est une algèbre sur K et $\text{Id}_E \neq 0$, et de même $\mathcal{M}_n(K)$ est une algèbre sur K et $I_n \neq 0$. Les notions de polynôme annulateur et de polynôme minimal s'appliquent donc à tout endomorphisme d'un espace vectoriel et à toute matrice carrée d'ordre n . Si $A \in \mathcal{M}_n(K)$, la famille $(A^k)_{0 \leq k \leq n^2}$ est liée dans l'espace vectoriel $\mathcal{M}_n(K)$ — il est de dimension finie égale à n^2 — donc A admet au moins un polynôme annulateur différent de zéro.

Soit f un endomorphisme d'un espace vectoriel E de dimension finie n et A la matrice représentative de f dans une base $\mathcal{B}$ de E . Pour tout polynôme P , la matrice représentative de $P(f)$ dans $\mathcal{B}$ est $P(A)$, donc l'idéal annulateur de f est celui de A . Par conséquent f admet donc au moins un polynôme annulateur différent de zéro et f et A ont le même polynôme minimal. Il en résulte que des matrices A et B , carrées d'ordre n et semblables, ont le même polynôme minimal.

THÉORÈME 4.8. — Théorème de Cayley-Hamilton⁹.

Si f est un endomorphisme d'un espace vectoriel de dimension finie, le polynôme minimal de f divise son polynôme caractéristique, et si A est une matrice carrée, le polynôme minimal de A divise son polynôme caractéristique.

Le théorème de Cayley-Hamilton exprime que si f (resp. A) est un endomorphisme d'un espace vectoriel de dimension finie (resp. une matrice carrée), le polynôme caractéristique de f (resp. de A) en est un polynôme annulateur.

Si deux matrices carrées d'ordre n sont semblables, elles ont le même polynôme caractéristique et le même polynôme minimal. Cependant la réciproque est fausse.

4.23. Deux matrices non semblables ayant le même polynôme caractéristique et le même polynôme minimal.

Nous considérons les matrices carrées d'ordre 4 à coefficients dans $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$:

$$M = \begin{pmatrix} 1 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix} \quad \text{et} \quad N = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

et nous notons f et g les endomorphismes de l'espace vectoriel $E = \mathbb{K}^4$ dont les matrices représentatives dans la base canonique (e_1, e_2, e_3, e_4) sont respectivement

9. En recherchant l'inverse d'un quaternion, William Hamilton démontre, en 1853, le résultat pour la dimension 4 sans vraiment l'exprimer. Arthur Cayley énonce le résultat pour des matrices carrées d'ordre n , le démontre pour $n = 2$, prétend l'avoir fait pour $n = 3$ et dit qu'il ne lui semble pas nécessaire de le démontrer dans le cas général... Georg Frobenius fournit la première démonstration générale en 1878. Pour une preuve de ce théorème, voir [ARN1], §XV.4.

M et N . Remarquons que les sous-espaces vectoriels F , engendré par (e_1, e_2) , et G , engendré par (e_3, e_4) , sont supplémentaires dans E et stables par f et par g . Les restrictions $f|_F$ et $f|_G$ ont la même matrice représentative, à savoir :

$$A = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix},$$

respectivement dans les bases (e_1, e_2) de F et (e_3, e_4) de G . On en déduit que le polynôme minimal de f est celui de la matrice A . Le calcul donne $(A - I_2)^2 = 0$, donc $(X - 1)^2$ est un polynôme annulateur de A . Par conséquent, le polynôme minimal de A est un diviseur unitaire de $(X - 1)^2$, mais ce n'est ni 1, ni $X - 1$ — puisque $A \neq I_2$ —, c'est donc $(X - 1)^2$. Il en résulte que le polynôme minimal de f est $(X - 1)^2$, et c'est celui de M . Les restrictions $g|_F$ et $g|_G$ ont pour matrices représentatives respectives dans les bases (e_1, e_2) de F et (e_3, e_4) de G :

$$B = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \text{ et } A = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix},$$

donc leurs polynôme minimaux respectifs sont $X - 1$ et $(X - 1)^2$. Par suite $(X - 1)^2$ est un polynôme annulateur de g , donc le polynôme minimal de g est un diviseur unitaire de $(X - 1)^2$ et, comme $g \neq \text{Id}_E$, on conclut comme ci-dessus que le polynôme minimal de g est $(X - 1)^2$: c'est le même que celui de f .

Les matrices M et N étant triangulaires et de diagonale $(1, 1, 1, 1)$, M et N ont le même polynôme caractéristique, à savoir $(1 - X)^4 = (X - 1)^4$.

Si les matrices M et N étaient semblables, les sous-espaces propres $V(M, 1)$ et $V(N, 1)$ des matrices M et N relatifs à la valeur propre 1 seraient isomorphes, donc de même dimension. Or l'étude des systèmes linéaires homogènes :

$$(M - I_4) \begin{pmatrix} x \\ y \\ z \\ t \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \end{pmatrix} \text{ et } (N - I_4) \begin{pmatrix} x \\ y \\ z \\ t \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}$$

montre facilement que $V(M, 1)$ est de dimension 2 et $V(N, 1)$ de dimension 3, donc les matrices M et N ne sont pas semblables. $\square$

Si f et g sont des endomorphismes d'un espace vectoriel E de dimension finie, les endomorphismes $f \circ g$ et $g \circ f$ ont le même polynôme caractéristique ; en revanche, $f \circ g$ et $g \circ f$ n'ont pas forcément le même polynôme minimal.

4.24. Endomorphismes f et g tels que $f \circ g$ et $g \circ f$ n'ont pas le même polynôme minimal.

Nous considérons les endomorphismes f et g de l'espace vectoriel $E = \mathbb{K}^2$ sur $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$ de matrices représentatives respectives dans la base canonique :

$$A = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \text{ et } B = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}$$

Le calcul donne : $AB = B^2 = 0$ et $BA = B$. Par conséquent, $f \circ g = 0$, donc le polynôme minimal de $f \circ g$ est un diviseur unitaire de X différent de 1, c'est donc X . Comme $g^2 = 0$, le polynôme minimal de g est un diviseur unitaire de X^2 , et ce n'est ni 1, ni X — puisque $g \neq 0$ —, c'est donc X^2 . Or $g \circ f = g$, donc le polynôme minimal de $g \circ f$ est X^2 . $\square$

En dimension finie, tout endomorphisme possède un polynôme minimal.

4.25. Endomorphisme sans polynôme minimal.

Nous considérons l'espace vectoriel $E = \mathbb{K}[X]$ sur le corps $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$ et sa dérivation canonique, c'est-à-dire l'endomorphisme $D : P \mapsto D(P) = P'$ de E .

Supposons que D admette un polynôme annulateur non nul et de degré n , à savoir :

$$A = \sum_{k=0}^n a_k X^k.$$

Alors $a_n \neq 0$ et l'égalité $A(D) = 0$ s'écrit dans $\mathcal{L}(E)$: (1) $\sum_{k=0}^{n-1} a_k D^k + a_n D^n = 0$.

En appliquant (1) à $P = X^n$, on obtient : (2) $\sum_{k=0}^{n-1} a_k P^{(k)} + n! a_n = 0$.

Or, pour tout entier naturel $k < n$, 0 est une racine d'ordre $n - k \geq 1$ du polynôme $P^{(k)}$ donc, en substituant 0 à X dans l'égalité (2), il vient : $n! a_n = 0$, ce qui est impossible puisque $n! \neq 0$ et $a_n \neq 0$.

Il en résulte que D n'admet aucun polynôme annulateur différent de zéro. Par suite l'endomorphisme D ne possède pas de polynôme minimal. $\square$

Terminons ce paragraphe par le rappel d'un théorème qui sera utilisé dans la suite.

THÉORÈME 4.9. — Si un endomorphisme f d'un espace vectoriel de dimension finie admet un polynôme annulateur scindé et à racines simples, en particulier si son polynôme caractéristique est scindé et à racines simples, alors f est diagonalisable.

■ Matrices

DÉFINITION 4.5. — Une matrice carrée est diagonalisable (resp. trigonalisable) si elle est semblable à une matrice diagonale (resp. triangulaire).

Une matrice carrée A est diagonalisable (resp. trigonalisable) si, et seulement si, tout endomorphisme de matrice représentative A dans une base est diagonalisable (resp. trigonalisable), ce qui équivaut à l'existence d'au moins un endomorphisme diagonalisable (resp. trigonalisable) de matrice représentative A dans une base. Il en résulte que le théorème 4.9 s'applique en remplaçant l'endomorphisme f par une matrice carrée et que, sur un corps algébriquement clos, par exemple $\mathbb{C}$, toute matrice carrée est trigonalisable.

Nous vérifions que la somme ou le produit de matrices diagonalisables n'a aucune raison d'être diagonalisable.

4.26. Matrices A et B diagonalisables telles que la matrice somme $A + B$ n'est pas diagonalisable.

Nous considérons les matrices carrées d'ordre 2, à coefficients dans un corps K :

$$A = \begin{pmatrix} 1 & 1 \\ 0 & 2 \end{pmatrix} \text{ et } B = \begin{pmatrix} -1 & 0 \\ 0 & -2 \end{pmatrix}.$$

La matrice B est évidemment diagonalisable et, le polynôme caractéristique de A étant $(X - 2)(X - 1)$, polynôme scindé et à racines simples, A est diagonalisable.

Le polynôme caractéristique de la matrice :

$$A + B = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$$

est X^2 , dont la seule racine est 0, donc, si $A + B$ était diagonalisable, la matrice $A + B$ serait semblable à la matrice nulle, et par suite elle serait nulle, ce qui est bien sûr faux. En conclusion, la matrice somme $A + B$ n'est pas diagonalisable. $\square$

4.27. Matrices A et B diagonalisables telles que la matrice produit AB n'est pas diagonalisable.

Nous introduisons les matrices carrées d'ordre 2, à coefficients dans un corps K :

$$A = \begin{pmatrix} 1 & 1 \\ 0 & 0 \end{pmatrix} \text{ et } B = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}.$$

La matrice B est évidemment diagonalisable et, comme le polynôme caractéristique $X^2 - X = X(X - 1)$ de A est scindé et à racines simples, A est diagonalisable. Le calcul donne :

$$AB = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$$

donc (voir l'exemple précédent 4.26) la matrice AB n'est pas diagonalisable. $\square$

THÉORÈME 4.10. — Théorème spectral.

Toute matrice carrée et symétrique à coefficients réels est diagonalisable et, si une telle matrice représente un endomorphisme f d'un espace vectoriel euclidien dans une base orthonormale, il existe une base orthonormale dans laquelle la matrice représentative de f est diagonale.

Une matrice carrée et symétrique à coefficients dans un corps autre que $\mathbb{R}$ n'est pas forcément diagonalisable.

4.28. Matrice carrée symétrique non diagonalisable.

Il suffit de travailler sur le corps $\mathbb{Q}$ des nombres rationnels et de considérer une matrice carrée dont les valeurs propres réelles sont irrationnelles. Prenons par exemple la matrice, à coefficients rationnels :

$$M = \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}.$$

Son polynôme caractéristique est $(1 - X)(-1 - X) - 1 = (X^2 - 1) - 1 = X^2 - 2$, qui n'admet pas de racine rationnelle (voir l'exemple 5.3, page 84), donc la matrice M n'a pas de valeurs propres. Par suite M n'est pas diagonalisable, ni même trigonalisable. $\square$

Dans l'exemple précédent, M est diagonalisable sur un corps « plus gros », puisque $\mathbb{Q}$ est un sous-corps de $\mathbb{R}$ et que M est diagonalisable sur $\mathbb{R}$ — son polynôme caractéristique $X^2 - 2 = (X - \sqrt{2})(X + \sqrt{2})$ est scindé sur $\mathbb{R}$ et à racines simples.

Nous examinons un exemple où une matrice carrée symétrique est trigonalisable sans être diagonalisable et un autre, sur un corps qui n'est pas un sous-corps ou un sur-corps de $\mathbb{R}$, où une matrice carrée symétrique n'est pas diagonalisable.

4.29. Matrice carrée symétrique qui est trigonalisable mais qui n'est pas diagonalisable.

Nous travaillons sur le corps $\mathbb{C}$ des nombres complexes. Nous posons $M = \begin{pmatrix} 2 & i \\ i & 0 \end{pmatrix}$.

Comme $\mathbb{C}$ est algébriquement clos, M est trigonalisable. Son polynôme caractéristique est $X^2 - 2X + 1 = (X - 1)^2$. Si la matrice M était diagonalisable, elle serait semblable à la matrice unité I_2 , donc égale à I_2 , ce qui est clairement faux. $\square$

4.30. Matrice carrée symétrique à coefficients dans le corps $\mathbb{F}_5$ des entiers modulo 5 qui n'est pas diagonalisable.

Nous considérons le corps $\mathbb{F}_5 = \mathbb{Z}/5\mathbb{Z}$ des entiers modulo 5 et nous posons :

$$M = \begin{pmatrix} 0 & 2 \\ 2 & 4 \end{pmatrix}.$$

Le calcul montre que le polynôme caractéristique de M est $\Gamma_M(X) = X^2 + X + 1$. En substituant à X chacun des cinq éléments de $\mathbb{F}_5$, on vérifie que $\Gamma_M(X)$ n'a pas de racine dans $\mathbb{F}_5$. On en déduit que M n'admet pas de valeur propre, donc que M n'est pas trigonalisable. *A fortiori* M n'est pas diagonalisable. $\square$

Le rang $\text{rg}(A)$ d'une matrice A à p lignes et n colonnes à coefficients dans K est le rang de l'application linéaire $\Phi_A : X \mapsto \Phi_A(X) = AX$ de K^n dans K^p canoniquement associé à A , et le rang de A s'interprète de trois manières : rang de toute application linéaire dont une matrice représentative est A ; rang de la famille des vecteurs-colonnes de A ; rang de la famille des vecteurs-lignes de A .

On a, quelles que soient les matrices A et B pour lesquelles le produit AB existe, $\text{rg}(AB) \leq \min(\text{rg}(A), \text{rg}(B))$. De plus, $\text{rg}({}^tAA) = \text{rg}(A)$ pour toute matrice A à coefficients réels. Sur un corps quelconque, on a bien sûr $\text{rg}({}^tAA) \leq \text{rg}(A)$, mais l'égalité est fautive dans le cas général.

4.31. Matrice A à coefficients complexes telle que $\text{rg}({}^tAA) \neq \text{rg}(A)$.

Nous considérons le corps $\mathbb{C}$ des nombres complexes et nous posons $A = \begin{pmatrix} 1 & 1 \\ i & i \end{pmatrix}$.

On voit que tAA est la matrice nulle, donc $\text{rg}({}^tAA) = 0 < 1 = \text{rg}(A)$. $\square$

4.32. Matrice A à coefficients dans le corps $\mathbb{F}_5 = \mathbb{Z}/5\mathbb{Z}$ des entiers modulo 5 telle que $\text{rg}({}^tAA) \neq \text{rg} A$.

Nous travaillons sur le corps $\mathbb{F}_5 = \mathbb{Z}/5\mathbb{Z}$ et nous posons $A = \begin{pmatrix} 1 & 1 \\ 2 & 2 \end{pmatrix}$.

Comme dans l'exemple précédent 4.31, ${}^tAA = 0$ donc $\text{rg}({}^tAA) = 0 < 1 = \text{rg}(A)$. $\square$

4.33. Matrices A et B telles que $\text{rg}(AB) \neq \text{rg}(BA)$.

Nous considérons les matrices $A = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}$ et $B = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}$.

Alors $AB = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}$ et $BA = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}$, donc $\text{rg}(AB) = 0 < 1 = \text{rg}(BA)$. $\square$

DÉFINITION 4.6. — La trace d'une matrice carrée A est la somme $\text{Tr}(A)$ de sa diagonale.

THÉORÈME 4.11. — Si A et B sont des matrices carrées de même ordre, on a :

$$\text{Tr}(AB) = \text{Tr}(BA).$$

Cependant, pour $n \geq 3$, on ne conserve pas la trace en opérant une permutation quelconque sur le produit de n matrices carrées.

4.34. Matrices carrées A, B et C de même ordre telles que :

$$\text{Tr}(ABC) \neq \text{Tr}(ACB).$$

Nous considérons les matrices $A = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}$, $B = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}$ et $C = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$.

La matrice AB est nulle, donc ABC est nulle ; par suite $\text{Tr}(ABC) = 0$.

Par ailleurs, $ACB = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}$, donc $\text{Tr}(ACB) = 1$. $\square$

Etant donné une (n, p) -matrice A et une (p, n) -matrice B , les matrices AB et BA ont les mêmes valeurs propres non nulles et, si A et B sont des matrices carrées, les matrices AB et BA ont les mêmes valeurs propres.

4.35. Matrices A et B telles que 0 est une valeur propre de AB mais pas de BA .

Le corps de base étant $\mathbb{R}$, nous posons $A = \begin{pmatrix} 1 & 1 \end{pmatrix}$ et $B = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$.

Comme $AB = (2)$, AB est inversible, donc 0 n'est pas une valeur propre de AB .

Or $BA = \begin{pmatrix} 1 & 1 \\ 1 & 1 \end{pmatrix}$, donc $(BA)\begin{pmatrix} 1 \\ -1 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$: 0 est une valeur propre de BA . $\square$

■ Modules

La notion d'espace vectoriel se généralise en celle de *module* quand on remplace, dans la définition 4.1 (page 60), le corps par un anneau commutatif non nul, les sous-espaces vectoriels devenant alors les sous-modules. Comme dans le cas des espaces vectoriels, pour un module M , les scalaires sont les éléments de l'anneau de base et les vecteurs ceux de M . Les définitions des applications linéaires, endomorphismes, automorphismes... s'appliquent sans changement aux modules. Cependant, un certain nombre de propriétés des espaces vectoriels ne se généralisent pas aux modules, en particulier la plus grande partie de ce qui concerne les familles et les parties libres ou liées et les bases, car on ne peut plus diviser par un scalaire non nul quelconque. Par exemple, certains modules ne possèdent pas de base. Un module qui admet au moins une base s'appelle un *module libre*.

Si A est un anneau commutatif non nul, son addition et, pour loi de composition externe de domaine A , sa multiplication, font de A un module sur lui-même ; cette structure de module est appelée la structure canonique de module sur A de l'anneau

A , les sous-modules du module A sont les idéaux de l'anneau A et A est un module libre — en effet, $\{1\}$ est clairement une base du module A sur lui-même.

Un module sur $\mathbb{Z}$ est d'abord un groupe abélien additif. Nous avons rappelé à la page 42 la définition et les propriétés, si $(M, +)$ est un groupe abélien additif, de l'application :

$$\begin{array}{l} \mathbb{Z} \times M \longrightarrow M \\ (n, x) \longmapsto n \cdot x \text{ (c'est-à-dire } n \text{ fois } x), \end{array}$$

loi de composition externe de domaine $\mathbb{Z}$ sur M qui fait clairement de M un module sur $\mathbb{Z}$; cette structure d'un groupe abélien additif est appelée sa structure canonique de module sur $\mathbb{Z}$. Les modules sur l'anneau $\mathbb{Z}$ des entiers relatifs sont donc tout simplement les groupes abéliens additifs avec la loi de composition externe de domaine $\mathbb{Z}$ définie ci-dessus.

Si M est un sous-groupe additif de $\mathbb{C}$ — en particulier $M = \mathbb{Z}, \mathbb{Q}, \mathbb{R}$ ou $\mathbb{C}$ — que l'on munit de sa structure canonique de module sur $\mathbb{Z}$, alors, pour tout entier relatif n et tout élément x de M , $n \cdot x$ est le produit « ordinaire » nx de n par x — il suffit de le prouver lorsque $n \in \mathbb{N}$, ce qui se fait par une récurrence immédiate.

Dans les exemples qui suivent, nous nous intéressons aux *parties* libres, liées ou génératrices, de préférence aux familles de vecteurs libres, liées ou génératrices. Les éventuelles bases d'un module sont donc ses parties libres et génératrices.

4.36. Module ne possédant pas de base.

Nous considérons un entier $p \geq 2$ et le groupe abélien additif $M = \mathbb{Z}/p\mathbb{Z}$, que nous munissons de sa structure canonique de module sur $\mathbb{Z}$. Pour tout $x \in M$, $p \cdot x = 0$ — en effet, si x est la classe de $n \in \mathbb{Z}$ modulo p , $p \cdot x$ est celle de pn . On en déduit que, pour tout vecteur a de M , la partie $\{a\}$ est liée, donc que toute partie non vide de M est liée. Le sous-module de M engendré par $\emptyset$ est $\{0\}$ et $M \neq \{0\}$, donc $\emptyset$ n'engendre pas M . En conclusion, ce module ne possède pas de base. $\square$

Dans l'exemple précédent 4.36, nous avons utilisé en fait un module de torsion, c'est-à-dire un module M tel que, pour tout vecteur a de M , la partie $\{a\}$ est liée : il existe un scalaire λ différent de zéro tel que $\lambda a = 0$.

DÉFINITION 4.7. — Un module sans torsion est un module M tel que, pour tout scalaire λ et tout vecteur x de M , $\lambda x = 0$ implique $\lambda = 0$ ou $x = 0$, c'est-à-dire un module dans lequel, pour tout vecteur x non nul, $\{x\}$ est une partie libre.

Il ne faut pas croire qu'un module sans torsion admette forcément une base.

4.37. Module sans torsion ne possédant pas de base.

Nous munissons $\mathbb{Q}$ de sa structure canonique de module sur $\mathbb{Z}$. Si n est un entier relatif (un scalaire) et r un nombre rationnel (un vecteur), $n \cdot r$ est le produit « ordinaire » nr , donc $n \cdot r = 0$ implique $n = 0$ ou $r = 0$. Par conséquent $\mathbb{Q}$ est un module sans torsion. Comme $\mathbb{Q} \neq \{0\}$, $\emptyset$ n'est pas une base de $\mathbb{Q}$.

Soit r un nombre rationnel. Nous choisissons un représentant (p, q) de r . On a, pour tout $n \in \mathbb{Z}$:

$$n \cdot r = nr = \frac{np}{q}.$$

Or il n'existe aucun entier relatif n tel que $\frac{1}{2q} = \frac{np}{q}$, donc $\{r\}$ n'engendre pas $\mathbb{Q}$.

Il en résulte que le module $\mathbb{Q}$ sur $\mathbb{Z}$ n'admet aucune base de cardinal 1.

Soit r et s des nombres rationnels distincts. Notons (p, q) un représentant de r et (m, n) un représentant de s . Alors $(mq)r = mp = (pn)s$, donc $(mq)r - (pn)s = 0$. Or $r \neq s$, donc m ou p est différent de zéro, ce qui montre que les scalaires mq et pn ne sont pas tous les deux nuls ; par conséquent, la partie $\{r, s\}$ est liée. Il en résulte que toute partie de $\mathbb{Q}$ contenant au moins deux éléments est liée.

En conclusion, $\mathbb{Q}$ n'admet pas de base ; ce n'est donc pas un module libre. $\square$

Dans un espace vectoriel, les bases sont les parties libres maximales et les parties génératrices minimales. Ceci devient faux dans un module.

4.38. Partie libre maximale qui n'est pas une base.

Nous reprenons le module $\mathbb{Q}$ sur $\mathbb{Z}$ étudié dans l'exemple précédent 4.37, dont nous utilisons les résultats, et nous choisissons un nombre rationnel r (un vecteur) différent de zéro. Alors $\{r\}$ est libre et $\{r\}$ n'engendre pas $\mathbb{Q}$, donc $\{r\}$ n'est pas une base de $\mathbb{Q}$. Or toute partie de $\mathbb{Q}$ de cardinal supérieur ou égal à 2 est liée, donc $\{r\}$ est une partie libre maximale du module $\mathbb{Q}$ sur $\mathbb{Z}$. $\square$

4.39. Partie génératrice minimale qui n'est pas une base.

Nous munissons $\mathbb{Z}$ de sa structure canonique de module sur lui-même. Alors, si n est un entier relatif (un scalaire) et p un entier relatif (un vecteur), $n \cdot p$ est le produit $n \times p$. On a, pour tout entier relatif n , $n = 3n - 2n = n \times 3 - n \times 2 = n \cdot 3 - n \cdot 2$, donc $\{2, 3\}$ est une partie génératrice de $\mathbb{Z}$. L'entier 2 ne divise pas 1 dans l'anneau $\mathbb{Z}$, donc $\{2\}$ n'engendre pas $\mathbb{Z}$; de même, $\{3\}$ n'est pas une partie génératrice de $\mathbb{Z}$. Enfin, $\mathbb{Z} \neq \{0\}$, donc $\emptyset$ n'engendre pas $\mathbb{Z}$. Par suite, la paire $\{2, 3\}$ est une partie génératrice minimale de $\mathbb{Z}$. Comme $0 = 3 \times 2 - 2 \times 3 = 3 \cdot 2 - 2 \cdot 3$ et que $3 \neq 0$, $\{2, 3\}$ n'est pas une partie libre de $\mathbb{Z}$, donc $\{2, 3\}$ n'est pas une base de $\mathbb{Z}$. $\square$

Dans un espace vectoriel de dimension finie E , un sous-espace vectoriel de même dimension que E , c'est-à-dire isomorphe à E , est égal à E . Ceci devient faux pour les sous-modules d'un module libre.

4.40. Module libre ayant un sous-module strict de même « dimension ».

Nous munissons, comme dans l'exemple précédent 4.39, $\mathbb{Z}$ de sa structure canonique de module sur lui-même. Alors $\mathbb{Z}$ est un module libre et $\{1\}$ une base de $\mathbb{Z}$. Nous considérons son sous-module $\mathbb{P} = 2\mathbb{Z}$, ensemble des nombre entiers relatifs pairs. Le singleton $\{2\}$ est clairement une base du module $\mathbb{P}$, donc l'application :

$$\left| \begin{array}{l} f : \mathbb{Z} \longrightarrow \mathbb{P} \\ n \longmapsto f(n) = 2n = n \times 2 = n \cdot 2 \end{array} \right.$$

est un isomorphisme de module de $\mathbb{Z}$ sur $\mathbb{P}$. Par suite, $\mathbb{Z}$ et $\mathbb{P}$ sont isomorphes en tant que modules¹⁰ sur $\mathbb{Z}$, alors que $\mathbb{P}$ est un sous-module strict de $\mathbb{Z}$. $\square$

10. Remarquons que $\mathbb{Z}$ est un anneau, mais que $\mathbb{P}$ n'est pas un sous-anneau de $\mathbb{Z}$.

4.41. Module libre ayant au moins un sous-module qui n'est pas un module libre.

Nous munissons l'anneau $\mathbb{Z}[X]$ des polynômes à coefficients entiers de sa structure de module sur lui-même. Alors $\mathbb{Z}[X]$ est un module libre et ses sous-modules sont ses idéaux. Si $P, Q \in \mathbb{Z}[X]$ et si $P \neq Q$, $PQ + (-Q)P = 0$ et P ou Q est différent de zéro, donc la paire $\{P, Q\}$ est liée. Ceci montre que toute partie de $\mathbb{Z}[X]$ contenant au moins deux éléments est liée. Par conséquent, si $\mathcal{M}$ est un sous-module libre différent de $\{0\}$ de $\mathbb{Z}[X]$, alors $\mathcal{M}$ possède une base de cardinal 1, donc $\mathcal{M}$ est un idéal principal de l'anneau $\mathbb{Z}[X]$. Or l'idéal $\mathcal{J}$ de $\mathbb{Z}[X]$ engendré par $\{2X^0, X\}$ n'est pas un idéal principal (exemple 3.29, page 49), donc $\mathcal{J}$ n'est pas un sous-module libre du module $\mathbb{Z}[X]$ sur lui-même. $\square$

Dans un espace vectoriel de dimension finie, tout sous-espace vectoriel possède au moins un supplémentaire, ce qui est encore vrai en dimension infinie si l'on accepte l'axiome du choix. Ceci est faux, en général, dans un module.

4.42. Sous-module sans supplémentaire.

Nous munissons $\mathbb{Z}$ de sa structure de module sur lui-même. Supposons que le sous-module $M = 2\mathbb{Z}$ de $\mathbb{Z}$ possède un supplémentaire N dans $\mathbb{Z}$. Comme N est un sous-module, c'est un sous-groupe de $(\mathbb{Z}, +)$, donc il existe un entier naturel n tel que $N = n\mathbb{Z}$; de plus, $\mathbb{Z} = M \oplus N$ et $M \subsetneq \mathbb{Z}$, donc $N \neq \{0\}$, ce qui montre que $n \geq 1$. Alors $2n \in M \cap N$ et $2n \neq 0$, en contradiction avec $M \cap N = \{0\}$. Ainsi le sous-module $2\mathbb{Z}$ de $\mathbb{Z}$ n'admet pas de supplémentaire dans le module $\mathbb{Z}$. $\square$

■ Formes bilinéaires symétriques et formes quadratiques

DÉFINITION 4.8. — Si E est un espace vectoriel, une forme bilinéaire symétrique sur E est une application f de $E \times E$ dans son corps de base K telle que :

- (I) pour tout vecteur x de E , $y \mapsto f(x, y)$ est une forme linéaire sur E ,
- (II) quels que soient les vecteurs x et y de E , $f(y, x) = f(x, y)$,

et l'application $q : x \mapsto q(x) = f(x, x)$ de E dans le corps K s'appelle la forme quadratique sur E associée à la forme bilinéaire symétrique f .

L'assertion (II) exprime que l'application f est symétrique, et il résulte alors de l'assertion (I) que f est bien une application bilinéaire de $E \times E$ dans K .

THÉORÈME 4.12. — Si K n'est pas de caractéristique 2, une application q d'un espace vectoriel E dans K est une forme quadratique sur E si, et seulement si, les deux assertions suivantes sont vérifiées :

- (I) L'application $(x, y) \mapsto \frac{1}{2}(q(x+y) - q(x) - q(y))$ de $E \times E$ dans K est une forme bilinéaire symétrique sur E .
- (II) Pour tout vecteur x de E , $q(2x) = 4q(x)$.

L'assertion (II) du théorème 4.12 est évidemment nécessaire pour obtenir une forme quadratique.

4.43. Application vérifiant l'assertion (i) du théorème 4.12 qui n'est pas une forme quadratique.

Il suffit d'ajouter une forme linéaire non nulle à une forme quadratique donnée. Par exemple l'application $q : x = (x_1, x_2) \mapsto q(x) = (x_1)^2 + (x_2)^2 + x_1 + x_2$ de $E = \mathbb{R}^2$ dans $K = \mathbb{R}$ vérifie clairement l'assertion (i) du théorème 4.12, mais pas la proposition (ii), puisque $q(2(1, 0)) = q((2, 0)) = 6 \neq 8 = 4q((1, 0))$. $\square$

DÉFINITION 4.9. — Une forme bilinéaire f sur un espace vectoriel E est alternée si $f(x, x) = 0$ pour tout vecteur x de E et antisymétrique si, quels que soient les vecteurs x et y de E , $f(y, x) = -f(x, y)$.

THÉORÈME 4.13. — Une forme bilinéaire alternée est antisymétrique et la réciproque est vraie si le corps de base K n'est pas de caractéristique 2.

4.44. Forme bilinéaire antisymétrique qui n'est pas alternée.

Nous considérons l'espace vectoriel $E = K^2$ sur le corps $K = \mathbb{F}_2 = \mathbb{Z}/2\mathbb{Z}$ des entiers modulo 2 et l'application f de $E \times E$ dans K définie de la manière suivante : si $x = (x_1, x_2)$ et $y = (y_1, y_2)$ sont des vecteurs de E , $f(x, y) = x_1y_1 + x_2y_2$. Il est clair que f est une forme bilinéaire symétrique sur E .

Dans K , tout élément est égal à son opposé, donc, quels que soient les vecteurs x et y de E , $f(y, x) = f(x, y) = -f(x, y)$. Il en résulte que f est une forme bilinéaire antisymétrique¹¹ sur l'espace vectoriel E .

On a, pour tout vecteur $x = (x_1, x_2)$ de E , $f(x, x) = (x_1)^2 + (x_2)^2$. En particulier, pour le vecteur $z = (1, 0)$, $f(z, z) = 1^2 + 0^2 = 1 \neq 0$. Ceci montre que la forme bilinéaire f n'est pas alternée. $\square$

DÉFINITION 4.10. — Si f est une forme bilinéaire symétrique sur un espace vectoriel E , on associe à tout vecteur x de E la forme linéaire $f_x : y \mapsto f_x(y) = f(x, y)$ sur E , le noyau de f , noté $\text{Ker}(f)$, est le noyau de l'application linéaire :

$$\left| \begin{array}{l} \Phi_f : E \longrightarrow \mathcal{L}(E, K) \\ x \longmapsto \Phi_f(x) = f_x \end{array} \right.$$

de E dans le dual $\mathcal{L}(E, K)$ de E , la forme bilinéaire symétrique f est dégénérée si $\text{Ker}(f) \neq \{0\}$ et non dégénérée si $\text{Ker}(f) = \{0\}$, et le cône isotrope de la forme quadratique q associée à f est $C(q) = \{x \in E \mid q(x) = 0\}$.

Le noyau d'une forme bilinéaire symétrique f sur un espace vectoriel E est donc $\text{Ker}(f) = \{x \in E \mid f(x, y) = 0 \text{ pour tout } y \in E\}$ et, de manière évidente, le noyau de f est inclus dans le cône isotrope de la forme quadratique associée.

4.45. Cône isotrope différent du noyau.

Nous considérons l'espace vectoriel $E = K^2$ sur le corps $K = \mathbb{R}$ et l'application f de $E \times E$ dans K définie de la manière suivante : si $x = (x_1, x_2)$ et $y = (y_1, y_2)$ sont des vecteurs de E :

$$f(x, y) = \frac{1}{2}(x_1y_2 + x_2y_1).$$

11. Sur un corps de caractéristique 2, les formes symétriques et antisymétriques coïncident.

Clairement, f est une forme bilinéaire symétrique sur E et sa forme quadratique associée est l'application $q : x = (x_1, x_2) \mapsto q(x) = x_1x_2$ de E dans K . Le cône isotrope $C(q)$ de q est donc la réunion des droites vectorielles D_1 engendrée par $e_1 = (1, 0)$ et D_2 engendrée par $e_2 = (0, 1)$.

Si un vecteur $x = (x_1, x_2)$ de E appartient au noyau de f , on a en particulier, en posant $y = (x_2, x_1)$, $0 = f(x, y) = (x_1)^2 + (x_2)^2$, donc $x = (0, 0)$. De plus, $(0, 0)$ appartient à $\text{Ker}(f)$, donc le noyau de f est $\{(0, 0)\} \subsetneq C(q)$. $\square$

4.46. Forme bilinéaire symétrique non dégénérée telle que la restriction à un sous-espace vectoriel différent de $\{0\}$ de la forme quadratique associée est nulle.

Nous considérons l'espace vectoriel $E = K^4$ sur le corps $K = \mathbb{R}$ et l'application f de $E \times E$ dans K définie comme suit : si $x = (x_1, x_2, x_3, x_4)$ et $y = (y_1, y_2, y_3, y_4)$ sont des vecteurs de E , $f(x, y) = x_1y_1 + x_2y_2 - x_3y_3 - x_4y_4$. Clairement, f est une forme bilinéaire symétrique sur E . Sa forme quadratique associée est l'application :

$$\left| \begin{array}{l} q : E \longrightarrow K \\ x = (x_1, x_2, x_3, x_4) \longmapsto q(x) = (x_1)^2 + (x_2)^2 - (x_3)^2 - (x_4)^2. \end{array} \right.$$

Nous notons enfin $\mathcal{B} = (e_1, e_2, e_3, e_4)$ la base canonique de $E = K^4$.

Si $x = (x_1, x_2, x_3, x_4)$ est un vecteur de E , la décomposition de la forme linéaire $f_x : y \mapsto f_x(y) = f(x, y)$ sur E dans la base duale $\mathcal{B}^* = (e^{*1}, e^{*2}, e^{*3}, e^{*4})$ de la base $\mathcal{B}$ s'écrit $f_x = x_1e^{*1} + x_2e^{*2} - x_3e^{*3} - x_4e^{*4}$, donc $f_x = 0$ si, et seulement si, $x = (0, 0, 0, 0)$, ce qui montre que f n'est pas dégénérée.

Nous posons $a = e_1 + e_3 = (1, 0, 1, 0)$ et $b = e_2 + e_4 = (0, 1, 0, 1)$, et nous notons F le sous-espace vectoriel de E engendré par $\{a, b\}$, de dimension 2 puisque la famille (a, b) est libre. Le calcul montre que $q(a) = q(b) = 0$ et $f(a, b) = 0$, donc, par bilinéarité, la restriction de q à F est nulle. $\square$

Ceci devient impossible si E est de dimension finie et si la dimension de F est strictement plus grande que la moitié de $\dim E$, ou si la forme bilinéaire symétrique f est un produit scalaire sur un espace vectoriel réel de dimension finie¹² — qui devient ainsi un espace vectoriel euclidien (définition 17.13, page 336).

DÉFINITION 4.11. — Soit f une forme bilinéaire symétrique sur un espace vectoriel E .

Des vecteurs x et y de E sont orthogonaux relativement à f si $f(x, y) = 0$, et l'orthogonal d'une partie F de E relativement à f est l'ensemble des vecteurs de E orthogonaux relativement à f à tous les vecteurs de F , c'est-à-dire l'ensemble des vecteurs x de E tels que $f(x, z) = 0$ pour tout vecteur z de F .

Si f est une forme bilinéaire symétrique sur un espace vectoriel E de dimension finie $n \geq 1$ et q la forme quadratique associée, une base orthogonale pour f (ou pour q) est une base $\mathcal{B} = (b_1, \dots, b_n)$ de E dont les vecteurs sont deux à deux orthogonaux relativement à f — ce qui signifie que, quels que soient les indices distincts $i, j \in [1, n]$, $f(b_i, b_j) = 0$ — et si $\mathcal{B} = (b_1, \dots, b_n)$ est une base orthogonale

12. Plus précisément, pour une forme quadratique q non dégénérée de signature (p_1, p_2) les sous-espaces vectoriels F de E pour lesquels la restriction de q à F est nulle sont de dimension maximum $\text{Min}(p_1, p_2)$.

pour f , alors, en posant $\beta_i = f(b_i, b_i)$ pour tout $i \in \llbracket 1, n \rrbracket$, on a, quels que soient les vecteurs $x = x_1 b_1 + \cdots + x_n b_n$ et $y = y_1 b_1 + \cdots + y_n b_n$ de E :

$$f(x, y) = \sum_{i=1}^n \beta_i x_i y_i \text{ et } q(x) = \sum_{i=1}^n \beta_i (x_i)^2.$$

Lorsque qu'une forme bilinéaire symétrique f sur un espace vectoriel réel E de dimension finie est un produit scalaire, alors, pour tout sous-espace vectoriel F de E , l'orthogonal $F^\perp$ de F relativement à f est un supplémentaire¹³ de F dans E . Ceci devient faux pour une forme bilinéaire symétrique quelconque.

4.47. Sous-espace vectoriel différent de $\{0\}$ égal à son orthogonal.

Nous reprenons les hypothèses, définitions et notations de l'exemple précédent 4.46. En raison de la bilinéarité de f , un vecteur $x = (x_1, x_2, x_3, x_4)$ de E appartient à l'orthogonal de F relativement à f si, et seulement si, $f(x, a) = f(x, b) = 0$, ce qui équivaut à $x_2 = x_4$ et $x_1 = x_3$, ou encore à l'appartenance de x à F . Ainsi, F est son propre orthogonal relativement à f . $\square$

Soit E un espace vectoriel de dimension finie $n \geq 1$. Diagonaliser une forme bilinéaire symétrique f sur E ou sa forme quadratique associée q , c'est déterminer, s'il en existe, une base de E orthogonale pour f (ou q). Tenter la diagonalisation simultanée de deux formes quadratiques sur E , c'est chercher une éventuelle base de E orthogonale simultanément pour les deux formes. Ce problème de diagonalisation simultanée se résout pour $\mathbb{R}$ si l'une des deux formes provient d'un produit scalaire, résultat qui ne se généralise pas pour des formes quelconques.

4.48. Deux formes quadratiques qui ne sont pas simultanément diagonalisables.

Nous considérons l'espace vectoriel $E = K^2$ sur le corps $K = \mathbb{R}$ et les applications f et f' de $E \times E$ dans K définies de la manière suivante : si $x = (x_1, x_2)$ et $y = (y_1, y_2)$ sont des vecteurs de E , $f(x, y) = x_1 y_1 - x_2 y_2$ et $f'(x, y) = x_1 y_2 + x_2 y_1$. Ce sont clairement des formes bilinéaires symétriques sur E et, si l'on note respectivement q et q' leurs formes quadratiques associées, on a, pour tout vecteur $x = (x_1, x_2)$ de E , $q(x) = (x_1)^2 - (x_2)^2$ et $q'(x) = 2x_1 x_2$.

Soit $a = (a_1, a_2)$ et $b = (b_1, b_2)$ des vecteurs de E tels que $f(a, b) = f'(a, b) = 0$. On a alors :

$$(1) \quad a_1 b_1 = a_2 b_2 \quad \text{et} \quad (2) \quad a_1 b_2 + a_2 b_1 = 0.$$

En multipliant (2) par b_1 , on obtient : $a_1 b_1 b_2 + a_2 (b_1)^2 = 0$, donc, en remplaçant $a_1 b_1$ par $a_2 b_2$ grâce à l'égalité (1), il vient : $a_2 (b_2)^2 + a_2 (b_1)^2 = 0$, d'où l'on déduit que $a_2 ((b_1)^2 + (b_2)^2) = 0$, ce qui entraîne $a_2 = 0$ ou $b = (0, 0)$. Si $a_2 = 0$, (1) et (2) donnent $a_1 b_1 = a_2 b_2 = 0$, donc $a_1 = 0$ ou $b = (0, 0)$. On en déduit que l'un des deux vecteurs a ou b est nul.

Cette étude montre que si des vecteurs a et b sont tous deux différents de zéro, on ne peut avoir à la fois $f(a, b) = 0$ et $f'(a, b) = 0$, ce qui montre qu'il n'existe aucune base de E orthogonale simultanément pour q et pour q' . $\square$

13. Voir les exemples 17.23 à 17.26, pages 334, 335 et 336.

Chapitre 5

Nombres réels

La notion intuitive que nous avons de l'ensemble $\mathbb{R}$ des nombres réels n'est pas seulement ensembliste : $]0, 1[$, $[0, 1]$, $\mathbb{R}$ et même $\mathbb{R}^2$ ont le même cardinal, ce qui signifie qu'il existe une bijection de l'un sur l'autre, généralisation aux ensembles infinis du fait, pour les ensembles finis, d'avoir le même nombre d'éléments. Pourtant nous nous en faisons une idée différente. Ceci est dû au fait que l'on a présentes à l'esprit, en même temps, la structure d'ensemble ordonné de $\mathbb{R}$ et sa topologie usuelle compatible avec celle-là. Approcher un point de $\mathbb{R}$ a quelque chose de linéaire, c'est une approche ordonnée ; ceci n'est pas le cas dans $\mathbb{R}^2$.

Malgré une utilisation très ancienne des nombres réels, il faut attendre le XIX^e siècle pour qu'ils soient définis de manière rigoureuse. Pendant longtemps, on s'est contenté de justifications intuitives en évoquant «l'évidence géométrique». Le besoin de définir précisément les notions de continuité et de limite apparaît vers 1820 avec Bernhart Bolzano et Augustin-Louis Cauchy, et se développe avec Karl Weierstrass vers 1850. Ce dernier propose en 1863 la première construction des nombres réels. Il ne la publie qu'en 1872, après que Charles Méray et Georg Cantor (suites de Cauchy de nombres rationnels) et Richard Dedekind (coupures dans l'ensemble des nombres rationnels) en aient détaillé d'autres.

■ Ecriture décimale des nombres réels

5.1. Nombre possédant deux écritures décimales.

En posant $a = 0,999999\dots$, on a :

$$a = \sum_{n=1}^{+\infty} \frac{9}{10^n} = \frac{9}{10} \frac{1}{1 - \frac{1}{10}} = 1$$

et comme $1 = 1,00000000\dots$, nous obtenons ainsi deux développements décimaux pour le nombre réel 1. $\square$

Plus généralement tous les nombres décimaux, et seulement eux, possèdent deux développements décimaux ; dans le premier, appelé le développement décimal propre, les décimales sont nulles à partir d'un certain rang, dans le second les décimales sont égales à 9 à partir d'un certain rang et on passe du premier au second en enlevant une unité au dernier chiffre non nul et en le faisant suivre d'une suite infinie de 9.

5.2. Nombre ayant dans ses décimales toutes les séquences de chiffres possibles.

Nous posons $\gamma = 0,12345678910111213\dots$, nombre dont la suite des décimales est obtenue en concaténant tous les entiers naturels non nuls écrits en base 10 en ordre croissant ; γ s'appelle le nombre de Champernowne¹. Toute séquence de chiffres est clairement dans ce développement décimal car elle représente elle-même un nombre. En particulier, n'importe quelle œuvre littéraire codée en chiffre par quelque méthode que ce soit, se trouve écrite dans les décimales de γ . $\square$

■ Différents ensembles de nombres réels

L'ensemble $\mathbb{Q}$ des nombres rationnels — les « fractions » — est strictement inclus dans $\mathbb{R}$; en effet, le nombre réel e , base des logarithmes népériens, est irrationnel². Entre les deux se situent en particulier deux ensembles de nombres, l'ensemble $\mathbb{A}$ des nombres algébriques et l'ensemble Γ des nombres constructibles, dont nous rappelons les définitions ; pour celle des nombres constructibles, on identifie $\mathbb{R}^2$ « au plan », ce qui signifie que $\mathbb{R}^2$ est muni de sa structure canonique de plan affine euclidien et de son repère canonique, qui est un repère orthonormal.

DÉFINITION 5.1. — Un nombre réel est algébrique s'il est racine d'un polynôme non nul à coefficients entiers et transcendant dans le cas contraire. Si x est un nombre réel algébrique, le polynôme minimal de x est le générateur unitaire de l'idéal de l'anneau $\mathbb{Q}[X]$ constitué des polynômes P à coefficients rationnels tels que $P(x) = 0$, et le degré de x est celui de son polynôme minimal.

DÉFINITION 5.2. — Un nombre réel x est constructible si $X = (x, 0)$ est un point constructible, c'est-à-dire un point de $\mathbb{R}^2$ que l'on peut construire à la règle et au compas à partir des points $O = (0, 0)$ et $I = (1, 0)$.

Les ensembles $\mathbb{Q}$, Γ et $\mathbb{A}$ sont des sous-corps du corps $\mathbb{R}$ des nombres réels et nous avons les inclusions strictes $\mathbb{Q} \subsetneq \Gamma \subsetneq \mathbb{A} \subsetneq \mathbb{R}$.

5.3. Nombre réel constructible qui n'est pas rationnel.

Le nombre réel $\sqrt{2}$, racine du polynôme $X^2 - 2$, est un nombre algébrique.

Nous prouvons que $\sqrt{2}$ est irrationnel. Il s'agit de montrer que $\sqrt{2}$ n'appartient pas à $\mathbb{Q}$, donc qu'aucun nombre rationnel n'a pour carré 2. Supposons qu'il existe un rationnel r tel que $r^2 = 2$. Nous choisissons un représentant irréductible (p, q) de $r : p$ et q sont des entiers relatifs non nuls premiers entre eux et $r = p/q$. Alors, dans l'anneau $\mathbb{Z}$, $p^2 = 2q^2$, donc 2 divise p^2 , ce qui, puisque 2 est un nombre premier, montre que 2 divise p , d'où l'existence d'un entier p' tel que $p = 2p'$. On en déduit que $4p'^2 = 2q^2$, donc $2p'^2 = q^2$. Par suite 2 divise q^2 , donc aussi q , en contradiction avec le fait que p et q sont premiers entre eux.

1. Le nombre réel γ est introduit en 1933 par l'étudiant anglais, devenu économiste, David Champernowne (1912-2000) et son compatriote Kurt Mahler (1903-1988) prouve en 1961 que γ est transcendant — voir la définition 5.1.

2. Comme le prouve vers 1737 le génial mathématicien suisse Leonhard Euler (1707-1783).

Nous montrons que $\sqrt{2}$ est constructible. Nous construisons grâce au compas la symétrique $I' = (-1, 0)$ de $I = (1, 0)$ par rapport à $O = (0, 0)$, nous traçons les deux cercles de centres respectifs I et I' et de rayon $II' (= 2)$, qui se coupent en deux points B et B' situés sur la perpendiculaire à (OI) passant par O — c'est la droite (Oy) du repère —, nous reportons à l'aide du compas la longueur $1 = OI$ à partir de O sur la droite $(OB) = (Oy)$, nous notons C l'un des deux points obtenus — alors $C = (0, 1)$ ou $C = (0, -1)$ — et nous construisons enfin le point D en reportant, grâce au compas, la longueur IC sur la droite $(OI) = (Ox)$, vers I et à partir de O ; comme $IC = \sqrt{2}$ (diagonale d'un carré de côté 1), $D = (\sqrt{2}, 0)$, donc $\sqrt{2}$ est constructible. En conclusion, $\sqrt{2}$ est constructible et irrationnel. $\square$

On dispose du théorème suivant³.

THÉORÈME 5.1. — Un nombre réel constructible est algébrique et son degré est une puissance de 2.

5.4. Nombre réel algébrique qui n'est pas constructible.

Le nombre $\alpha = \sqrt[3]{2}$ est algébrique comme racine du polynôme $X^3 - 2$. Il n'existe aucun rationnel r tel que $r^3 = 2$ — voir l'exemple 3.18 (page 41). Ainsi α est irrationnel et $X^3 - 2$ n'a pas de racine rationnelle. Le polynôme $X^3 - 2$, de degré 3, est irréductible sur $\mathbb{Q}$ (théorème 3.4, page 40), donc le nombre algébrique α est de degré 3. Le théorème 5.1 montre alors que α n'est pas constructible⁴. $\square$

5.5. Nombre réel qui n'est pas algébrique.

Nous prouvons que le nombre réel e n'est pas algébrique, c'est-à-dire que e est un nombre transcendant⁵.

Supposons e algébrique. Il existe un entier naturel q et des nombres entiers relatifs $b_0, b_1, \dots, b_q$ non tous nuls tels que $b_0 + b_1 e + \dots + b_q e^q = 0$. En notant k le plus petit et ℓ le plus grand des indices i tels que $b_i \neq 0$, et en divisant cette égalité par e^k , elle s'écrit $a_0 + a_1 e + \dots + a_n e^n = 0$ où n est un entier naturel et $a_0, a_1, \dots, a_n$ des entiers relatifs tels que $a_0 \neq 0$ et $a_n \neq 0$ — le réel e étant irrationnel, n est supérieur ou égal à 2. Nous posons $R = a_0 + a_1 X + \dots + a_n X^n$. Nous associons à tout entier $p \geq 2$ le polynôme de degré $m = p - 1 + np$:

$$P = \frac{1}{(p-1)!} X^{p-1} \prod_{k=1}^n (X - k)^p.$$

Soit p un entier supérieur ou égal à 2.

Nous posons, pour tout réel α et tout polynôme A , $I_\alpha(A) = \int_0^1 \alpha e^{-\alpha x} A(\alpha x) dx$.

Soit α un nombre réel. Si A est un polynôme, une intégration par parties donne :

$$I_\alpha(A) = [-e^{-\alpha x} A(\alpha x)]_{x=0}^{x=1} + \int_0^1 \alpha e^{-\alpha x} A'(\alpha x) dx = [-e^{-\alpha x} A(\alpha x)]_{x=0}^{x=1} + I_\alpha(A').$$

3. Voir [BOUA], chapitre 8, §8.2.2, théorème 8.1.

4. Ceci démontre l'impossibilité du problème posé par les Grecs de la duplication du cube.

5. Ce résultat est établi en 1873 par le mathématicien français Charles Hermite (1822-1901). Le mathématicien allemand Carl Lindemann (1852-1939) prouve en 1882 que π est transcendant, ce qui démontre l'impossibilité de la quadrature du cercle, mettant ainsi fin à des siècles de recherches.

En appliquant ceci aux polynômes $P, P', P'' \dots$ et en remarquant que $P^{(m+1)} = 0$, on obtient de proche en proche :

$$I_\alpha(P) = \sum_{i=0}^m [-e^{-\alpha x} P^{(i)}(\alpha x)]_{x=0}^{x=1} = [-e^{-\alpha x} Q(\alpha x)]_{x=0}^{x=1} \text{ où } Q = \sum_{i=0}^m P^{(i)}.$$

On a $I_\alpha(P) = Q(0) - e^{-\alpha} Q(\alpha)$, donc $e^\alpha Q(0) = Q(\alpha) + e^\alpha I_\alpha(P)$.

Par conséquent $e^k Q(0) = Q(k) + e^k I_k(P)$ pour tout $k \in \llbracket 1, n \rrbracket$ et, comme $R(e) = 0$:

$$0 = Q(0)R(e) = a_0 Q(0) + \sum_{k=1}^n a_k e^k Q(0) = a_0 Q(0) + \sum_{k=1}^n a_k Q(k) + \sum_{k=1}^n a_k e^k I_k(P).$$

Nous posons $N_p = a_0 Q(0) + \sum_{k=1}^n a_k Q(k)$, ou encore $N_p = - \sum_{k=1}^n a_k e^k I_k(P)$.

On a, pour tout entier k compris entre 1 et n , $Q(k) = \sum_{i=0}^m P^{(i)}(k)$.

Pour tout $k \in \llbracket 1, n \rrbracket$, k est une racine d'ordre p de P donc $P^{(i)}(k) = 0$ pour tout $i \in \llbracket 0, p-1 \rrbracket$. De même 0 est une racine d'ordre $p-1$ de P , donc $P^{(i)}(0) = 0$ pour tout $i \in \llbracket 0, p-2 \rrbracket$. On déduit de la formule de Leibniz, appliquée à l'ordre $p-1$ au produit $P = AB$ où :

$$A = \frac{1}{(p-1)!} X^{p-1} \text{ et } B = \prod_{k=1}^n (X - k)^p,$$

que $P^{(p-1)} = \sum_{i=0}^{p-1} \binom{p-1}{i} A^{(i)} B^{(p-1-i)}$. Or $A^{(p-1)} = 1$ et $A^{(i)}(0) = 0$ pour $i < p-1$, donc :

$$P^{(p-1)}(0) = B(0) = (-1)^{np} (n!)^p.$$

Le polynôme $(p-1)! P$ est à coefficients entiers, son degré est m et sa valuation $p-1$, donc :

$$P = \sum_{q=p-1}^m \frac{c_q}{(p-1)!} X^q$$

où $c_q \in \mathbb{Z}$ pour tout $q \in \llbracket p-1, m \rrbracket$. Pour tout entier $i \geq p$, le polynôme dérivé d'ordre i de X^q est nul pour $q < i$ donc :

$$P^{(i)} = \sum_{i \leq q \leq m} \frac{c_q}{(p-1)!} \frac{q!}{(q-i)!} X^{q-i} = \sum_{i \leq q \leq m} \frac{i!}{(p-1)!} \binom{q}{i} c_q X^{q-i} = \sum_{i \leq q \leq m} p \frac{i!}{p!} \binom{q}{i} c_q X^{q-i}$$

et comme $i!/p!$ est un nombre entier, le polynôme $P^{(i)}$ est à coefficients entiers divisibles par p .

En conclusion de cette étude, nous voyons qu'il existe un nombre entier relatif ℓ_p tel que $N_p = (-1)^{np} a_0 (n!)^p + p \ell_p$. En particulier, N_p est un nombre entier.

Nous étudions la suite $(N_p)_{p \geq 2}$.

Soit p un entier ≥ 2 . On a $N_p = - \sum_{k=1}^n a_k e^k I_k(P)$.

Soit $t \in [0, n]$. Alors $|t| = t \leq n$ et, pour tout $i \in \llbracket 1, n \rrbracket$, $0 \leq t \leq n$ et $-n \leq -i \leq 0$, donc $|t-i| \leq n$, d'où l'on déduit que :

$$|P(t)| = \frac{1}{(p-1)!} |t|^{p-1} |t-1|^p \dots |t-n|^p \leq \frac{1}{(p-1)!} n^{p-1} \underbrace{n^p \dots n^p}_{n \text{ facteurs}} = \frac{n^{p-1+np}}{(p-1)!}.$$

Il en résulte que, pour tout $k \in \llbracket 1, n \rrbracket$:

$$\begin{aligned} |I_k(P)| &= \left| \int_0^1 k e^{-kx} P(kx) dx \right| = \left| \int_0^k e^{-t} P(t) dt \right| \leq \int_0^k \underbrace{|e^{-t}|}_{=e^{-t} \leq 1} |P(t)| dt \\ &\leq \int_0^k |P(t)| dt \leq \int_0^n |P(t)| dt \leq n \frac{n^{p-1+n}}{(p-1)!} = \frac{n^{p+n}}{(p-1)!} = \frac{(n^{n+1})^p}{(p-1)!}, \end{aligned}$$

ce qui donne :

$$|N_p| \leq \sum_{k=1}^n |a_k| e^k |I_k(P)| \leq \underbrace{\left(\sum_{k=1}^n |a_k| e^k \right)}_{\text{constante}} \frac{(n^{n+1})^p}{(p-1)!}.$$

Pour tout nombre réel x , la suite $\left(\frac{x^q}{q!}\right)_{q \in \mathbb{N}}$ converge vers 0, donc la suite $(u_p)_{p \geq 2}$, de terme général :

$$u_p = \frac{(n^{n+1})^p}{(p-1)!} = n^{n+1} \times \frac{(n^{n+1})^{p-1}}{(p-1)!},$$

admet pour limite 0. Il en est donc de même de la suite $(N_p)_{p \geq 2}$; or c'est une suite d'entiers, donc il existe un entier $p_0 \geq 2$ tel que $N_p = 0$ pour tout entier $p \geq p_0$. Nous posons $q_0 = \text{Max}(p_0, |a_0|, n)$ et nous choisissons **un nombre premier** $p > q_0$ — un tel nombre premier p existe puisque l'ensemble des nombres premiers est infini. Comme $p > |a_0|$ et que $p > n$, p ne divise ni a_0 , ni aucun des entiers compris entre 1 et n , donc p ne divise pas $N_p = p \ell_p \pm a_0(n!)^p$, ce qui montre que N_p n'est pas nul, en contradiction avec $p > p_0$. En conclusion, e n'est pas algébrique, donc e est transcendant. $\square$

■ Ordre canonique de $\mathbb{R}$

Si une partie de $\mathbb{R}$ admet un plus grand (resp. un plus petit) élément m , m est sa borne supérieure (resp. inférieure) dans $\mathbb{R}$ (chapitre 1, au bas de la page 9).

5.6. Partie de $\mathbb{R}$ ayant une borne supérieure mais n'admettant pas de plus grand élément.

L'ensemble des majorants de $A = [0, 1[$ dans $\mathbb{R}$ est $[1, +\infty[$, donc 1 est la borne supérieure de A dans $\mathbb{R}$, alors que A ne possède pas de plus grand élément. $\square$

Dans $\mathbb{R}$, toute partie non vide majorée possède une borne supérieure — c'est la propriété fondamentale de l'ordre « canonique » de $\mathbb{R}$, celle qui permet par exemple de prouver que toute suite croissante et majorée de nombre réels converge dans $\mathbb{R}$, le théorème des valeurs intermédiaires, etc., en fait tous les théorèmes essentiels de l'analyse réelle. Cette propriété est fausse dans un ensemble totalement ordonné quelconque, par exemple $\mathbb{Q}$.

5.7. Partie non vide majorée de $\mathbb{Q}$ sans borne supérieure.

Nous munissons $\mathbb{Q}$ de l'ordre habituel et nous notons A l'ensemble des rationnels positifs r tels que $r^2 \leq 2$. Soit a un réel positif tel que $a < \sqrt{2}$. L'intervalle $]a, \sqrt{2}[$ contient au moins un rationnel r , on a $a^2 < r^2 < 2$, donc r appartient à A et

$r > a$, ce qui montre que a ne majore pas A . Nous en déduisons que tout majorant de A dans $\mathbb{R}$ est supérieur ou égal à $\sqrt{2}$, donc que tout majorant de A dans $\mathbb{Q}$ est strictement plus grand que $\sqrt{2}$ — en effet, 2 n'est le carré d'aucun nombre rationnel (exemple 5.3, page 84). Soit b un majorant de A dans $\mathbb{Q}$. Alors $\sqrt{2} < b$, ce qui justifie le choix d'un rationnel c tel que $\sqrt{2} < c < b$. On a, pour tout élément r de A , $r^2 \leq 2 < c^2$ donc $r < c$. Par conséquent c majore A dans $\mathbb{Q}$; de plus $c < b$. Ainsi l'ensemble des majorants de A dans $\mathbb{Q}$ n'admet pas de plus petit élément, donc A ne possède pas de borne supérieure dans $\mathbb{Q}$. $\square$

La borne supérieure de la réunion de deux parties non vides et majorées de $\mathbb{R}$ est la plus grande des bornes supérieures de ces deux parties. On ne dispose pas d'un résultat analogue pour l'intersection.

5.8. Parties non vides majorées de $\mathbb{R}$ dont la borne supérieure de l'intersection diffère de la plus petite des bornes supérieures.

Nous posons $A = \{0\} \cup \left\{ \frac{1}{n} \mid n \in \mathbb{N}^* \right\}$ et $B = \{0\} \cup \left\{ \frac{\pi}{n} \mid n \in \mathbb{N}^* \right\}$.

La borne supérieure de A est son plus grand élément 1 et celle de B est son plus grand élément π (obtenus pour $n = 1$). Comme π est irrationnel, $A \cap B = \{0\}$ donc la borne supérieure de $A \cap B$ est $0 \neq 1 = \min(1, \pi)$. $\square$

Pour des applications bornées f et g d'une partie A de $\mathbb{R}$ dans $\mathbb{R}$, on a :

$$\sup_{x \in A} (f(x) + g(x)) \leq \sup_{x \in A} f(x) + \sup_{x \in A} g(x).$$

Cette inégalité peut être stricte.

5.9. Fonctions f et g bornées sur une partie A de $\mathbb{R}$ telles que :

$$\sup_{x \in A} (f(x) + g(x)) < \sup_{x \in A} f(x) + \sup_{x \in A} g(x).$$

Nous définissons sur $A = [0, 1]$ la fonction f par $f(x) = 1 - x$ et la fonction g par $g(x) = x$. La borne supérieure de f sur A est égale à 1, ainsi que celle de g , donc la somme de ces deux bornes supérieures vaut 2. Pour tout point x de A , $f(x) + g(x) = 1$, donc la borne supérieure de $f + g$ sur A est égale à $1 < 2$. $\square$

■ Topologie de $\mathbb{R}$

DÉFINITION 5.3. — Une partie O de $\mathbb{R}$ est un ouvert de $\mathbb{R}$ si, pour tout élément x de O , il existe un réel $\varepsilon > 0$ tel que $]x - \varepsilon, x + \varepsilon[\subset O$.

On dit qu'une partie de $\mathbb{R}$ est ouverte si c'est un ouvert de $\mathbb{R}$. L'ensemble vide et les intervalles ouverts sont des ouverts de $\mathbb{R}$, toute réunion d'ouverts de $\mathbb{R}$ est ouverte et l'intersection d'une famille *finie* d'ouverts est ouverte. Ce dernier résultat ne se généralise pas à une intersection quelconque.

5.10. Intersection d'ouverts de $\mathbb{R}$ qui n'est pas ouverte.

Nous considérons la suite $(I_n)_{n \in \mathbb{N}^*}$ d'ouverts de $\mathbb{R}$, de terme général $I_n = \left] -\frac{1}{n}, \frac{1}{n} \right[$.

Pour tout réel $x \neq 0$, il existe au moins un entier $n \geq 1$ tel que $1/n < |x|$, donc l'intersection de cette suite est $\{0\}$, et $\{0\}$ n'est pas un ouvert de $\mathbb{R}$. $\square$

DÉFINITION 5.4. — Une partie de $\mathbb{R}$ est un fermé de $\mathbb{R}$ si son complémentaire dans $\mathbb{R}$ est un ouvert de $\mathbb{R}$.

On dit qu'une partie de $\mathbb{R}$ est fermée si c'est un fermé de $\mathbb{R}$. Pour toute partie X de $\mathbb{R}$, $\mathcal{C}_{\mathbb{R}}(\mathcal{C}_E X) = X$, donc une partie de $\mathbb{R}$ est fermée (resp. ouverte) si, et seulement si, son complémentaire est un ouvert (resp. un fermé) de $\mathbb{R}$. Les intervalles fermés sont les intervalles qui sont des fermés de $\mathbb{R}$; ce sont l'ensemble vide, $\mathbb{R}$, les segments et les intervalles $[a, +\infty[$ et $]-\infty, a]$ pour $a \in \mathbb{R}$. Une réunion *finie* de parties fermées de $\mathbb{R}$ est fermée; ceci devient faux pour une réunion quelconque.

5.11. Réunion de fermés de $\mathbb{R}$ qui n'est pas fermée.

Nous considérons la suite $(I_n)_{n \geq 2}$ de fermés de $\mathbb{R}$, de terme général $I_n = \left[\frac{1}{n}, 1 - \frac{1}{n}\right]$. La réunion de cette suite est l'intervalle $]0, 1[$, qui n'est pas un fermé de $\mathbb{R}$. $\square$

L'intersection d'une suite de segments emboîtés n'est pas vide. Ceci devient faux pour des intervalles fermés quelconques.

5.12. Suite décroissante d'intervalles fermés non vides dont l'intersection est vide.

Nous définissons la suite $(I_n)_{n \in \mathbb{N}}$ d'intervalles fermés de $\mathbb{R}$ par $I_n = [n, +\infty[$. Pour tout entier naturel n , $I_n \neq \emptyset$ et $I_{n+1} \subset I_n$, donc c'est une suite décroissante d'intervalles fermés non vides, alors que son intersection est vide. $\square$

Si A est une partie de $\mathbb{R}$, la réunion de tous les ouverts inclus dans A est clairement le plus grand pour l'inclusion des ouverts inclus dans A .

DÉFINITION 5.5. — L'intérieur $\overset{\circ}{A}$ d'une partie A de $\mathbb{R}$ est le plus grand pour l'inclusion des ouverts de $\mathbb{R}$ inclus dans A .

Par exemple, si a et b sont des nombres réels et si $a < b$, l'intérieur de chacun des intervalles $]a, b[$, $[a, b[$, $]a, b]$ et $[a, b]$ est l'intervalle ouvert $]a, b[$.

DÉFINITION 5.6. — Si A est une partie de $\mathbb{R}$, un nombre réel x est intérieur à A s'il existe un réel $\varepsilon > 0$ tel que $]x - \varepsilon, x + \varepsilon[\subset A$.

L'intérieur d'une partie A de $\mathbb{R}$ est l'ensemble des nombres réels intérieurs à A et une partie de $\mathbb{R}$ est ouverte si, et seulement si, elle est égale à son intérieur.

Si A et B sont des parties de $\mathbb{R}$, on a l'égalité $\overset{\circ}{A \cap B} = \overset{\circ}{A} \cap \overset{\circ}{B}$.

Pour la réunion, seule l'inclusion $\overset{\circ}{A} \cup \overset{\circ}{B} \subset \overset{\circ}{A \cup B}$ subsiste.

5.13. Parties A et B de $\mathbb{R}$ telles que $\overset{\circ}{A} \cup \overset{\circ}{B} \neq \overset{\circ}{A \cup B}$.

Nous posons $A = [0, 1]$ et $B = [1, 2]$. Comme $A \cup B = [0, 2]$, on a :

$$\overset{\circ}{A} \cup \overset{\circ}{B} =]0, 1[\cup]1, 2[\neq]0, 2[= \overset{\circ}{A \cup B}. \quad \square$$

Dans l'exemple précédent 5.13, seul 1 est dans l'un des ensembles et pas dans l'autre. Voici un exemple où la différence est plus conséquente.

5.14. Autres parties A et B de $\mathbb{R}$ telles que $\overset{\circ}{A} \cup \overset{\circ}{B} \neq \overset{\circ}{A \cup B}$.

Nous posons $A = \mathbb{Q}$ et $B = \mathbb{R} \setminus \mathbb{Q}$: A est l'ensemble des nombres rationnels et B celui des nombres irrationnels. Si x est un nombre réel et ε un réel > 0 , l'intervalle ouvert $]x - \varepsilon, x + \varepsilon[$ contient des rationnels et des irrationnels, donc il n'est inclus ni dans A ni dans B . Il en résulte que l'intérieur de A est vide, ainsi que l'intérieur de B . De plus $A \cup B = \mathbb{R}$ donc :

$$\overset{\circ}{A} \cup \overset{\circ}{B} = \emptyset \text{ et } \overset{\circ}{A \cup B} = \mathbb{R}. \quad \square$$

Dans le cas d'une famille finie de parties de $\mathbb{R}$, l'intérieur de la réunion est la réunion des intérieurs. Ceci devient faux pour une réunion quelconque.

5.15. Intérieur d'une réunion différent de la réunion des intérieurs.

Nous définissons la suite $(A_n)_{n \in \mathbb{N}^*}$ de parties de $\mathbb{R}$ par $A_n = \left] -\frac{1}{n}, 1 \right[$.

Pour tout entier $n \geq 1$, A_n est un ouvert de $\mathbb{R}$ donc $\overset{\circ}{A_n} = A_n$. Il en résulte que :

$$\bigcap_{n \in \mathbb{N}^*} \overset{\circ}{A_n} = \bigcap_{n \in \mathbb{N}^*} A_n = [0, 1[$$

alors que l'intérieur de $\bigcap_{n \in \mathbb{N}^*} A_n$ est l'intervalle ouvert $]0, 1[$. $\square$

Si A est une partie de $\mathbb{R}$, l'intersection de tous les fermés de $\mathbb{R}$ contenant A est clairement le plus petit pour l'inclusion des fermés contenant A .

DÉFINITION 5.7. — L'adhérence $\bar{A}$ d'une partie A de $\mathbb{R}$ est le plus petit pour l'inclusion des fermés de $\mathbb{R}$ contenant A .

Par exemple, si $a, b \in \mathbb{R}$ et $a < b$, l'adhérence de chacun des intervalles $]a, b[$, $[a, b[$, $]a, b]$ et $[a, b]$ est le segment $[a, b]$.

DÉFINITION 5.8. — Si A est une partie de $\mathbb{R}$, un nombre réel x est adhérent à A si, pour tout réel $\varepsilon > 0$, $]x - \varepsilon, x + \varepsilon[$ rencontre A .

L'adhérence d'une partie A de $\mathbb{R}$ est l'ensemble des réels adhérents à A . Une partie de $\mathbb{R}$ est fermée si, et seulement si, elle est égale à son adhérence. Si A est une partie de $\mathbb{R}$, un réel a est adhérent à A si, et seulement si, a est la limite dans $\mathbb{R}$ d'une suite de points de A .

DÉFINITION 5.9. — Une partie A de $\mathbb{R}$ est dense dans $\mathbb{R}$ si $\bar{A} = \mathbb{R}$.

Nous rappelons un théorème fondamental⁶ concernant les sous-groupes de $(\mathbb{R}, +)$.

THÉORÈME 5.2. — Si S est un sous-groupe du groupe additif $(\mathbb{R}, +)$ de $\mathbb{R}$, il existe un réel $a \geq 0$ tel que $S = a\mathbb{Z}$ ($= \{ka \mid k \in \mathbb{Z}\}$) ou bien S est dense dans $\mathbb{R}$.

Le passage au complémentaire des propriétés de l'intérieur donne : $\overline{A \cup B} = \bar{A} \cup \bar{B}$, mais seulement l'inclusion $\overline{A \cap B} \subset \bar{A} \cap \bar{B}$, qui peut être stricte.

6. Voir [BOUA], chapitre 21, §21.2, théorème 21.1.

5.16. Parties A et B de $\mathbb{R}$ telles que $\overline{A \cap B} \neq \bar{A} \cap \bar{B}$.

Nous posons $A =]0, 1[$ et $B =]1, 2[$. Alors $A \cap B$ est vide donc $\overline{A \cap B} = \emptyset$. Par ailleurs $\bar{A} \cap \bar{B} = [0, 1] \cap [1, 2] = \{1\} \neq \emptyset = \overline{A \cap B}$. $\square$

5.17. Autres parties A et B de $\mathbb{R}$ telles que $\overline{A \cap B} \neq \bar{A} \cap \bar{B}$.

Nous posons, comme dans l'exemple 5.14 (page 90), $A = \mathbb{Q}$ et $B = \mathbb{R} \setminus \mathbb{Q}$. Pour tout nombre réel x et tout réel $\varepsilon > 0$, l'intervalle ouvert $]x - \varepsilon, x + \varepsilon[$ contient des rationnels et des irrationnels, donc tout réel est adhérent à A et à B , ce qui montre que $\bar{A} = \bar{B} = \mathbb{R}$. De plus $A \cap B$ est vide donc :

$$\bar{A} \cap \bar{B} = \mathbb{R} \text{ et } \overline{A \cap B} = \emptyset. \square$$

DÉFINITION 5.10. — La frontière $\text{Fr}(A)$ d'une partie A de $\mathbb{R}$ est l'ensemble des nombres réels adhérents à A qui ne sont pas intérieurs à A .

La frontière d'une partie A de $\mathbb{R}$ est donc $\text{Fr}(A) = \bar{A} \setminus \overset{\circ}{A}$.

Pour toute partie A de $\mathbb{R}$, la frontière de A contient la frontière de $\overset{\circ}{A}$ et celle de $\bar{A}$. Ces inclusions peuvent être strictes.

5.18. Partie A de $\mathbb{R}$ telle que $\text{Fr}(\overset{\circ}{A}) \neq \text{Fr}(A)$ et $\text{Fr}(\bar{A}) \neq \text{Fr}(A)$.

Nous reprenons la partie $A = \mathbb{Q}$ (exemples 5.14 et 5.17). Comme $\overset{\circ}{A} = \emptyset$ et que $\bar{A} = \mathbb{R}$, on a $\text{Fr}(A) = \mathbb{R} \setminus \emptyset = \mathbb{R}$. Or $\mathbb{R}$ et $\emptyset$ sont à la fois ouverts et fermés donc :

$$\overline{\overset{\circ}{A}} = \overline{\emptyset} = \emptyset \text{ et } \bar{\bar{A}} = \bar{\mathbb{R}} = \mathbb{R},$$

ce qui montre que $\text{Fr}(\bar{A}) = \mathbb{R} \setminus \mathbb{R} = \emptyset$ et $\text{Fr}(\overset{\circ}{A}) = \emptyset \setminus \emptyset = \emptyset$. $\square$

5.19. Partie A de $\mathbb{R}$ telle que $\overline{\overset{\circ}{A}} \subsetneq A$.

Nous posons $A = \mathbb{Z}$. Aucun intervalle ouvert non vide n'étant inclus dans $\mathbb{Z}$, $\overset{\circ}{A} = \emptyset$ donc :

$$\overline{\overset{\circ}{A}} = \bar{\emptyset} = \emptyset \subsetneq \mathbb{Z} = A. \square$$

5.20. Partie A de $\mathbb{R}$ telle que $\overline{\overset{\circ}{A}} \subsetneq A \subsetneq \overline{\bar{A}}$.

Posons de nouveau $A = \mathbb{Q}$. Alors $\overset{\circ}{A} = \emptyset$ et $\bar{A} = \mathbb{R}$, donc $\overline{\overset{\circ}{A}} = \emptyset \subsetneq \mathbb{Q} \subsetneq \mathbb{R} = \overline{\bar{A}}$. $\square$

5.21. Partie A de $\mathbb{R}$ telle que $A \subsetneq \overline{\bar{A}}$.

Notons A le complémentaire de $\mathbb{Z}$ dans $\mathbb{R}$. C'est un ouvert comme réunion de la famille $(]n, n+1[)_{n \in \mathbb{Z}}$ d'intervalles ouverts, donc l'intérieur de A est égal à A . Si p est un entier relatif alors, pour tout réel $\varepsilon > 0$, $x = p + \frac{1}{2} \text{Min}(1, \varepsilon)$ appartient à A et à $]p - \varepsilon, p + \varepsilon[$, donc p est adhérent à A ; par conséquent $\bar{A} = \mathbb{R}$. On a donc :

$$A \subsetneq \mathbb{R} = \bar{A} = \overline{\bar{A}}. \square$$

DÉFINITION 5.11. — Si A est une partie de $\mathbb{R}$, un réel a est un point d'accumulation de A s'il est adhérent à $A \setminus \{a\}$.

Il en résulte que si A est une partie de $\mathbb{R}$, un réel a est un point d'accumulation de A si, et seulement si, c'est la limite dans $\mathbb{R}$ d'une suite d'éléments de $A \setminus \{a\}$.

DÉFINITION 5.12. — Si A est une partie de $\mathbb{R}$, l'ensemble dérivé A' de A est l'ensemble des points d'accumulation de A .

Soit A une partie de $\mathbb{R}$. On pose bien sûr $A'' = (A')'$, $A''' = (A'')'$, etc. Si a est un point d'accumulation de A , c'est la limite dans $\mathbb{R}$ d'une suite d'éléments de A , donc a est adhérent à A . Par suite $A' \subset \bar{A}$, mais cette inclusion peut être stricte.

5.22. Partie A de $\mathbb{R}$ telle que $A' \neq \bar{A}$.

Nous posons $A = [0, 1] \cup \mathbb{Z}$. C'est un fermé de $\mathbb{R}$ comme réunion de deux fermés (voir l'exemple précédent 5.21) donc $\bar{A} = A$. Si p est un entier relatif différent de 0 et de 1, l'intervalle ouvert $]p - \frac{1}{2}, p + \frac{1}{2}[$ ne contient aucun point de A différent de p , donc p n'est la limite d'aucune suite d'éléments de $A \setminus \{p\}$, ce qui montre que p n'est pas un point d'accumulation de A . Par conséquent $A' = [0, 1] \subsetneq A$. $\square$

5.23. Partie A de $\mathbb{R}$ telle que $A' \neq \emptyset$ et $A \cap A' = \emptyset$.

Nous posons $A = \left\{ \frac{1}{n} \mid n \in \mathbb{N}^* \right\}$. Comme 0 est la limite de la suite $\left(\frac{1}{n} \right)$, $0 \in A'$. Clairement, $]1/2, 2[$ ne contient aucun point de A différent de $1 = 1/1$ et, pour tout entier $p \geq 2$, on a $1/(p+1) < 1/p < 1/(p-1)$ et l'intervalle ouvert :

$$\left] \frac{1}{p+1}, \frac{1}{p-1} \right[$$

ne contient aucun point de A différent de $1/p$. Par suite, pour tout entier $p \geq 1$, $1/p$ ne peut pas être la limite dans $\mathbb{R}$ d'une suite de points de $A \setminus \{1/p\}$, donc $1/p$ n'est pas un point d'accumulation de A . Ainsi $A' = \{0\}$; de plus $A' \cap A = \emptyset$. $\square$

5.24. Partie A de $\mathbb{R}$ telle que $A' \neq A''$ et $A'' \neq A'''$.

Nous considérons la suite double $(a_{p,n})_{p,n \in \mathbb{N}^*}$, de terme général $a_{p,n} = \frac{1}{p} + \frac{1}{2^n p^2}$, et nous posons :

$$A = \{a_{p,n} \mid p, n \in \mathbb{N}^*\} \text{ et } B = \left\{ \frac{1}{p} \mid p \in \mathbb{N}^* \right\}.$$

Pour p fixé dans $\mathbb{N}^*$, la suite $(a_{p,n})_{n \in \mathbb{N}^*}$ est strictement décroissante et, pour n fixé dans $\mathbb{N}^*$, la suite $(a_{p,n})_{p \in \mathbb{N}^*}$ est strictement décroissante. Par suite, quels que soient $p, n \in \mathbb{N}^*$:

$$a_{p,n+1} < a_{p,n}, \quad a_{p+1,n} < a_{p,n} \text{ et } a_{p,n} \leq a_{1,n} \leq a_{1,1} = \frac{3}{2}.$$

En particulier, $A \subset [0, 3/2]$. Or $[0, 3/2]$ est fermé, donc $A' \subset \bar{A} \subset [0, 3/2]$.

On a, si $p, n \in \mathbb{N}^*$ et $p \geq 2$, $\frac{1}{2^n p^2} < \frac{1}{p^2} < \frac{1}{p(p-1)} = \frac{1}{p-1} - \frac{1}{p}$, donc $\frac{1}{p} < a_{p,n} < \frac{1}{p-1}$.

Si $p \in \mathbb{N}^*$, $\frac{1}{p} = \lim_{n \rightarrow \infty} a_{p,n}$ et $a_{p,n} \neq \frac{1}{p}$ pour tout entier $n \geq 1$, donc $\frac{1}{p}$ appartient à A' .

De plus $0 = \lim_{n \rightarrow \infty} a_{n,n}$ et $a_{n,n} \neq 0$ pour tout n , donc 0 appartient à A' . Par suite :

$$B \cup \{0\} \subset A'.$$

Nous démontrons que cette inclusion est une égalité.

Soit x un réel n'appartenant pas à $B \cup \{0\}$. Si $x < 0$ ou $x > 3/2$, $x \notin [0, 3/2]$ donc $x \notin A'$. Nous supposons que $0 \leq x \leq 3/2$. Comme $x \neq 0$, on a $0 < x \leq 3/2 (= a_{1,1})$.

1^{er} cas : $x > 1$. On voit que $a_{2,1} < 1$, d'où l'on déduit que $a_{p,n} < 1$ pour tout entier $p \geq 2$ et tout entier $n \geq 1$. On pose $q = \text{Max}\{k \in \mathbb{N}^* \mid x \leq a_{1,k}\}$. Alors $q \in \mathbb{N}^*$

et $a_{1,q+1} < x \leq a_{1,q}$. Si $x \neq a_{1,q}$, l'intervalle ouvert $]a_{1,q+1}, a_{1,q}[$ contient x mais aucun élément de A et, si $x = a_{1,q}$, alors, en posant $c = a_{1,q-1}$ si $q \geq 2$ et $c = 2$ si $q = 1$, l'intervalle ouvert $]a_{1,q+1}, c[$ ne contient pas d'autre élément de A que x .
 2^e cas : $x \leq 1$. On pose $p = \text{Min}\{k \in \mathbb{N}^* \mid (1/k) < x\}$. Comme $(1/p) < 1$, on a $p \geq 2$, donc $1/p < x \leq 1/(p-1)$; de plus $x \notin B$, donc $1/p < x < 1/(p-1)$. On prouve alors que : si $x > a_{p,1}$, l'intervalle ouvert $]a_{p,1}, 1/(p-1)[$ contient x mais aucun élément de A ; si $x = a_{p,1}$, l'intervalle ouvert $]a_{p,2}, 1/(p-1)[$ ne contient pas d'autre élément de A que x ; si $x < a_{p,1}$, on pose $n = \text{Min}\{k \in \mathbb{N}^* \mid a_{p,k} \leq x\}$, on voit que $n \geq 2$ et que $a_{p,n} \leq x < a_{p,n-1}$, et alors : si $x \neq a_{p,n}$, l'intervalle ouvert $]a_{p,n}, a_{p,n-1}[$ contient x mais aucun élément de A ; si x est égal à $a_{p,n}$, l'intervalle ouvert $]a_{p,n+1}, a_{p,n-1}[$ ne contient pas d'autre élément de A que x .

Ainsi, dans tous les cas, x n'appartient pas à A' . On a prouvé que $A' \subset B \cup \{0\}$, donc que $A' = B \cup \{0\}$. Nous avons vu dans l'exemple précédent 5.23 que $B' = \{0\}$. Il est alors immédiat que $A'' = \{0\}$, ce qui montre que $A''' = \emptyset$. $\square$

Si A et B sont des parties de $\mathbb{R}$, leur somme est $A + B = \{x + y \mid x \in A \text{ et } y \in B\}$.

Si des parties A et B de $\mathbb{R}$ sont des ouverts de $\mathbb{R}$, il en est de même de $A + B$. Ceci devient faux pour des fermés comme le montre l'exemple suivant.

5.25. Fermés A et B de $\mathbb{R}$ tels que $A + B$ n'est pas un fermé.

Les parties $A = \mathbb{Z}$ et $B = \pi\mathbb{Z}$ de $\mathbb{R}$ sont fermées puisque leurs complémentaires :

$$\mathbb{R} \setminus \mathbb{Z} = \bigcup_{n \in \mathbb{Z}}]n, n+1[\text{ et } \mathbb{R} \setminus (\pi\mathbb{Z}) = \bigcup_{n \in \mathbb{Z}}]n\pi, (n+1)\pi[$$

sont des réunions d'intervalles ouverts. Par ailleurs $A + B$ est le sous-groupe du groupe additif $(\mathbb{R}, +)$ engendré par 1 et π . Supposons l'existence d'un réel $a \geq 0$ tel que $A + B = a\mathbb{Z}$. Il existe des entiers relatifs k et ℓ tels que $1 = ka$ et $\pi = \ell a$. Alors $a \neq 0$, $k \neq 0$ et $\pi = \ell/k$, donc π est rationnel, ce qui est faux. On déduit donc du théorème 5.2 (page 90) que $A + B$ est dense dans $\mathbb{R}$. Comme $A + B$ est dénombrable et que $\mathbb{R}$ ne l'est pas, $A + B \neq \mathbb{R} = \overline{A + B}$. $\square$

5.26. Fermées A et B de $\mathbb{R}$ tels que $\overset{\circ}{A} + \overset{\circ}{B} \neq \overset{\circ}{A + B}$.

Nous posons $A = B = \mathbb{R} \setminus \mathbb{Q}$. Alors $\overset{\circ}{A} = \overset{\circ}{B} = \emptyset$ (exemple 5.14) donc $\overset{\circ}{A} + \overset{\circ}{B} = \emptyset$. Nous démontrons que $A + B = \mathbb{R}$. Soit x un nombre réel. Si x est rationnel, alors $x - \pi$ est irrationnel et $x = (x - \pi) + \pi$ donc $x \in A + B$ et, si x est irrationnel, alors $x/2$ est irrationnel et $x = (x/2) + (x/2)$ donc $x \in A + B$. Par suite :

$$\overset{\circ}{A + B} = \overset{\circ}{\mathbb{R}} = \mathbb{R} \neq \emptyset = \overset{\circ}{A} + \overset{\circ}{B}. \quad \square$$

■ Distance d'un point à une partie et distance entre deux parties

DÉFINITION 5.13. — Si A est une partie non vide de $\mathbb{R}$ et a un nombre réel, la distance de a à A est le réel positif ou nul $\delta(a, A) = \inf_{x \in A} |x - a|$.

Comme il s'agit d'une borne inférieure, cette distance n'est pas forcément atteinte.

5.27. Distance à une partie qui n'est pas atteinte.

Nous posons $A =]0, 1[$ et $a = 0$. Alors $\{|x - a| \mid x \in A\} =]0, 1[$, donc $\delta(0, A) = 0$. Si cette distance était atteinte par un élément x_0 de A , on aurait $0 = |x_0 - 0|$ donc $x_0 = 0$, alors que 0 n'appartient pas à A . $\square$

En fait, si $A \subset \mathbb{R}$ et si $a \in \mathbb{R}$, $\delta(a, A) = 0$ si, et seulement si, a est adhérent à A . Si A est un fermé de $\mathbb{R}$, alors, pour tout réel x , la distance de x à A est atteinte.

DÉFINITION 5.14. — Si A et B sont des parties non vides de $\mathbb{R}$, la distance entre A et B est le réel positif ou nul $d(A, B) = \inf_{x \in A, y \in B} |x - y|$.

5.28. Distance entre deux fermés de $\mathbb{R}$ qui n'est pas atteinte.

Nous posons $A = \mathbb{N}^*$ et $B = \left\{ n - \frac{1}{n} \mid n \in \mathbb{N}^* \right\}$.

Comme A est l'intersection des deux fermés $\mathbb{Z}$ et $\left[\frac{1}{2}, +\infty[\right]$, A est un fermé de $\mathbb{R}$.

Les suites $(n)_{n \geq 1}$ et $(-1/n)_{n \geq 1}$ sont strictement croissantes, donc la suite $(a_n)_{n \geq 1}$, de terme général $a_n = n - (1/n)$, est strictement croissante. De plus $a_1 = 0$, donc le complémentaire de B dans $\mathbb{R}$ est la réunion de la suite $(\Omega_n)_{n \in \mathbb{N}}$ d'ouverts de $\mathbb{R}$ définie par $\Omega_0 =]-\infty, 0[$ et $\Omega_n =]a_n, a_{n+1}[$ pour tout entier $n \geq 1$. Par suite $\complement_{\mathbb{R}} B$ est un ouvert de $\mathbb{R}$, donc B est un fermé de $\mathbb{R}$. Pour tout entier $n \geq 1$, n appartient à A et a_n à B et la suite de terme général $|n - a_n| = 1/n$ admet pour limite 0, donc $d(A, B) = 0$.

Par ailleurs l'ensemble B ne contient aucun entier sauf 0, donc $A \cap B$ est vide, ce qui montre que, quels que soient $x \in A$ et $y \in B$, $|x - y| \neq 0$. La distance entre A et B n'est donc pas atteinte. $\square$

■ Endomorphismes du groupe additif $(\mathbb{R}, +)$ de $\mathbb{R}$

Le corps $\mathbb{Q}$ des nombres rationnels étant un sous-corps de $\mathbb{R}$, $\mathbb{R}$ est canoniquement un espace vectoriel sur $\mathbb{Q}$; les scalaires sont alors les nombres rationnels et les vecteurs les nombres réels.

Si $\mathbb{R}$ était de dimension finie n sur $\mathbb{Q}$, on aurait $n \geq 1$ puisque $\mathbb{R}$ est différent de $\{0\}$ et l'espace vectoriel $\mathbb{R}$ sur $\mathbb{Q}$ serait isomorphe à $\mathbb{Q}^n$, donc $\mathbb{R}$ serait équipotent⁷ à $\mathbb{Q}^n$, ce qui est faux puisque $\mathbb{Q}^n$ est dénombrable alors que $\mathbb{R}$ ne l'est pas. Il en résulte que l'espace vectoriel $\mathbb{R}$ sur $\mathbb{Q}$ n'est pas de dimension finie.

Grâce à l'axiome du choix⁸, on prouve, pour les espaces vectoriels de dimension infinie sur un corps commutatif — c'est-à-dire ceux qui ne sont pas de dimension finie sur ce corps — l'essentiel des théorèmes établis en dimension finie, en

7. Voir le chapitre 1, page 4.

8. Voir le chapitre 1, page 14.

particulier l'existence de bases⁹, le fait que toutes les bases d'un espace vectoriel ont le même cardinal, d'où la notion de dimension¹⁰, le théorème de la base incomplète et l'existence de supplémentaires pour les sous-espaces vectoriels.

Les endomorphismes de groupe du groupe additif $(\mathbb{R}, +)$ de $\mathbb{R}$, c'est-à-dire les applications f de $\mathbb{R}$ dans $\mathbb{R}$ telles que $f(x + y) = f(x) + f(y)$ quels que soient les réels x et y , sont les endomorphismes de l'espace vectoriel $\mathbb{R}$ sur $\mathbb{Q}$ — en effet, si f est un endomorphisme du groupe $(\mathbb{R}, +)$, alors $f(nx) = nf(x)$ pour $x \in \mathbb{R}$ et $n \in \mathbb{Z}$ et, si $x \in \mathbb{R}$ et $r \in \mathbb{Q}$, en choisissant un représentant (n, d) de r , on a $df(rx) = f(d(rx)) = f((dr)x) = f(nx) = nf(x)$, donc $f(rx) = rf(x)$.

THÉORÈME 5.3. — Si f est endomorphisme du groupe additif $(\mathbb{R}, +)$ de $\mathbb{R}$ et si l'application f est continue ou monotone, f est une homothétie de l'espace vectoriel $\mathbb{R}$ sur $\mathbb{R}$ — ce qui signifie qu'il existe un nombre réel α tel que $f = \alpha \text{Id}_{\mathbb{R}}$.

Nous démontrons ce résultat. Soit f endomorphisme du groupe $(\mathbb{R}, +)$, monotone ou continu. Nous savons que f est $\mathbb{Q}$ -linéaire. Nous posons $\alpha = f(1)$. On a, pour tout rationnel r , $f(r) = f(r1) = rf(1) = \alpha r$. Soit x un réel. Notons $(r_n)_{n \geq 0}$ (resp. $(s_n)_{n \geq 0}$) la suite des valeurs décimales approchées par défaut (resp. par excès) de x . Pour tout entier naturel n , r_n et s_n sont des nombres rationnels et $r_n \leq x < s_n$, et les suites (r_n) et (s_n) admettent pour limite x . Si f est continue, le passage à la limite dans l'égalité $f(r_n) = \alpha r_n$, valable pour tout $n \in \mathbb{N}$, montre que $f(x) = \alpha x$. Si f est croissante (resp. décroissante), alors, pour tout entier naturel n , $f(r_n) \leq f(x) \leq f(s_n)$ (resp. $f(s_n) \leq f(x) \leq f(r_n)$), donc $\alpha r_n \leq f(x) \leq \alpha s_n$ (resp. $\alpha s_n \leq f(x) \leq \alpha r_n$), et en passant à la limite, on obtient l'égalité $f(x) = \alpha x$. Dans les deux cas, $f(x) = \alpha x$. En conclusion, $f = \alpha \text{Id}_{\mathbb{R}}$. $\square$

Le théorème 5.3 est très important au point de vue théorique car il permet de démontrer que les fonctions exponentielles, logarithmes et puissances réelles sont caractérisées par leur équation fonctionnelle et la continuité ou la monotonie.

Cependant les endomorphismes du groupe additif $(\mathbb{R}, +)$ ne sont pas tous des homothéties de l'espace vectoriel $\mathbb{R}$ sur $\mathbb{R}$.

5.29. Endomorphisme du groupe additif $(\mathbb{R}, +)$ qui n'est pas une homothétie de l'espace vectoriel $\mathbb{R}$ sur $\mathbb{R}$.

L'espace vectoriel $\mathbb{Q}$ sur $\mathbb{Q}$ est un sous-espace vectoriel de l'espace vectoriel $\mathbb{R}$ sur $\mathbb{Q}$. Nous choisissons un supplémentaire S de $\mathbb{Q}$ dans $\mathbb{R}$ — il en existe; voir ci-dessus — et nous notons f la projection sur $\mathbb{Q}$ parallèlement à S . Alors f est un endomorphisme de l'espace vectoriel $\mathbb{R}$ sur $\mathbb{Q}$, donc un endomorphisme du groupe additif $(\mathbb{R}, +)$ de $\mathbb{R}$. Si $S = \{0\}$, $\mathbb{R} = \mathbb{Q} \oplus S = \mathbb{Q} \oplus \{0\} = \mathbb{Q}$, ce qui est faux. Par suite $S \neq \{0\}$. Nous choisissons un élément non nul s de S . Si f était une homothétie de l'espace vectoriel $\mathbb{R}$ sur $\mathbb{R}$, il existerait un réel α tel que $f = \alpha \text{Id}_{\mathbb{R}}$ et, comme $0 = f(s) = \alpha s$, α serait nul, donc on aurait $1 = f(1) = \alpha 1 = 0$, en contradiction avec $1 \neq 0$. $\square$

9. Les bases de l'espace vectoriel $\mathbb{R}$ sur $\mathbb{Q}$ sont appelées les bases de Hamel, du nom du mathématicien allemand Georg Hamel (1877-1954).

10. Par exemple la dimension de l'espace vectoriel $\mathbb{R}$ sur $\mathbb{Q}$ est la puissance du continu, c'est-à-dire le cardinal de $\mathbb{R}$.

■ Mesurabilité de parties de $\mathbb{R}$

L'intégrale de Lebesgue, plus générale que l'intégrale de Riemann, se définit à l'aide de la notion de mesure. Cette notion¹¹ ne sera pas définie ici et nous nous contenterons d'étudier quelques propriétés simples de la mesure de Lebesgue sur $\mathbb{R}$, que nous notons m . Pour une partie A de $\mathbb{R}$, les trois assertions :

- (I) A est finie ou dénombrable,
- (II) $m(A) = 0$,
- (III) l'intérieur de A est vide,

ne sont pas équivalentes. Les implications (I) implique (II) et (II) implique (III) sont vraies ; nous prouvons que les réciproques sont fausses. Pour ceci nous aurons besoin de la notion d'équipotence et du théorème de Cantor-Bernstein¹².

Rappelons que $\mathbb{R}$ est équipotent à tous ses intervalles d'intérieur non vide, c'est-à-dire lui-même, les $]a, b[$, $[a, b[$, $]a, b]$ et $[a, b]$ pour des réels a et b tels que $a < b$ et les $]a, +\infty[$, $[a, +\infty[$, $]-\infty, a[$ et $]-\infty, a]$ où a est un nombre réel¹³.

5.30. Partie de mesure nulle équipotente à $\mathbb{R}$.

Pour tout entier $n \geq 1$, nous notons K_n l'ensemble des entiers $k \in [0, 3^n - 1]$ ne contenant aucun 1 dans leur écriture en base 3. Le cardinal de K_n est 2^n puisque ses éléments sont les entiers qui s'écrivent en base 3 avec n chiffres appartenant à la paire $\{0, 2\}$. On a par exemple $K_1 = \{0, 2\}$ et $K_2 = \{0, 2, 6, 8\}$. Nous posons $C_0 = [0, 1]$ et, pour tout entier $n \geq 1$:

$$C_n = \bigcup_{k \in K_n} \left[\frac{k}{3^n}, \frac{k+1}{3^n} \right].$$

En particulier, $C_1 = \left[0, \frac{1}{3}\right] \cup \left[\frac{2}{3}, 1\right]$ et $C_2 = \left[0, \frac{1}{8}\right] \cup \left[\frac{2}{9}, \frac{1}{3}\right] \cup \left[\frac{2}{3}, \frac{7}{9}\right] \cup \left[\frac{8}{9}, 1\right]$.

Soit n un entier ≥ 2 . Soit $k \in K_n$. Le dernier chiffre de l'écriture de k en base 3 est 0 ou 2. Si c'est 0, on a $k = 3\ell$ où $\ell \in K_{n-1}$; si c'est 2, $k = 3\ell + 2$ où $\ell \in K_{n-1}$. Ainsi, dans les deux cas :

$$\left[\frac{k}{3^n}, \frac{k+1}{3^n} \right] \subset \left[\frac{\ell}{3^{n-1}}, \frac{\ell+1}{3^{n-1}} \right] \subset C_{n-1}.$$

Par conséquent $C_n \subset C_{n-1}$ — en fait on passe de C_{n-1} à C_n en enlevant le tiers central ouvert de chacun des segments disjoints dont C_{n-1} est la réunion.

Pour tout $n \in \mathbb{N}^*$, la mesure de C_n est $\sum_{k \in K_n} m\left(\left[\frac{k}{3^n}, \frac{k+1}{3^n}\right]\right) = 2^n \times \frac{1}{3^n} = \left(\frac{2}{3}\right)^n$.

Nous posons $C = \bigcap_{n \in \mathbb{N}} C_n$. L'ensemble C est l'ensemble triadique de Cantor¹⁴.

L'ensemble C est un compact de $\mathbb{R}$ comme intersection de compacts et, comme $m(C) \leq m(C_n)$ pour tout entier naturel n , C est de mesure nulle.

11. Pour un exposé de la mesure de Lebesgue sur $\mathbb{R}$, voir [BOUA], chapitre 22.

12. Voir dans le chapitre 1 la définition 1.1 page 4 et le théorème 1.2 page 7.

13. Compte tenu du théorème de Cantor-Bernstein, il suffit d'établir que si $a, b \in \mathbb{R}$ et si $a < b$, $\mathbb{R}$ est équipotent à $]a, b[$, ce qui est fait au chapitre 1, page 5, à la suite de l'exemple 1.11.

14. Il a été introduit par Cantor.

Nous démontrons que C est équipotent à $\mathbb{R}$.

Les nombres triadiques sont les quotients $a/3^n$ où a est un nombre entier et n un entier naturel. Si $(c_i)_{i \geq 1}$ est une suite de chiffres en base 3 — ce qui signifie que $c_i \in \{0, 1, 2\}$ pour tout i — on pose :

$$0, c_1 c_2 c_3 \dots c_n \dots = \sum_{i=1}^{+\infty} \frac{c_i}{3^i},$$

le réel x ainsi défini appartient à $[0, 1]$ et l'écriture $x = 0, c_1 c_2 c_3 \dots c_n \dots$ est appelée un développement illimité de x en base 3. Nous rappelons les résultats suivants. Tout point x de $[0, 1]$ qui n'est pas un nombre triadique admet un unique développement illimité $x = 0, c_1 c_2 c_3 \dots c_n \dots$ en base 3 — et alors il n'existe aucun entier $p \geq 1$ tel que $c_i = 2$ pour tout $i \geq p$ —, l'unique développement illimité de 0 en base 3 est $0 = 0, 000 \dots 0 \dots$ et l'unique développement illimité de 1 en base 3 s'écrit $1 = 0, 222 \dots 2 \dots$. Enfin un nombre triadique x appartenant à $]0, 1[$ admet deux développements illimités en base 3 : le premier, dit *propre*, s'écrit $x = 0, b_1 b_2 b_3 \dots b_n \dots$, les chiffres $b_1, b_2, b_3, \dots, b_n \dots$ étant tous nuls à partir d'un certain rang et le second, dit *impropre*, s'écrit $x = 0, c_1 c_2 c_3 \dots c_n \dots$, les chiffres $c_1, c_2, c_3, \dots, c_n, \dots$ étant tous égaux à 2 à partir d'un certain rang, et on passe du premier au second de la manière suivante : si k est le plus grand des indices i tel que $b_i \neq 0$, alors $c_i = b_i$ pour $i < k$, $c_k = b_k - 1$ et $c_i = 2$ pour tout $i > k$.

On démontre alors¹⁵ qu'un point x du segment $[0, 1]$ appartient à C si, et seulement si, il admet un développement illimité $x = 0, c_1 c_2 c_3 \dots c_n \dots$ en base 3 dont tous les chiffres $c_1, c_2, c_3, \dots, c_n, \dots$ sont différents de 1, et que tout point x de C admet un unique développement illimité en base 3 de ce type. Par exemple :

$$\frac{1}{3} = 0, 1 = 0, 02222 \dots 2 \dots \text{ et } \frac{2}{3} = 0, 2 = 0, 12222 \dots 2 \dots$$

appartiennent à C et un seul de leurs deux développements est du type demandé.

On définit de la même façon les développements illimités en base 2 des points du segment $[0, 1]$, en remplaçant 3 par 2 et 2 par 1, les nombres dyadiques $a/2^n$ (où a est un nombre entier et n un entier naturel) se substituant aux nombres triadiques.

Si x est un point du segment $[0, 1]$ et si son développement illimité en base 2 s'écrit $x = 0, a_1 a_2 a_3 \dots a_n \dots$ — il s'agit, si x est un nombre dyadique appartenant à $]0, 1[$, du développement propre —, nous posons :

$$f(x) = \sum_{i=1}^{+\infty} \frac{c_i}{3^i} \text{ où } c_i = 2a_i \text{ pour tout entier } i \geq 1,$$

ce qui, puisque $c_i = 0$ ou 2 pour tout i , montre que $f(x)$ appartient à C . On a ainsi construit une application f de $[0, 1]$ dans C , et tout point de C admettant un unique développement illimité en base 3 dont tous les chiffres sont différents de 1, f est injective. Par suite f , l'injection canonique de C dans $[0, 1]$ et le théorème de Cantor-Bernstein montrent que C est équipotent à $[0, 1]$. Or $[0, 1]$ est équipotent à $\mathbb{R}$, donc C est équipotent à $\mathbb{R}$. $\square$

15. Voir [BOUA], chapitre 23, § 23.3, théorème 23.1.

5.31. Partie de $\mathbb{R}$ d'intérieur vide et de mesure strictement positive.

Nous notons A l'ensemble des nombres irrationnels appartenant au segment $[0, 1]$. Alors $m(A) = m([0, 1]) - m([0, 1] \cap \mathbb{Q})$; or $[0, 1] \cap \mathbb{Q}$ est dénombrable, donc de mesure nulle, ce qui montre que $m(A) = 1$. Cependant l'intérieur de A est vide puisque tout intervalle ouvert non vide contient au moins un nombre rationnel. $\square$

5.32. Ouvert de $\mathbb{R}$ contenant $\mathbb{Q}$ et de mesure aussi petite que l'on veut.

Soit un réel $\varepsilon > 0$. L'ensemble $\mathbb{Q}$ est dénombrable donc $\mathbb{Q} = \{r_1, r_2, \dots, r_n, \dots\}$ où les r_i sont deux à deux distincts. Pour tout entier $n \geq 1$, l'ensemble :

$$O_n = \left] r_n - \frac{\varepsilon}{2^{n+1}}, r_n + \frac{\varepsilon}{2^{n+1}} \right[$$

est un ouvert de mesure $\varepsilon/2^n$ contenant r_n , donc la réunion O de la suite $(O_n)_{n \in \mathbb{N}^*}$ est un ouvert contenant $\mathbb{Q}$, et on a :

$$m(O) \leq \sum_{n=1}^{+\infty} m(O_n) = \sum_{n=1}^{+\infty} \frac{\varepsilon}{2^n} = \varepsilon. \quad \square$$

Bien que $\mathbb{Q}$ soit de mesure nulle, on ne peut pas trouver d'ouvert W de mesure nulle contenant $\mathbb{Q}$; en effet, la mesure d'un ouvert non vide n'est jamais nulle.

Nous examinons pour terminer deux exemples plus délicats. Le premier montre que l'axiome du choix prouve l'existence d'au moins une partie de $\mathbb{R}$ qui n'est pas mesurable. Signalons que Robert Solovay¹⁶ a établi que la proposition « *Toute partie de $\mathbb{R}$ est mesurable* » est indécidable dans la théorie axiomatique ZF_0 obtenue en supprimant l'axiome du choix de la liste des axiomes de ZF — la théorie des ensembles Zermelo-Fraenkel¹⁷ —, ce qui signifie que ni elle ni sa négation ne sont démontrables dans ZF_0 .

5.33. Partie de $\mathbb{R}$ qui n'est pas mesurable.

Nous définissons sur $[0, 1]$ la relation d'équivalence $\mathcal{R}$ par : $x \mathcal{R} y$ si $x - y \in \mathbb{Q}$. L'axiome du choix justifie la construction d'une partie A de $[0, 1]$ obtenue en choisissant un élément et un seul dans chaque classe d'équivalence modulo $\mathcal{R}$. Nous associons à tout rationnel a l'ensemble $A_a = \{a + x \mid x \in A\}$.

Soit a et b des nombres rationnels distincts. Si z appartient à $A_a \cap A_b$, il existe des éléments x et y de A tels que $z = a + x = b + y$, donc $y - x = a - b$ est un rationnel différent de zéro, en contradiction avec le fait que, dans A , on ne peut avoir deux éléments distincts équivalents modulo $\mathcal{R}$. Par suite $A_a \cap A_b$ est vide.

Supposons A mesurable.

Nous posons $Q = \mathbb{Q} \cap [0, 1]$. Pour tout élément a de Q , $A_a \subset [0, 2]$, donc l'ensemble $B = \bigcup_{a \in Q} A_a$ est inclus dans $[0, 2]$. La mesure de Lebesgue étant invariante par translation, on a $m(A) = m(A_a)$ pour tout élément a de Q .

16. Mathématicien américain, spécialiste de théorie des ensembles.

17. Voir le chapitre 1, pages 3 et 14.

Supposons que $m(A) > 0$. L'ensemble Q est dénombrable et $(A_a)_{a \in Q}$ est une famille de parties mesurables deux à deux disjointes, donc B est mesurable et :

$$m(B) = m\left(\bigcup_{a \in Q} A_a\right) = \sum_{a \in Q} m(A_a) = +\infty,$$

ce qui est impossible puisque B est inclus dans $[0, 2]$. Il en résulte que $m(A) = 0$. Soit z un point du segment $[0, 1]$. Alors z appartient à une classe d'équivalence modulo $\mathbb{R}$, donc il existe $a \in \mathbb{Q}$ et $x \in A$ tels que $z = a + x$; par conséquent $z \in A_a$; de plus $a = z - x$, $z \in [0, 1]$ et $x \in [0, 1]$, donc $a \in [-1, 1]$. Finalement, en posant $D = \mathbb{Q} \cap [-1, 1]$, le segment $[0, 1]$ est inclus dans $E = \bigcup_{a \in D} A_a$. Or D est dénombrable et les A_a sont tous mesurables et de mesure nulle, donc E est mesurable et de mesure nulle, ce qui est absurde puisqu'il contient le segment $[0, 1]$ dont la mesure est égale à 1. En conclusion, A n'est pas mesurable. $\square$

5.34. Application continue qui envoie un ensemble de mesure nulle sur un ensemble de mesure 1.

Nous avons défini dans l'exemple 5.30 (pages 96 et 97) une suite $(C_n)_{n \in \mathbb{N}}$ de parties de $[0, 1]$ et leur intersection C (l'ensemble triadique de Cantor). Nous introduisons la fonction f définie dans l'exemple 10.3 (pages 193 à 196), limite uniforme d'une suite (f_n) de fonctions affines par morceaux continues sur le segment $[0, 1]$.

Nous posons $D = [0, 1] \setminus C$. Soit x un point de D . Il existe un entier $n \geq 1$ tel que $x \in C_{n-1}$ et $x \notin C_n$. Alors x appartient à un intervalle ouvert sur lequel f_n est constante, la valeur de la constante étant $k/2^n$ où $k \in \llbracket 1, 2^n - 1 \rrbracket$ — voir le dessin de l'exemple 10.3 —, et de plus, pour tout entier $q \geq n$, $f_q(x) = f_n(x)$, donc :

$$f(x) = \frac{k}{2^n}.$$

Par suite $f(D)$ est inclus dans l'ensemble $B = \left\{ \frac{k}{2^n} \mid n \in \mathbb{N}^* \text{ et } k \in \llbracket 0, 2^n \rrbracket \right\}$.

Or B est dénombrable, donc aussi $f(D)$, ce qui montre que $f(D)$ est de mesure nulle. De plus $f([0, 1]) = [0, 1]$ — en effet, f est continue et à valeurs dans $[0, 1]$, $f(0) = 0$ et $f(1) = 1$ — donc $[0, 1] \setminus f(D) \subset f(C) \subset [0, 1]$, et comme $[0, 1] \setminus f(D)$ et $[0, 1]$ sont de mesure 1, la mesure de $f(C)$ est égale à 1. Pourtant C est de mesure nulle comme nous l'avons vu dans l'exemple 5.30. $\square$

Remarquons qu'une application absolument continue¹⁸ envoie les ensembles de mesure nulle sur des ensembles de mesure nulle¹⁹.

18. Définition 8.9, page 144.

19. Ceci redémontre que l'application f de l'exemple 5.34 n'est pas absolument continue sur le segment $[0, 1]$, comme établi dans l'exemple 10.3, pages 193 à 196.

Chapitre 6

Suites numériques

La notion de suite apparaît très tôt en mathématiques. Le premier exemple célèbre est la suite de Fibonacci que ce dernier introduit au début du XIII^e siècle. Cependant la formalisation précise de la notion de convergence d'une suite date du début du XIX^e siècle.

L'importance des suites en analyse provient du fait que les problèmes topologiques de $\mathbb{R}$, et plus généralement d'un espace métrique, peuvent être traités à l'aide de suites. Nous nous intéresserons dans ce chapitre aux problèmes concernant les suites de nombres réels.

■ Convergence et divergence

DÉFINITION 6.1. — Une suite numérique est une application de $\mathbb{N}$ dans $\mathbb{R}$.

L'image d'un entier naturel n par une suite se note le plus souvent indiciellement x_n au lieu de $f(n)$. Pour alléger l'écriture et lorsqu'il n'y a pas de confusion possible, nous notons (x_n) la suite de terme général x_n , c'est-à-dire l'application $n \mapsto x_n$ de $\mathbb{N}$ dans $\mathbb{R}$. Souvent une suite (x_n) n'est définie « qu'à partir d'un certain rang », par exemple le rang k , ce qui signifie que x_n n'existe à coup sûr que pour tous les entiers $n \geq k$ et si cette précision est utile, on la note $(x_n)_{n \geq k}$; on pourra donc étudier des suites $(x_n)_{n \geq 0} = (x_n)_{n \in \mathbb{N}}$, $(y_n)_{n \geq 1} = (y_n)_{n \in \mathbb{N}^*}$, $(z_n)_{n \geq 2}$, etc.

► Dans tout ce qui suit, (x_n) , (y_n) , (u_n) , (v_n) ... sont des suites numériques.

DÉFINITION 6.2. — La suite (x_n) converge vers le nombre réel λ si, pour tout réel $\varepsilon > 0$, il existe un entier naturel N tel que $|x_n - \lambda| < \varepsilon$ pour tout entier $n \geq N$.

Une suite (x_n) converge dans $\mathbb{R}$ s'il existe au moins un nombre réel λ tel que (x_n) converge vers λ , et il y a alors unicité de λ , que l'on appelle *la limite* (dans $\mathbb{R}$) de la suite (x_n) , notée $\lim_{n \rightarrow \infty} x_n$ ou $\lim_{(n \rightarrow \infty)} x_n$.

Si une suite (x_n) converge vers un nombre réel λ , on dit aussi que (x_n) admet pour limite λ (dans $\mathbb{R}$) ou que (x_n) tend vers λ (dans $\mathbb{R}$).

Une suite qui converge dans $\mathbb{R}$, que l'on appelle une suite convergente (sous-entendu : dans $\mathbb{R}$) est bornée, ce qui signifie qu'elle est minorée et majorée, ou encore qu'il existe une constante réelle $M \geq 0$ telle que $|x_n| \leq M$ pour tout n .

Une suite qui diverge dans $\mathbb{R}$, appelée aussi une suite divergente (sous-entendu : dans $\mathbb{R}$) est une suite qui ne converge pas dans $\mathbb{R}$.

DÉFINITION 6.3. — La suite (x_n) tend vers $+\infty$ (resp. vers $-\infty$) si, pour tout nombre réel A (resp. B), il existe un entier naturel N tel que $x_n > A$ (resp. $x_n < B$) pour tout entier $n \geq N$.

Une suite qui tend vers $+\infty$ (resp. vers $-\infty$) n'est pas majorée (resp. n'est pas minorée) dans $\mathbb{R}$, donc c'est une suite divergente ; dans ce cas on dit aussi que la suite (x_n) diverge vers $+\infty$ (resp. vers $-\infty$), ou encore que (x_n) admet pour limite $+\infty$ (resp. $-\infty$) — sous-entendu : dans la droite numérique achevée $\overline{\mathbb{R}}$ — et ceci se traduit en symboles par : $\lim_{n \rightarrow \infty} x_n = +\infty$ (resp. $\lim_{n \rightarrow \infty} x_n = -\infty$).

DÉFINITION 6.4. — Une suite (y_n) est une sous-suite — ou une suite extraite — de la suite (x_n) s'il existe une application strictement croissante φ de $\mathbb{N}$ dans $\mathbb{N}$ telle que $y_n = x_{\varphi(n)}$ pour tout n .

Si une suite converge vers un réel λ , toutes ses sous-suites convergent vers λ , et le résultat est encore valable pour une suite qui tend vers $+\infty$ ou vers $-\infty$.

THÉORÈME 6.1. — Théorème de Bolzano-Weierstrass¹.

De toute suite bornée on peut extraire une sous-suite convergente.

Ceci ne signifie évidemment pas que toute suite bornée soit convergente !

6.1. Suite bornée divergente.

Nous introduisons la suite $(x_n)_{n \geq 0}$ de terme général $x_n = (-1)^n$. Pour tout entier naturel n , $|x_n| = 1 \leq 1$, donc (x_n) est bornée dans $\mathbb{R}$. Soit λ un nombre réel. On a, pour tout entier naturel n , $|x_{n+1} - x_n| = 2$, donc $|x_n - \lambda| \geq 1$ ou $|x_{n+1} - \lambda| \geq 1$. Il en résulte que la suite (x_n) est divergente. $\square$

DÉFINITION 6.5. — La suite (x_n) est équivalente à la suite (y_n) s'il existe une suite (α_n) qui converge vers 1 et un entier naturel N tels que $y_n = \alpha_n x_n$ pour tout entier $n \geq N$.

On obtient ainsi une relation binaire, que l'on appelle la relation d'équivalence sur l'ensemble des suites numériques, et c'est une relation... d'équivalence. En raison de la symétrie, on dit que les suites (x_n) et (y_n) sont équivalentes pour exprimer que (x_n) est équivalente à (y_n) ou que (y_n) est équivalente à (x_n) . Si y_n est différent de zéro à partir d'un certain rang, les suites (x_n) et (y_n) sont équivalentes si, et seulement si, la suite quotient (x_n/y_n) converge vers 1. Si les suites (x_n) et (y_n) sont équivalentes et si (y_n) tend vers λ dans $\mathbb{R}$, ou si (y_n) diverge vers $+\infty$ ou vers $-\infty$, il en est de même de la suite (x_n) .

THÉORÈME 6.2. — Si des suites (x_n) et (y_n) à termes réels strictement positifs sont équivalentes et si elles convergent vers 0 ou divergent vers $+\infty$, leurs « suites logarithmes » $(\ln x_n)$ et $(\ln y_n)$ sont équivalentes. Ce résultat se généralise au cas où 1 n'est pas valeur d'adhérence de l'une des deux suites².

1. Ce théorème est démontré dans son cours en 1874 par Karl Weierstrass qui n'en publie pas la démonstration. Bernhard Bolzano l'avait énoncé vers 1830, mais ses travaux ne furent découverts qu'en 1930. Pour une démonstration de ce théorème, voir [ARN2], §III.3. ou [BOUA], §14.5.

2. Le rappel de la notion de valeur d'adhérence se trouve dans la suite de ce chapitre, page 107.

La conclusion devient fausse si les suites convergent vers 1.

6.2. Suites équivalentes dont les « suites logarithmes » ne sont pas équivalentes.

Nous considérons les suites $(x_n)_{n \geq 1}$ et $(y_n)_{n \geq 1}$ de termes généraux :

$$x_n = 1 + \frac{1}{n} \text{ et } y_n = 1 + \frac{1}{n^2}.$$

La suite quotient $\left(\frac{x_n}{y_n}\right)$ tend vers 1, donc les suites (x_n) et (y_n) sont équivalentes.

La fonction logarithme népérien est dérivable en 1, $\ln 1 = 0$ et $\ln'(1) = 1$, donc il existe une fonction ω , définie sur $] -1, +\infty[$, telle que $\omega(x)$ admet pour limite 0 quand x tend vers 0 et $\ln(1+x) = x + x\omega(x)$ pour tout point x de $] -1, +\infty[$. Pour tout entier $n \geq 2$, $\frac{1}{n}$ et $\frac{1}{n^2}$ appartiennent à $]0, 1[$ et y_n à $]1, 2[$, donc $\ln y_n \neq 0$ et :

$$u_n = \frac{\ln x_n}{\ln y_n} = \frac{\frac{1}{n} + \frac{1}{n}\omega\left(\frac{1}{n}\right)}{\frac{1}{n^2} + \frac{1}{n^2}\omega\left(\frac{1}{n^2}\right)} = n \frac{1 + \omega\left(\frac{1}{n}\right)}{1 + \omega\left(\frac{1}{n^2}\right)}.$$

Il en résulte que la suite (u_n) tend vers $+\infty$, ce qui montre que les suites $(\ln x_n)$ et $(\ln y_n)$ ne sont pas équivalentes. $\square$

DÉFINITION 6.6. — La suite (x_n) est une suite de Cauchy si, pour tout réel $\varepsilon > 0$, il existe un entier naturel N tel que, quels que soient les entiers $p \geq N$ et $q \geq N$:

$$|x_q - x_p| < \varepsilon.$$

Une suite à valeurs réelles converge dans $\mathbb{R}$ si, et seulement si, c'est une suite de Cauchy³. L'intérêt de cette propriété réside dans le fait que la connaissance préalable de la limite n'est pas nécessaire pour montrer qu'une suite converge. On peut ainsi introduire de nouveaux nombres comme limites de suites de Cauchy.

En prenant $p = n$ et $q = n + 1$ dans la définition, on voit que pour une suite de Cauchy (x_n) , la suite $(x_{n+1} - x_n)$ converge vers 0. Cependant la réciproque est fausse, même si l'on renforce cette affirmation de la manière suivante : pour tout entier $k \geq 1$, la suite $(x_{n+k} - x_n)$ converge vers 0.

6.3. Suite (x_n) divergente telle que, pour tout entier $k \geq 1$:

$$\lim_{n \rightarrow \infty} (x_{n+k} - x_n) = 0.$$

Nous introduisons la suite $(x_n)_{n \geq 1}$ de terme général $x_n = 1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{n} = \sum_{p=1}^n \frac{1}{p}$. On a, pour tout entier $n \geq 1$:

$$x_{2n} - x_n = \sum_{p=n+1}^{2n} \underbrace{\frac{1}{p}}_{\geq \frac{1}{2n}} \geq n \times \frac{1}{2n} = \frac{1}{2}$$

3. Bolzano utilise cette propriété des suites de nombres réels en 1817 dans sa démonstration du théorème des valeurs intermédiaires, mais c'est Cauchy qui introduit cette notion dans son cours d'Analyse de 1821 pour l'étude de la convergence des séries ([CAU1], chapitre VI, §1^{er}, pages 125 et 126 ; voir la note n° 3 du chapitre 16, page 309). Bolzano et Cauchy admettent cette propriété sans démonstration, que seule permet une construction précise des nombres réels ; pour ceci il faudra encore attendre un demi-siècle. Dans le langage actuel, on dit que $\mathbb{R}$, muni de sa distance canonique $(x, y) \mapsto |x - y|$, est un espace métrique complet ; voir le chapitre 16, page 312.

donc (x_n) n'est pas une suite de Cauchy — prendre $\varepsilon = \frac{1}{2}$ dans la définition.

Soit un entier $k \geq 1$. On obtient, pour tout entier $n \geq 1$:

$$0 < x_{n+k} - x_n = \sum_{p=n+1}^{n+k} \underbrace{\frac{1}{p}}_{\leq \frac{1}{n}} \leq \frac{k}{n}$$

donc la suite $(x_{n+k} - x_n)$ converge vers 0. $\square$

THÉORÈME 6.3. — Si λ est un nombre réel et f une fonction définie sur un voisinage de λ et continue en λ , alors, pour toute suite (x_n) convergeant vers λ , la suite image $(f(x_n))$ converge vers $f(\lambda)$.

Cependant, si une suite (x_n) converge vers un nombre réel λ , si $\mathcal{A}$ est une partie de $\mathbb{R}$, si x_n appartient à $\mathcal{A}$ pour tout n et si une fonction f est définie et continue sur $\mathcal{A}$, la suite image $(f(x_n))$ peut diverger — dans ce cas λ est bien sûr adhérent à $\mathcal{A}$, mais λ n'appartient pas à $\mathcal{A}$.

6.4. Suite convergente (x_n) et fonction f continue telles que la suite image $(f(x_n))$ diverge.

Nous considérons la suite $(x_n)_{n \geq 1}$, de terme général $x_n = \frac{1}{n}$, et la fonction f définie sur $\mathcal{A} = \mathbb{R} \setminus \{0\}$ par :

$$f(x) = \frac{1}{x}.$$

La suite (x_n) converge vers 0, x_n appartient à $\mathcal{A}$ pour tout entier $n \geq 1$ et f est continue sur $\mathcal{A}$. Or $f(x_n) = n$ pour tout $n \in \mathbb{N}^*$, donc la suite $(f(x_n))$ diverge. $\square$

Une suite de limite nulle n'a aucune raison d'être monotone, ni même monotone à partir d'un certain rang, contrairement à ce qu'écrivent souvent les débutants.

6.5. Suite de limite nulle qui n'est pas monotone.

Nous introduisons la suite $(u_n)_{n \geq 1}$, de terme général $u_n = \frac{(-1)^n}{n}$, qui tend vers 0. Pour tout entier $n \geq 1$, $u_{2n} > u_{2n+1}$ et $u_{2n} > u_{2n-1}$, donc la suite (u_n) n'est ni croissante ni décroissante à partir d'un certain rang. $\square$

Dans l'exemple précédent 6.5, la suite (u_n) s'alterne autour de sa limite 0. Ce n'est en fait pas nécessaire : le terme général de la suite peut être de signe constant.

6.6. Suite positive de limite nulle qui n'est pas monotone.

Nous considérons la suite $(u_n)_{n \geq 1}$, de terme général $u_n = \begin{cases} \frac{1}{n} & \text{si } n \text{ est pair,} \\ \frac{1}{n^2} & \text{si } n \text{ est impair.} \end{cases}$

On a, pour tout entier $n \geq 1$, $0 < u_n \leq \frac{1}{n}$, donc (u_n) tend vers 0. De plus, pour tout entier $n \geq 1$:

$$u_{2n} = \frac{1}{4n^2} < \frac{1}{2n+1} = u_{2n+1} \text{ et } u_{2n+1} = \frac{1}{2n+1} > \frac{1}{4n^2 + 8n + 4} = u_{2n+2}$$

donc la suite (u_n) n'est ni croissante ni décroissante à partir d'un certain rang. $\square$

On pourrait cependant penser que si une suite positive (u_n) tend vers 0, le nombre des indices n pour lesquels $u_{n+1} \leq u_n$ est suffisamment important. Il n'en est rien.

6.7. Suite positive (u_n) de limite nulle telle que :

$$\lim_{n \rightarrow \infty} \frac{\text{card}\{p \in \llbracket 1, n \rrbracket \mid u_{p+1} \leq u_p\}}{n} = 0.$$

Nous introduisons la suite $(k_n)_{n \geq 1}$ dont le terme général k_n est la partie entière de $\sqrt{n}$. On a alors, pour tout entier $n \geq 1$, $n = k_n^2 + q_n$ où $q_n \in \llbracket 0, 2k_n \rrbracket$. Nous considérons la suite $(u_n)_{n \geq 1}$, de terme général :

$$u_n = \begin{cases} \frac{1}{n} = \frac{1}{k_n^2} & \text{si } q_n = 0, \\ \frac{1}{(k_n + 1)^2 - q_n} & \text{si } q_n \neq 0. \end{cases}$$

Par exemple $u_4 = \frac{1}{4}$, $u_5 = \frac{1}{8}$, $u_6 = \frac{1}{7}$, $u_7 = \frac{1}{6}$, $u_8 = \frac{1}{5}$ et $u_9 = \frac{1}{9}$.

Si n est un entier ≥ 2 , l'inégalité $u_n < u_{n-1}$ n'est vérifiée que si n est un carré. Pour tout $n \in \mathbb{N}^*$, il n'y a que k_n carrés plus petits que n , donc :

$$c_n = \text{card}\{p \in \llbracket 1, n \rrbracket \mid u_{p+1} \leq u_p\} = k_n \leq \sqrt{n}$$

ce qui montre que :

$$0 \leq \frac{c_n}{n} \leq \frac{\sqrt{n}}{n} = \frac{1}{\sqrt{n}}.$$

Par conséquent la suite $\left(\frac{c_n}{n}\right)$ converge vers 0. $\square$

■ Convergence au sens de Cesàro

THÉORÈME 6.4. — Théorème de Cesàro.

Si une suite $(x_n)_{n \geq 1}$ converge vers un réel λ , la suite $(y_n)_{n \geq 1}$, de terme général :

$$y_n = \frac{x_1 + x_2 + \cdots + x_n}{n},$$

converge dans $\mathbb{R}$ vers la même limite λ , et le résultat reste valable si (x_n) tend vers $+\infty$ ou vers $-\infty$.

La réciproque est vraie dans le cas où la suite (x_n) est monotone, mais elle devient fausse dans le cas général.

Quand la suite (y_n) converge on dit que la suite (x_n) converge au sens de Cesàro⁴.

6.8. Suite divergente qui converge au sens de Cesàro.

Nous considérons les suites $(x_n)_{n \geq 1}$ et $(y_n)_{n \geq 1}$, de termes généraux :

$$x_n = (-1)^n \text{ et } y_n = \frac{x_1 + x_2 + \cdots + x_n}{n}.$$

4. Cette notion est introduite par Georg Frobenius en 1880 pour donner une valeur à une série entière à la frontière du disque de convergence. Elle est généralisée par Otto Hölder, en 1882, puis par Ernesto Cesàro, en 1890.

Si $n \in \mathbb{N}^*$, alors $y_n = 0$ si n est pair et $y_n = \frac{1}{n}$ si n est impair, donc $0 \leq y_n \leq \frac{1}{n}$. Par conséquent la suite (y_n) converge vers 0, alors que la suite (x_n) diverge, comme établi dans l'exemple 6.1 (page 101). $\square$

6.9. Suite qui ne diverge pas vers $+\infty$ mais qui tend vers $+\infty$ au sens de Cesàro.

Nous considérons les suites $(x_n)_{n \geq 1}$ et $(y_n)_{n \geq 1}$, de termes généraux :

$$x_n = n(1 + (-1)^n) \text{ et } y_n = \frac{x_1 + x_2 + \cdots + x_n}{n}.$$

On a $x_{2n} = (2n) \times 2 = 4n$ et $x_{2n+1} = 0$ pour tout n . Ainsi la sous-suite (x_{2n+1}) de (x_n) est constante égale à 0, donc la suite (x_n) ne diverge pas vers $+\infty$.

De plus, pour tout entier $n \geq 1$:

$$y_{2n} = \frac{x_2 + x_4 + \cdots + x_{2n}}{2n} = \frac{4(1 + 2 + \cdots + n)}{2n} = \frac{n(n+1)}{n} = n+1$$

et :

$$y_{2n+1} = \frac{x_2 + x_4 + \cdots + x_{2n}}{2n+1} = \frac{4(1 + 2 + \cdots + n)}{2n+1} = \frac{2n(n+1)}{2n+1} > n,$$

donc les sous-suites (y_{2n}) et (y_{2n+1}) de (y_n) divergent toutes deux vers $+\infty$, ce qui montre que la suite (y_n) tend vers $+\infty$. En conclusion, la suite (x_n) tend vers $+\infty$ au sens de Cesàro. $\square$

■ Limite supérieure et limite inférieure

Une suite (x_n) de réels admet une limite dans la droite numérique achevée $\overline{\mathbb{R}}$ si elle converge dans $\mathbb{R}$ ou si elle diverge vers $+\infty$ ou $-\infty$, et alors, dans chacun des trois cas, sa limite — unique — dans $\overline{\mathbb{R}}$ est ce qui a été à chaque fois noté $\lim_{n \rightarrow \infty} x_n$.

Soit (x_n) une suite de réels. Nous introduisons pour tout n les bornes inférieure y_n et supérieure z_n dans $\overline{\mathbb{R}}$ de l'ensemble $X_n = \{x_p \mid p \geq n\}$, c'est-à-dire :

$$y_n = \inf_{p \geq n} x_p \text{ et } z_n = \sup_{p \geq n} x_p.$$

Pour tout n , $X_{n+1} \subset X_n$, donc $y_n \leq y_{n+1}$ et $z_{n+1} \leq z_n$. Par conséquent (y_n) et (z_n) sont des suites à valeurs dans $\overline{\mathbb{R}}$, (y_n) est croissante et (z_n) décroissante, donc les suites (y_n) et (z_n) ont toutes deux une limite dans $\overline{\mathbb{R}}$. La limite de (y_n) dans $\overline{\mathbb{R}}$ est appelée la limite inférieure et la limite de (z_n) dans $\overline{\mathbb{R}}$ la limite supérieure de la suite (x_n) , et on pose⁵ :

$$\liminf_{n \rightarrow \infty} x_n = \lim_{n \rightarrow \infty} y_n \text{ et } \limsup_{n \rightarrow \infty} x_n = \lim_{n \rightarrow \infty} z_n.$$

Pour tout n , x_n appartient à X_n , donc $y_n \leq x_n \leq z_n$, ce qui montre que $y_n \leq z_n$. Le passage à la limite donne l'inégalité $\liminf_{n \rightarrow \infty} x_n \leq \limsup_{n \rightarrow \infty} x_n$.

Si (x_n) est bornée dans $\mathbb{R}$, les limites inférieure et supérieure de (x_n) sont des nombres réels, et (x_n) converge dans $\mathbb{R}$ si, et seulement si, $\liminf_{n \rightarrow \infty} x_n = \limsup_{n \rightarrow \infty} x_n$.

5. Cauchy introduit ces notions dans son cours d'Analyse de 1821, les appelant la *plus petite des limites* et la *plus grande des limites* de la suite (x_n) ; voir [WALT], §4.15.

Pour des suites (x_n) et (y_n) , on dispose des inégalités :

$$(1) \quad \begin{aligned} \liminf_{n \rightarrow \infty} x_n + \liminf_{n \rightarrow \infty} y_n &\leq \liminf_{n \rightarrow \infty} (x_n + y_n) \leq \liminf_{n \rightarrow \infty} x_n + \limsup_{n \rightarrow \infty} y_n \\ &\leq \limsup_{n \rightarrow \infty} (x_n + y_n) \leq \limsup_{n \rightarrow \infty} x_n + \limsup_{n \rightarrow \infty} y_n \end{aligned}$$

et pour des suites convergentes, ces inégalités sont des égalités. De même, si les suites (x_n) et (y_n) sont à termes strictement positifs, on a :

$$(2) \quad \begin{aligned} \liminf_{n \rightarrow \infty} x_n \times \liminf_{n \rightarrow \infty} y_n &\leq \liminf_{n \rightarrow \infty} (x_n y_n) \leq \limsup_{n \rightarrow \infty} (x_n y_n) \\ &\leq \limsup_{n \rightarrow \infty} x_n \times \limsup_{n \rightarrow \infty} y_n. \end{aligned}$$

6.10. Suites pour lesquelles les inégalités (1) sont strictes.

Nous considérons les suites $(x_n)_{n \geq 0}$ et $(y_n)_{n \geq 0}$, de période 4 — ce qui signifie que $x_{n+4} = x_n$ et $y_{n+4} = y_n$ pour tout $n \in \mathbb{N}$ —, définies par : $x_0 = 2$, $x_1 = x_3 = 0$ et $x_2 = 1$ d'une part, et $y_0 = 0$, $y_1 = 1$ et $y_2 = y_3 = 2$ d'autre part. Alors la suite $(z_n)_{n \geq 0}$, de terme général $z_n = x_n + y_n$, est de période 4, $z_0 = z_3 = 2$, $z_1 = 1$ et $z_2 = 3$, donc :

$$\begin{cases} \liminf_{n \rightarrow \infty} x_n = \liminf_{n \rightarrow \infty} y_n = 0, & \limsup_{n \rightarrow \infty} x_n = \limsup_{n \rightarrow \infty} y_n = 2, \\ \liminf_{n \rightarrow \infty} (x_n + y_n) = 1, & \limsup_{n \rightarrow \infty} (x_n + y_n) = 3 \end{cases}$$

ce qui montre que toutes les inégalités de (1) sont strictes pour ces deux suites. $\square$

6.11. Suites pour lesquelles les inégalités (2) sont strictes.

Nous introduisons les suites $(x_n)_{n \geq 0}$ et $(y_n)_{n \geq 0}$, de période 3, définies par : $x_0 = 1$, $x_1 = 2$ et $x_2 = 3$ d'une part et $y_0 = 3$, $y_1 = 2$ et $y_2 = 1$ d'autre part. Alors la suite $(z_n)_{n \geq 0}$, de terme général $z_n = x_n y_n$, est de période 3, $z_0 = z_2 = 3$ et $z_1 = 4$, donc :

$$\begin{cases} \liminf_{n \rightarrow \infty} x_n = \liminf_{n \rightarrow \infty} y_n = 1, & \limsup_{n \rightarrow \infty} x_n = \limsup_{n \rightarrow \infty} y_n = 3, \\ \liminf_{n \rightarrow \infty} (x_n y_n) = 3, & \limsup_{n \rightarrow \infty} (x_n y_n) = 4 \end{cases}$$

ce qui montre que toutes les inégalités de (2) sont strictes pour ces deux suites. $\square$

L'inégalité :

$$\limsup_{n \rightarrow \infty} (x_n + y_n) \leq \limsup_{n \rightarrow \infty} x_n + \limsup_{n \rightarrow \infty} y_n$$

ne se généralise pas, lorsque toutes les sommes ont un sens, au cas de la somme d'une infinité dénombrable de suites.

6.12. Famille $((a_{k,n})_{n \in \mathbb{N}})_{k \in \mathbb{N}}$ de suites telle que :

$$\sum_{k=0}^{+\infty} \limsup_{n \rightarrow \infty} a_{k,n} < \limsup_{n \rightarrow \infty} \sum_{k=0}^{+\infty} a_{k,n}.$$

Nous introduisons la suite double $(a_{k,n})_{k,n \in \mathbb{N}}$ définie par $a_{k,n} = \begin{cases} 0 & \text{si } n \neq k, \\ 1 & \text{si } n = k. \end{cases}$

Pour tout $k \in \mathbb{N}$, $a_{k,n} = 0$ pour tout $n \geq k + 1$, donc $\limsup_{n \rightarrow \infty} a_{k,n} = 0$. Par suite :

$$\sum_{k=0}^{+\infty} \limsup_{n \rightarrow \infty} a_{k,n} = 0.$$

On a $\sum_{k=0}^{+\infty} \underbrace{a_{k,n}}_{=0 \text{ si } k \neq n} = a_{n,n} = 1$ pour tout entier naturel n , donc $\limsup_{n \rightarrow \infty} \sum_{k=0}^{+\infty} a_{k,n} = 1$. $\square$

On a clairement, pour une suite (x_n) de signe constant :

- $\liminf_{n \rightarrow \infty} |x_n| = \liminf_{n \rightarrow \infty} x_n$ si $x_n \geq 0$ pour tout n ,
- $\liminf_{n \rightarrow \infty} |x_n| = -\limsup_{n \rightarrow \infty} x_n$ si $x_n \leq 0$ pour tout n .

Dans le cas d'une suite à termes de signe quelconque, $\liminf_{n \rightarrow \infty} |x_n|$ peut n'être égale ni à $\liminf_{n \rightarrow \infty} x_n$, ni à $-\limsup_{n \rightarrow \infty} x_n$.

6.13. Suite (x_n) de nombres réels pour laquelle :

$$\liminf_{n \rightarrow \infty} |x_n| \neq \text{Max} \left(\liminf_{n \rightarrow \infty} x_n, \limsup_{n \rightarrow \infty} x_n \right).$$

Nous introduisons la suite $(x_n)_{n \geq 0}$, de période 3, définie par : $x_0 = 1$, $x_1 = 0$ et $x_2 = -1$. Si n est un entier naturel, $|x_n|$ vaut 0 si n est congru à 1 modulo 3 et 1 sinon, donc $\liminf_{n \rightarrow \infty} |x_n| = 0$. Or $\liminf_{n \rightarrow \infty} x_n = -1$ et $\limsup_{n \rightarrow \infty} x_n = 1$, ce qui donne :

$$\text{Max} \left(\liminf_{n \rightarrow \infty} x_n, \limsup_{n \rightarrow \infty} x_n \right) = \text{Max}(1, 1) = 1. \square$$

6.14. Suites (x_n) et (y_n) pour lesquelles $\{n \in \mathbb{N} \mid x_n \geq y_{n+1}\}$ est un ensemble infini et $\limsup_{n \rightarrow \infty} x_n < \limsup_{n \rightarrow \infty} y_n$.

Nous considérons les suites $(x_n)_{n \geq 0}$ et $(y_n)_{n \geq 0}$, de période 2, définies par : $x_0 = 1$ et $x_1 = 2$ d'une part et $y_0 = 4$ et $y_1 = 0$ d'autre part. Alors :

$$\limsup_{n \rightarrow \infty} x_n = 2 < \limsup_{n \rightarrow \infty} y_n = 4$$

et, pour tout entier naturel impair n , $x_n = 2 > y_{n+1} = 0$. $\square$

■ Valeurs d'adhérence

DÉFINITION 6.7. — Une valeur d'adhérence d'une suite (x_n) est un nombre réel adhérent, pour tout n , à l'ensemble $X_n = \{x_p \mid p \geq n\}$.

L'ensemble des valeurs d'adhérence d'une suite (x_n) est donc l'intersection de la suite $(\overline{X_n})$, suite de fermés non vides de $\mathbb{R}$ décroissante pour l'inclusion.

Si une suite (x_n) est bornée, l'ensemble $\overline{X_n}$ est, pour tout n , un compact non vide de $\mathbb{R}$, donc l'ensemble des valeurs d'adhérence de (x_n) est un compact non vide de $\mathbb{R}$, dont le plus grand élément est la limite supérieure et le plus petit élément la limite inférieure de la suite (x_n) .

THÉORÈME 6.5. — Un nombre réel est une valeur d'adhérence d'une suite (x_n) si, et seulement si, c'est la limite d'une sous-suite convergente de (x_n) .

THÉORÈME 6.6. — Une suite bornée converge dans $\mathbb{R}$ si, et seulement si, elle admet une valeur d'adhérence et une seule, qui est alors sa limite.

6.15. Suite divergente ayant une valeur d'adhérence et une seule.

Nous considérons la suite $(x_n)_{n \geq 0}$ de terme général $x_n = n(1 + (-1)^n)$.

On a, pour tout entier naturel n , $x_{2n} = 4n$ et $x_{2n+1} = 0$, donc la sous-suite (x_{2n}) de (x_n) tend vers $+\infty$, alors que sa sous-suite (x_{2n+1}) converge vers 0. Ainsi la suite (x_n) diverge et 0 est clairement la seule valeur d'adhérence de (x_n) . $\square$

6.16. Suite dont tout entier naturel est une valeur d'adhérence.

Nous introduisons la suite $(x_n)_{n \geq 1}$ d'entiers naturels définie ainsi : pour tout entier $n \geq 1$, x_n est l'exposant de 2 dans la décomposition de n en produit de facteurs premiers. Les premiers termes de la suite (x_n) sont donc :

$$0, 1, 0, 2, 0, 1, 0, 3, 0, 1, 0, 2, 0, \dots$$

Pour tout entier naturel k , $\varphi_k : n \mapsto \varphi_k(n) = 2^k 3^n$ est une application strictement croissante de $\mathbb{N}$ dans $\mathbb{N}$, donc la suite $(x_{2^k 3^n})_{n \geq 1}$ est une sous-suite de (x_n) , constante égale à k , qui converge donc vers k . Il en résulte que tout entier naturel k est une valeur d'adhérence de (x_n) , et il n'y en a pas d'autres puisque la suite (x_n) ne prend que des valeurs entières positives ou nulles et que $\mathbb{N}$ est un fermé de $\mathbb{R}$. L'ensemble des valeurs d'adhérence de la suite (x_n) est donc exactement $\mathbb{N}$. $\square$

Nous avons vu que l'ensemble des valeurs d'adhérence d'une suite peut être infini ; nous constatons qu'il peut être infini et non dénombrable.

6.17. Suite dont l'ensemble des valeurs d'adhérence est le segment $[-1, 1]$.

Nous introduisons la suite $(x_n)_{n \geq 0}$ de terme général $x_n = \cos n$ et le sous-groupe $G = \{p + 2k\pi \mid p \in \mathbb{Z} \text{ et } k \in \mathbb{Z}\}$ du groupe additif $(\mathbb{R}, +)$ de $\mathbb{R}$. Comme π est irrationnel, les propriétés des sous-groupes additifs de $\mathbb{R}$ (théorème 5.2, page 90) montrent que l'adhérence $\overline{G}$ de G est égale à $\mathbb{R}$. Cosinus étant une application continue de $\mathbb{R}$ dans $\mathbb{R}$ et $[-1, 1]$ étant fermé, on a :

$$\cos(\overline{G}) \subset \overline{\cos(G)} \subset [-1, 1].$$

Or $\overline{G} = \mathbb{R}$ et $\cos(\mathbb{R}) = [-1, 1]$, donc $\overline{\cos(G)} = [-1, 1]$. Par conséquent l'adhérence de l'ensemble $\mathcal{A} = \{\cos p \mid p \in \mathbb{Z}\}$ est égale au segment $[-1, 1]$. De plus cosinus est paire, donc :

$$\mathcal{A} = \{\cos p \mid p \in \mathbb{N}\} = \{x_p \mid p \in \mathbb{N}\}.$$

Nous posons, pour tout entier naturel n , $X_n = \{x_p \mid p \geq n\}$ et $Y_n = \{x_p \mid p < n\}$ et nous démontrons que $\overline{X_n} = [-1, 1]$ pour tout $n \in \mathbb{N}$.

Soit n un entier naturel. L'ensemble Y_n étant fini, c'est un fermé donc $\overline{Y_n} = Y_n$. Comme $\mathcal{A} = X_n \cup Y_n$, on a :

$$[-1, 1] = \overline{\mathcal{A}} = \overline{X_n \cup Y_n} = \overline{X_n} \cup \overline{Y_n} = \overline{X_n} \cup Y_n,$$

ce qui montre que $\overline{X_n}$ est le segment $[-1, 1]$ privé éventuellement d'un ensemble fini Z_n . Supposons que l'ensemble Z_n ne soit pas vide. Nous choisissons un élément

a de Z_n . Pour tout nombre réel $\varepsilon > 0$, $]a - \varepsilon, a + \varepsilon[$ est infini donc contient des éléments de $\overline{X_n}$, d'où l'on déduit que a appartient à l'adhérence de $\overline{X_n}$, donc à $\overline{X_n}$, ce qui est contradictoire. Par conséquent Z_n est vide, donc $\overline{X_n} = [-1, 1]$.

L'ensemble des valeurs d'adhérence de la suite (x_n) étant l'intersection de la suite $(\overline{X_n})_{n \in \mathbb{N}}$, c'est le segment $[-1, 1]$. $\square$

6.18. Suite de nombres rationnels dont tout nombre réel est une valeur d'adhérence.

Nous notons, pour tout entier $n \geq 1$ et tout nombre premier p , $\nu_p(n)$ l'exposant de p dans la décomposition de n en produit de facteurs premiers, et nous considérons la suite $(x_n)_{n \geq 1}$, à valeurs dans $\mathbb{Q}$, de terme général :

$$x_n = (-1)^{\nu_5(n)} \frac{\nu_2(n)}{1 + \nu_3(n)}.$$

Soit r un nombre rationnel non nul. Si $r > 0$ et si (p, q) est le représentant irréductible de dénominateur positif de r , nous définissons l'application :

$$\left| \begin{array}{l} \varphi : \mathbb{N} \longrightarrow \mathbb{N} \\ n \longmapsto \varphi(n) = 2^p \times 3^{q-1} \times 7^n \end{array} \right.$$

de $\mathbb{N}$ dans $\mathbb{N}$, qui est clairement strictement croissante, et alors $x_{\varphi(n)} = r$ pour tout entier $n \geq 1$, donc la suite $(y_n)_{n \geq 1}$, de terme général $y_n = x_{\varphi(n)}$, est une sous-suite de (x_n) , constante égale à r . On opère de même pour un rationnel $r < 0$ en notant $(-p, q)$ son représentant irréductible de dénominateur positif et en définissant l'application $\varphi : n \mapsto \varphi(n) = 2^p \times 3^{q-1} \times 5 \times 7^n$ de $\mathbb{N}$ dans $\mathbb{N}$.

Il en résulte que tout nombre rationnel r différent de 0 est une valeur d'adhérence de la suite (x_n) . L'ensemble des valeurs d'adhérence de la suite (x_n) étant fermé, il contient l'adhérence de $\mathbb{Q} \setminus \{0\}$. Or $\overline{\mathbb{Q} \setminus \{0\}} = \mathbb{R}$, donc l'ensemble des valeurs d'adhérence de la suite (x_n) est $\mathbb{R}$. $\square$

6.19. Suite de nombres réels dont tout nombre réel est une valeur d'adhérence.

Nous notons f l'application réciproque de la fonction tangente hyperbolique $\tanh$. La fonction f est définie et continue sur l'intervalle ouvert $] -1, 1[$, et c'est une bijection de $] -1, 1[$ sur $\mathbb{R}$.

Nous considérons la suite $(x_n)_{n \geq 0}$ de terme général $x_n = \cos n$. Le nombre réel π étant irrationnel, aucun entier $n \geq 1$ n'appartient à $\{k\pi \mid k \in \mathbb{N}\}$. Par conséquent, pour tout entier $n \geq 1$, x_n appartient à $] -1, 1[$, ce qui justifie l'introduction de la suite $(y_n)_{n \geq 1}$ de terme général $y_n = f(x_n)$.

Soit ℓ un nombre réel. Nous posons $c = \tanh \ell$; alors c appartient à $] -1, 1[$ et $\ell = f(c)$. Comme c est une valeur d'adhérence de la suite (x_n) (exemple 6.17, pages 108 et 109) c est la limite d'une sous-suite convergente (u_n) de (x_n) . Il existe une application φ strictement croissante de $\mathbb{N}$ dans $\mathbb{N}$ telle que $u_n = x_{\varphi(n)}$ pour tout entier naturel n . La fonction f est continue en c et, pour tout entier $n \geq 1$, $y_{\varphi(n)} = f(x_{\varphi(n)}) = f(u_n)$, donc la sous-suite $(y_{\varphi(n)})$ de la suite (y_n) converge dans $\mathbb{R}$ vers $\ell = f(c)$, ce qui montre que ℓ est une valeur d'adhérence de (y_n) .

En conclusion, tout nombre réel est une valeur d'adhérence de la suite (y_n) . $\square$

Chapitre 7

Séries numériques

L'utilisation de la somme d'une série apparaît dès l'Antiquité avec Archimède pour les calculs d'aires et de volumes. Il faut attendre le XVII^e siècle pour voir se développer l'étude des séries, en particulier avec le développement du calcul différentiel et intégral et l'utilisation de la formule de Taylor, et c'est Augustin-Louis Cauchy qui établit le premier une théorie rigoureuse en 1821.

Dans tout le chapitre, (u_n) , (v_n) , (w_n) , (x_n) , (y_n) ... sont des suites numériques, c'est-à-dire des suites à valeurs réelles.

■ Convergence et convergence absolue

DÉFINITION 7.1. — La série de terme général u_n , que l'on appelle aussi la série $\sum u_n$, $\sum_{n \geq 0} u_n$ ou encore :

$$\sum_{n \geq 0} u_n,$$

converge si la suite (S_n) des sommes partielles de la série $\sum u_n$, de terme général :

$$S_n = \sum_{k=0}^n u_k = u_0 + u_1 + \cdots + u_n,$$

converge dans $\mathbb{R}$ et, si la série $\sum u_n$ converge, la limite de la suite (S_n) s'appelle la somme de la série et se note :

$$\sum_{n=0}^{+\infty} u_n.$$

Bien sûr, il peut s'agir d'une suite $(u_n)_{n \in \mathbb{N}^*} = (u_n)_{n \geq 1}$ définie à partir du rang 1. Dans ce cas, la série de terme général u_n , appelée aussi la série $\sum u_n$ ou $\sum_{n \geq 1} u_n$ converge si la suite $(S_n)_{n \geq 1}$ des sommes partielles de la série $\sum u_n$, définie à partir du rang 1 par :

$$S_n = \sum_{k=1}^n u_k = u_1 + \cdots + u_n,$$

converge et, si la série $\sum_{n \geq 1} u_n$ converge, sa somme est le réel $\sum_{n=1}^{+\infty} u_n = \lim_{n \rightarrow \infty} S_n$.

DÉFINITION 7.2. — Une série diverge si elle ne converge pas.

THÉORÈME 7.1. — Si la série $\sum u_n$ converge, alors la suite (u_n) tend vers 0.

En effet, avec les notations de la définition 7.1, $u_n = S_n - S_{n-1}$ pour tout entier $n \geq 1$ et la suite (S_n) converge dans $\mathbb{R}$. Cependant la réciproque est fausse.

7.1. Suite (u_n) qui converge vers 0 alors que la série $\sum u_n$ diverge.

Nous considérons la suite $(u_n)_{n \geq 1}$ de terme général $u_n = \frac{1}{n}$ et la suite $(S_n)_{n \geq 1}$ des sommes partielles de $\sum_n u_n$, de terme général :

$$S_n = \sum_{p=1}^n u_p = \sum_{p=1}^n \frac{1}{p}.$$

La suite (u_n) converge vers 0 et on a, pour tout entier $n \geq 1$:

$$S_{2n} - S_n = \sum_{p=n+1}^{2n} \underbrace{\frac{1}{p}}_{\geq \frac{1}{2n}} \geq n \times \frac{1}{2n} = \frac{1}{2},$$

donc la suite (S_n) ne vérifie pas le critère de Cauchy (voir l'exemple 6.3, pages 102 et 103), ce qui montre que la série $\sum_{n \geq 1} u_n$ diverge. $\square$

DÉFINITION 7.3. — La série divergente $\sum_{n \geq 1} \frac{1}{n}$ est appelée la série harmonique.

Nous précisons le résultat de l'exemple précédent 7.1 en rappelant la définition de la constante d'Euler. Pour ceci nous considérons la suite $(S_n)_{n \geq 1}$ de terme général :

$$(1) \quad S_n = \sum_{p=1}^n \frac{1}{p}$$

— voir l'exemple précédent 7.1 — et nous démontrons qu'il existe une constante réelle positive γ — c'est la constante d'Euler — et une suite (ε_n) qui converge vers 0 telles que, pour tout entier $n \geq 1$:

$$(2) \quad S_n = \ln(n) + \gamma + \varepsilon_n.$$

Nous introduisons la suite $(a_n)_{n \geq 1}$ de terme général $a_n = S_n - \ln(n)$. La fonction $x \mapsto 1/x$ étant continue et strictement décroissante sur $]0, +\infty[$, on a, pour tout entier $n \geq 1$:

$$\frac{1}{n+1} < \int_n^{n+1} \frac{1}{x} dx < \frac{1}{n},$$

ce qui s'écrit :

$$(3) \quad \frac{1}{n+1} < \ln \frac{n+1}{n} < \frac{1}{n}.$$

On déduit de (3) que, pour tout entier $n \geq 1$:

$$a_{n+1} - a_n = S_{n+1} - \ln(n+1) - S_n + \ln(n) = \frac{1}{n+1} - \ln \frac{n+1}{n} < 0$$

donc $a_{n+1} < a_n$. On a, pour tout $n \in \mathbb{N}^*$:

$$\ln(n) = \sum_{p=1}^{n-1} (\ln(p+1) - \ln(p)) = \sum_{p=1}^{n-1} \ln \frac{p+1}{p}$$

donc :

$$a_n = \sum_{p=1}^n \frac{1}{p} - \ln(n) = \frac{1}{n} + \underbrace{\sum_{p=1}^{n-1} \left(\frac{1}{p} - \ln \frac{p+1}{p} \right)}_{> 0 \text{ (voir (3))}} \geq \frac{1}{n} > 0.$$

La suite $(a_n)_{n \geq 1}$, décroissante et minorée par 0, converge dans $\mathbb{R}$ et sa limite γ est un réel positif ou nul (on a $\gamma = 0,577215665 \dots$). Il ne reste plus qu'à poser $(\varepsilon_n)_{n \geq 1} = (a_n - \gamma)_{n \geq 1}$ pour conclure. $\square$

On déduit de la formule (2) de la page précédente que la suite $(S_n)_{n \geq 1}$ des sommes partielles de la série harmonique est équivalente à la suite $(\ln n)$, ce qui redémontre que la série harmonique diverge.

A l'aide du critère de Cauchy pour les suites (définition 6.6, page 103), on prouve que si la série $\sum |u_n|$ converge, alors la série $\sum u_n$ converge.

DÉFINITION 7.4. — Une série $\sum u_n$ est absolument convergente si la série $\sum |u_n|$ converge, semi-convergente si elle converge alors que la série $\sum |u_n|$ diverge.

L'intérêt de la notion de série absolument convergente provient de la simplicité des critères de convergence des séries à termes positifs — les suites des sommes partielles associées sont alors croissantes — et au fait que la convergence absolue entraîne la convergence.

7.2. Série qui converge mais qui ne converge pas absolument.

Nous considérons la suite $(v_n)_{n \geq 1}$ de terme général $v_n = \frac{(-1)^{n+1}}{n}$.

La série $\sum |v_n|$ est la série harmonique, donc elle diverge. Nous notons $(T_n)_{n \geq 1}$ la suite des sommes partielles de la série $\sum v_n$, de terme général :

$$T_n = \sum_{k=1}^n v_k = \sum_{k=1}^n \frac{(-1)^{k+1}}{k}.$$

Nous introduisons les sous-suites $(x_n)_{n \geq 1}$ et $(y_n)_{n \geq 1}$ de la suite $(T_n)_{n \geq 1}$, de termes généraux $x_n = T_{2n}$ et $y_n = T_{2n+1}$. On a, pour tout entier $n \geq 1$:

$$y_n - x_n = \frac{(-1)^{2n+2}}{2n+1} = \frac{1}{2n+1},$$

$$x_{n+1} - x_n = T_{2n+2} - T_{2n} = \frac{1}{2n+1} - \frac{1}{2n+2} > 0$$

et :

$$y_{n+1} - y_n = T_{2n+3} - T_{2n+1} = -\frac{1}{2n+2} + \frac{1}{2n+3} < 0.$$

Par conséquent $(x_n)_{n \geq 1}$ est croissante et $(y_n)_{n \geq 1}$ décroissante et la suite $(y_n - x_n)$ converge vers 0, ce qui montre que les suites (x_n) et (y_n) sont adjacentes. Ainsi les sous-suites (T_{2n}) et (T_{2n+1}) de (T_n) convergent dans $\mathbb{R}$ vers la même limite, donc la suite (T_n) converge, d'où l'on déduit que la série $\sum v_n$ converge — sa somme est égale à $\ln 2$; voir l'exemple 7.22, page 123. En conclusion, $\sum v_n$ est une série semi-convergente. $\square$

Dans les exemples précédents 7.1 et 7.2, on pouvait conclure plus rapidement avec le critère de Riemann et celui des séries alternées.

THÉORÈME 7.2. — Critère de Riemann.

Soit α un nombre réel. La série $\sum_{n \geq 1} \frac{1}{n^\alpha}$ converge si, et seulement si, $\alpha > 1$.

THÉORÈME 7.3. — Critère des séries alternées¹ (ou de Leibniz).

Si $v_n \geq 0$ pour tout n et si les deux conditions suivantes sont réalisées :

- (I) la suite (v_n) est décroissante,
- (II) la suite (v_n) converge vers 0,

alors la série $\sum (-1)^n v_n$ converge.

Compte tenu du théorème 7.1 (page 111) la deuxième condition est nécessaire.

THÉORÈME 7.4. — Si on a $0 \leq u_n \leq v_n$ pour tout n , la convergence de la série $\sum v_n$ entraîne celle de $\sum u_n$.

Il suffit bien sûr que les inégalités $0 \leq u_n \leq v_n$ soit vraies à partir d'un certain rang.

Ceci devient faux pour des séries à termes de signes quelconques, même si l'on suppose que $|u_n| < |v_n|$ pour tout n .

7.3. Suites (u_n) et (v_n) telles que $|u_n| < |v_n|$ pour tout n , et pour lesquelles la série $\sum v_n$ converge alors que $\sum u_n$ diverge.

Nous considérons les suites $(u_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$ de termes généraux :

$$u_n = \frac{1}{n+1} \text{ et } v_n = \frac{(-1)^{n+1}}{n}.$$

Par le critère des séries alternées, la série $\sum v_n$ converge, alors que la série $\sum u_n$ diverge — en effet, (T_n) étant la suite des sommes partielles de la série harmonique, le terme général de la suite des sommes partielles de la série $\sum u_n$ est $S_n = T_{n+1} - 1$, donc (S_n) ne converge pas dans $\mathbb{R}$. De plus, pour tout entier $n \geq 1$:

$$|u_n| = \frac{1}{n+1} < \frac{1}{n} = |v_n|. \quad \square$$

7.4. Suites (u_n) et (v_n) pour lesquelles $|u_n| = o_{n \rightarrow \infty}(|v_n|)$, la série $\sum v_n$ converge et la série $\sum u_n$ diverge.

Nous introduisons les suites $(u_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$, de termes généraux :

$$u_n = \frac{1}{n} \text{ et } v_n = \frac{(-1)^{n+1}}{\sqrt{n}}.$$

La série harmonique $\sum u_n$ diverge et, par le critère des séries alternées (ou de Leibniz), la série $\sum v_n$ converge. Pourtant :

$$|u_n| = u_n = \frac{1}{n} = o_{n \rightarrow \infty}\left(\frac{1}{\sqrt{n}}\right) = o_{n \rightarrow \infty}(|v_n|). \quad \square$$

Dans l'exemple précédent 7.4, nous avons utilisé le fait que pour certaines valeurs de l'entier n , les réels u_n et v_n ne sont pas de même signe.

1. Ce critère est introduit par Leibniz en 1682 ; pour cette raison, on l'appelle aussi *critère de Leibniz*. On le trouve bien sûr en 1821 dans l'*Analyse algébrique* de Cauchy ([CAU1], chapitre VI, §3°, 3° théorème, page 144).

7.5. Suites (u_n) et (v_n) telles que u_n et v_n sont de même signe pour tout n , $|u_n| = o_{n \rightarrow \infty}(|v_n|)$, $\sum v_n$ converge et $\sum u_n$ diverge.

Nous considérons les suites $(u_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$ de termes généraux :

$$u_n = \begin{cases} \frac{1}{n} & \text{si } n \text{ est pair,} \\ -\frac{1}{n^2} & \text{si } n \text{ est impair} \end{cases} \quad \text{et } v_n = \frac{(-1)^n}{\sqrt{n}}.$$

Par le critère des séries alternées, la série $\sum v_n$ converge, et il est clair que u_n et v_n sont de même signe pour tout n .

Par ailleurs, pour tout entier $n \geq 1$, $|u_n| \leq \frac{1}{n} = \frac{1}{\sqrt{n}} |v_n|$, donc $|u_n| = o_{n \rightarrow \infty}(|v_n|)$.

Nous introduisons la suite $(S_n)_{n \geq 1}$ des sommes partielles de la série $\sum u_n$. On a, pour tout entier $n \geq 1$:

$$S_{2n} = -1 + \frac{1}{2} - \frac{1}{9} + \frac{1}{4} - \frac{1}{25} + \frac{1}{6} + \cdots - \frac{1}{(2n-1)^2} + \frac{1}{2n} = T_n - U_n$$

où :

$$T_n = \frac{1}{2} \left(1 + \frac{1}{2} + \frac{1}{3} + \cdots + \frac{1}{n} \right) \quad \text{et} \quad U_n = 1 + \frac{1}{9} + \frac{1}{25} + \cdots + \frac{1}{(2n-1)^2}.$$

La suite (T_n) tend vers $+\infty$ (série harmonique). De plus, pour tout entier $n \geq 3$:

$$U_n = 1 + \frac{1}{9} + \sum_{p=3}^n \underbrace{\frac{1}{(2p-1)^2}}_{\leq 1/(2p-2)^2} \leq \frac{10}{9} + \frac{1}{4} \sum_{k=2}^{n-1} \frac{1}{k^2} < \frac{31}{36} + \frac{1}{4} V_n \quad \text{où } V_n = \sum_{k=1}^n \frac{1}{k^2}.$$

La série $\sum_{n \geq 1} (1/n^2)$ converge (critère de Riemann), donc la suite (V_n) converge dans $\mathbb{R}$. Par conséquent (V_n) est majorée dans $\mathbb{R}$, donc aussi (U_n) , et comme (U_n) est croissante, (U_n) converge dans $\mathbb{R}$. Ainsi la sous-suite (S_{2n}) de (S_n) tend vers $+\infty$. En conclusion, (S_n) ne converge pas dans $\mathbb{R}$, donc la série $\sum u_n$ diverge. $\square$

■ Mise en défaut de certains critères de convergence

THÉORÈME 7.5. — Si, à partir d'un certain rang, $u_n \geq 0$ et si la suite (v_n) est équivalente à la suite (u_n) , alors $v_n \geq 0$ à partir d'un certain rang et les séries $\sum u_n$ et $\sum v_n$ sont de même nature — ce qui signifie que toutes les deux convergent ou bien que toutes les deux divergent.

Ceci devient faux pour des séries à termes de signes quelconques.

7.6. Suites équivalentes (u_n) et (v_n) telles que $\sum u_n$ diverge alors que $\sum v_n$ converge.

Nous considérons les suites $(u_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$, de termes généraux :

$$u_n = \frac{(-1)^n}{\sqrt{n}} + \frac{1}{n} \quad \text{et} \quad v_n = \frac{(-1)^n}{\sqrt{n}}.$$

Pour tout entier $n \geq 1$, $u_n = v_n \left(1 + \frac{(-1)^n}{\sqrt{n}} \right)$; or $\lim_{n \rightarrow \infty} \frac{(-1)^n}{\sqrt{n}} = 0$, donc $u_n \underset{n \rightarrow \infty}{\sim} v_n$.

Par le critère des séries alternées, la série $\sum v_n$ converge. La série harmonique $\sum (1/n)$ diverge et, pour tout entier $n \geq 1$, $u_n = v_n + (1/n)$, donc la série $\sum u_n$ diverge comme somme d'une série convergente et d'une série divergente. $\square$

Pour une suite (u_n) à termes positifs, les séries $\sum u_n$ et $\sum \ln(1 + u_n)$ sont de même nature, ce qui prouve, par passage au logarithme, l'équivalence entre la convergence de la série $\sum u_n$ et celle du produit infini $\prod_{n=0}^{+\infty} (1 + u_n)$.

7.7. Suite (u_n) telle que $\sum u_n$ converge et $\sum \ln(1 + u_n)$ diverge.

Nous définissons la suite $(u_n)_{n \geq 1}$ par $u_1 = 0$ et, pour $n \geq 2$, $u_n = \frac{(-1)^n}{\sqrt{n}}$.

Par le critère des séries alternées, la série $\sum u_n$ converge.

La suite (u_n) tend vers 0 et $\ln(1 + x) = x - \frac{1}{2}x^2 + O(x^3)$, donc :

$$(1) \quad \ln(1 + u_n) = u_n - \frac{1}{2}(u_n)^2 + o((u_n)^3) = u_n - \frac{1}{2} \times \frac{1}{n} + v_n$$

pour une suite (v_n) dominée par la suite $\left(\frac{1}{n\sqrt{n}}\right)$.

Par le critère de Riemann (avec $\alpha = 3/2 > 1$), la série $\sum_{n \geq 1} \frac{1}{n\sqrt{n}}$ converge.

Il existe une constante réelle positive M telle que, à partir d'un certain rang :

$$|v_n| \leq \frac{M}{n\sqrt{n}},$$

donc la série $\sum v_n$ converge (absolument). De plus la série $\sum u_n$ converge et la série harmonique diverge, donc (1) montre que la série $\sum \ln(1 + u_n)$ diverge. $\square$

7.8. Suite (u_n) telle que $\sum u_n$ diverge et $\sum \ln(1 + u_n)$ converge.

Nous considérons la suite $(u_n)_{n \geq 1}$ de terme général $u_n = -1 + \exp\left(\frac{(-1)^n}{\sqrt{n}}\right)$. On a, pour tout entier $n \geq 1$:

$$\ln(1 + u_n) = \frac{(-1)^n}{\sqrt{n}}$$

donc la série $\sum \ln(1 + u_n)$ converge (critère de Leibniz). La suite (v_n) , de terme général $v_n = (-1)^n/\sqrt{n}$, tend vers 0 et $\exp(x) = 1 + x + \frac{1}{2}x^2 + O(x^3)$, donc :

$$(2) \quad u_n = v_n + \frac{1}{2}(v_n)^2 + o((v_n)^3) = v_n + \frac{1}{2} \times \frac{1}{n} + w_n$$

pour une suite (w_n) dominée par la suite $\left(\frac{1}{n\sqrt{n}}\right)$.

En raisonnant comme dans l'exemple 7.7, les séries $\sum v_n$ et $\sum w_n$ convergent ; or la série harmonique diverge, donc (2) montre que la série $\sum u_n$ diverge. $\square$

7.9. Série divergente $\sum (-1)^n v_n$ où $v_n \geq 0$ pour tout n et où la suite (v_n) converge vers 0.

Nous reprenons la suite $(u_n)_{n \geq 1}$ définie dans l'exemple 7.5 (page 114). On a, pour tout entier $n \geq 1$:

$$u_n = (-1)^n v_n \quad \text{où} \quad v_n = \begin{cases} \frac{1}{n} & \text{si } n \text{ est pair,} \\ \frac{1}{n^2} & \text{si } n \text{ est impair.} \end{cases}$$

Seule la condition (i) du critère des séries alternées n'est pas vérifiée par (v_n) , et nous avons vu dans l'exemple 7.5 que la série $\sum u_n$ diverge. $\square$

Dans le critère des séries alternées (ou de Leibniz), on ne peut remplacer la suite (v_n) par une suite équivalente comme le montre l'exemple suivant.

7.10. Suites (v_n) et (w_n) à termes positifs, équivalentes et telles que la série $\sum (-1)^n w_n$ vérifie le critère des séries alternées, alors que la série $\sum (-1)^n v_n$ diverge.

Nous considérons les suites (v_n) et (w_n) , de termes généraux :

$$v_n = \frac{1}{\sqrt{n}} + \frac{(-1)^n}{n} \text{ et } w_n = \frac{1}{\sqrt{n}}.$$

Pour tout entier $n \geq 2$, $0 < n < n^2$, donc $0 < \sqrt{n} < n$, ce qui montre que $v_1 = 0$ et que $v_n > 0$ pour tout entier $n \geq 2$; de plus, $w_n > 0$ pour tout n .

La suite de terme général $(-1)^n v_n$ est la suite (u_n) de l'exemple 7.6 (pages 114 et 115), donc la série $\sum (-1)^n v_n$ diverge (voir cet exemple). Enfin les suites (v_n) et (w_n) sont équivalentes (même démonstration que dans l'exemple 7.6) et la série $\sum (-1)^n w_n$ vérifie clairement le critère des séries alternées. $\square$

THÉORÈME 7.6. — Règle d'Abel.

Si (ε_n) et (v_n) sont des suites à valeurs réelles, si (u_n) et (S_n) sont les suites de termes généraux :

$$u_n = \varepsilon_n v_n \text{ et } S_n = \sum_{k=0}^n \varepsilon_k,$$

si (S_n) est bornée dans $\mathbb{R}$, si (v_n) converge vers 0 et si la série $\sum |v_{n+1} - v_n|$ converge, alors la série $\sum u_n$ converge.

Bien sûr les suites peuvent être définies à partir du rang 1, et dans ce cas (S_n) est la suite de terme général $S_n = \sum_{k=1}^n \varepsilon_k$.

En prenant la suite (ε_n) de terme général $\varepsilon_n = (-1)^n$ et en supposant que la suite (v_n) est décroissante et de limite nulle, on retrouve le critère des séries alternées comme cas particulier de la règle d'Abel.

7.11. Série vérifiant la règle d'Abel mais ne vérifiant pas le critère des séries alternées.

Nous utilisons la suite $(u_n)_{n \geq 1}$ de terme général $u_n = \frac{(-1)^{n+1}}{n + (-1)^{n+1}}$.

Nous montrons que la série $\sum u_n$ vérifie la règle d'Abel pour les suites (ε_n) et (v_n) de termes généraux :

$$\varepsilon_n = (-1)^{n+1} \text{ et } v_n = \frac{1}{n + (-1)^{n+1}}.$$

Pour tout entier $n \geq 1$, $S_n = \sum_{k=1}^n \varepsilon_k$ est égal à 0 ou à 1 suivant que n est pair ou impair, donc la suite (S_n) est bornée dans $\mathbb{R}$.

Pour tout $n \in \mathbb{N}^*$:

$$\begin{aligned} v_{n+1} - v_n &= \frac{1}{n+1+(-1)^n} - \frac{1}{n+(-1)^{n+1}} = \frac{n+(-1)^{n+1} - n-1-(-1)^n}{(n+1+(-1)^n)(n+(-1)^{n+1})} \\ &= \frac{-1+2(-1)^{n+1}}{(n+1+(-1)^n)(n+(-1)^{n+1})} \text{ et } -1+2(-1)^{n+1} = 1 \text{ ou } -3, \end{aligned}$$

$$\text{donc } |v_{n+1} - v_n| \leq 3w_n \text{ où } w_n = \frac{1}{(n+1+(-1)^n)(n+(-1)^{n+1})}.$$

Or $(n+1+(-1)^n) \underset{n \rightarrow \infty}{\sim} n$ et $(n+(-1)^{n+1}) \underset{n \rightarrow \infty}{\sim} n$, donc $w_n \underset{n \rightarrow \infty}{\sim} \frac{1}{n^2}$.

Par le critère de Riemann, la série $\sum (1/n^2)$ converge, donc aussi la série $\sum w_n$ (théorème 7.5, page 114), ce qui montre que la série $\sum |v_{n+1} - v_n|$ converge.

Comme la suite (v_n) converge vers 0, la règle d'Abel est vérifiée². Cependant la suite (v_n) n'est pas décroissante, puisque, pour tout entier $p \geq 1$:

$$v_{2p} = \frac{1}{2p-1} > \frac{1}{2p} = v_{2p-1}. \quad \square$$

■ Séries à termes positifs

THÉORÈME 7.7. — Si une suite (u_n) est décroissante et à termes positifs et si la série $\sum u_n$ converge, la suite (nu_n) converge vers 0.

D'une part la réciproque est fausse, et d'autre part la décroissance de la suite (u_n) est nécessaire comme le montrent les exemples suivants.

7.12. Suite (u_n) décroissante à termes positifs telle que la suite (nu_n) converge vers 0 alors que la série $\sum u_n$ diverge.

Nous considérons la suite $(u_n)_{n \geq 2}$ de terme général $u_n = \frac{1}{n \ln n}$.

La suite des sommes partielles³ de la série $\sum_{n \geq 2} u_n$ est la suite $(S_n)_{n \geq 2}$ de terme général :

$$S_n = \sum_{k=2}^n u_k.$$

Pour tout entier $n \geq 2$, $\ln n < \ln(n+1)$ donc $-\ln n > -\ln(n+1)$, ce donne :

$$(n+1) \ln(n+1) - n \ln n > (n+1) \ln(n+1) - n \ln(n+1) = \ln(n+1) > 0.$$

Il en résulte que (u_n) est strictement décroissante. On a, pour tout entier $n \geq 2$:

$$S_{n^2} - S_n = \sum_{k=n+1}^{n^2} \underbrace{u_k}_{\geq \frac{1}{k \ln(n^2)}} \geq \frac{1}{2 \ln n} \sum_{k=n+1}^{n^2} \frac{1}{k} = \frac{v_n}{2 \ln n} \text{ où } v_n = \sum_{k=n+1}^{n^2} \frac{1}{k}.$$

Nous notons $(H_n)_{n \geq 1}$ la suite des sommes partielles de la série harmonique. En désignant par γ la constante d'Euler, il existe une suite $(\varepsilon_n)_{n \geq 1}$ convergeant

2. On peut directement démontrer la convergence de cette série en effectuant un développement limité à l'ordre 1 et en utilisant le critère des séries alternées.

3. On peut ici utiliser directement le critère de Bertrand : la série $\sum_{n \geq 2} \frac{1}{n(\ln n)^\alpha}$ converge si, et seulement si, $\alpha > 1$.

vers 0 telle que, pour tout entier $n \geq 1$, $H_n = \ln(n) + \gamma + \varepsilon_n$ — voir, dans ce qui suit la définition 7.3, la formule (2) de la page 111. On a, pour tout entier $n \geq 1$:

$$v_n = H(n^2) - H(n) = \ln(n^2) + \gamma + \varepsilon_{n^2} - \ln(n) - \gamma - \varepsilon_n = \ln(n) + \omega_n$$

où $\omega_n = \varepsilon_{n^2} - \varepsilon_n$. Par conséquent la suite de terme général $v_n/(2 \ln n)$ converge vers $1/2$, donc, à partir d'un certain rang, $v_n/(2 \ln n) > 1/4$ donc $S_{n^2} - S_n > 1/4$, ce qui montre que (S_n) n'est pas une suite de Cauchy. On en déduit que la suite (S_n) ne converge pas dans $\mathbb{R}$, donc que la série $\sum u_n$ diverge, alors que $u_n > 0$ pour tout $n \geq 2$ et que la suite (nu_n) , de terme général $1/(\ln n)$, converge vers 0. $\square$

7.13. Suite (u_n) à termes positifs telle que (nu_n) ne converge pas vers 0 alors que la série $\sum u_n$ converge.

Nous considérons la suite $(u_n)_{n \geq 0}$ de terme général :

$$u_n = \begin{cases} \frac{1}{n} & \text{s'il existe } q \in \mathbb{N} \text{ tel que } n = 2^q, \\ 0 & \text{sinon} \end{cases}$$

et la suite $(v_n)_{n \geq 0}$ de terme général $v_n = nu_n$. Pour tout entier naturel n , $u_n \geq 0$, donc la suite $(S_n)_{n \geq 0}$ des sommes partielles de la série $\sum u_n$ est croissante.

Pour tout $n \in \mathbb{N}$, $n < 2^n$ donc $S_n \leq S_{2^n} = 1 + \frac{1}{2} + \frac{1}{4} + \cdots + \frac{1}{2^n} = \frac{1 - \frac{1}{2^{n+1}}}{1 - \frac{1}{2}} < 2$.

La suite (S_n) , croissante et majorée dans $\mathbb{R}$, converge dans $\mathbb{R}$, donc la série $\sum u_n$ converge. Nous posons $\mathcal{A} = \{2^k \mid k \in \mathbb{N}\}$; alors $v_n = 1$ si $n \in \mathcal{A}$ et $v_n = 0$ si $n \in \mathbb{N} \setminus \mathcal{A}$. En particulier, la sous-suite (v_{2^n}) de (v_n) converge vers 1. L'application $\varphi : n \mapsto \varphi(n) = 1 + 2^n$ de $\mathbb{N}$ dans $\mathbb{N}$ est strictement croissante. Pour tout entier $n \geq 1$, $2^n < \varphi(n) < 2^{n+1}$ donc $\varphi(n) \notin \mathcal{A}$, ce qui montre que la sous-suite $(v_{\varphi(n)})$ de (v_n) converge vers 0. La suite (nu_n) ne converge donc pas dans $\mathbb{R}$. $\square$

■ Règles de convergence

THÉORÈME 7.8. — Règle de d'Alembert⁴.

Si $u_n > 0$ pour tout n et si la suite (u_{n+1}/u_n) admet une limite ℓ dans la droite numérique achevée $\overline{\mathbb{R}}$, la série $\sum u_n$ converge si $\ell < 1$ et diverge si $\ell > 1$.

THÉORÈME 7.9. — Règle de Cauchy⁵.

Si $u_n \geq 0$ pour tout n et si la suite $(\sqrt[n]{u_n})$ admet une limite ℓ dans la droite numérique achevée $\overline{\mathbb{R}}$, la série $\sum u_n$ converge si $\ell < 1$ et diverge si $\ell > 1$.

Pour chacune de ces deux règles, on ne peut conclure dans le cas où $\ell = 1$, comme nous allons le voir sur des exemples. La règle de Cauchy, bien que moins pratique, est plus efficace ; en effet, la convergence de la suite quotient (u_{n+1}/u_n) entraîne celle de la suite $(\sqrt[n]{u_n})$. Cependant, la réciproque est fautive.

4. Ce critère est énoncé par d'Alembert vers 1750, mais il n'est correctement démontré qu'en 1821 par Cauchy ([CAU1], chapitre VI, §2^e, 2^e théorème, pages 134 et 135).

5. Ce critère est énoncé et prouvé par Cauchy en 1821 ([CAU1], chapitre VI, §2^e, 1^{er} théorème, pages 132 à 134) ; pour une démonstration « moderne », voir [ARN2], §IX.2, théorème IX.2.1.

7.14. Suite (u_n) à termes strictement positifs telle la série $\sum u_n$ converge et :

$$\lim_{n \rightarrow \infty} \sqrt[n]{u_n} = \lim_{n \rightarrow \infty} \frac{u_{n+1}}{u_n} = 1.$$

Nous considérons la suite $(u_n)_{n \geq 1}$ de terme général $u_n = \frac{1}{n^2}$.

Par le critère de Riemann, la série $\sum u_n$ converge. Pour tout entier $n \geq 1$:

$$v_n = \frac{u_{n+1}}{u_n} = \frac{n^2}{(n+1)^2} = \left(\frac{n}{n+1}\right)^2.$$

Il en résulte que la suite quotient (v_n) converge vers 1. Pour tout entier $n \geq 1$, $\sqrt[n]{u_n} = (1/\sqrt[n]{n})^2$; comme $\lim_{n \rightarrow \infty} \sqrt[n]{n} = 1$, la suite $(\sqrt[n]{u_n})$ tend également vers 1. $\square$

7.15. Suite (u_n) à termes strictement positifs telle la série $\sum u_n$ diverge et :

$$\lim_{n \rightarrow \infty} \sqrt[n]{u_n} = \lim_{n \rightarrow \infty} \frac{u_{n+1}}{u_n} = 1.$$

Nous utilisons la suite $(u_n)_{n \geq 1}$ de terme général $u_n = \frac{1}{n}$.

La série harmonique $\sum u_n$ diverge. On a, pour tout entier $n \geq 1$:

$$v_n = \frac{u_{n+1}}{u_n} = \frac{n}{n+1}.$$

Il en résulte que la suite quotient (v_n) converge vers 1. Pour tout entier $n \geq 1$, $\sqrt[n]{u_n} = 1/\sqrt[n]{n}$; comme $\lim_{n \rightarrow \infty} \sqrt[n]{n} = 1$, la suite $(\sqrt[n]{u_n})$ tend également vers 1. $\square$

7.16. Suite à termes strictement positifs pour laquelle la règle de Cauchy permet de conclure, alors que la règle de d'Alembert ne le permet pas.

Nous considérons la suite $(u_n)_{n \geq 1}$ de terme général $u_n = \begin{cases} \frac{1}{3^n} & \text{si } n \text{ est pair,} \\ \frac{4}{3^n} & \text{si } n \text{ est impair.} \end{cases}$

Pour tout entier $n \geq 1$, $\sqrt[n]{u_n} = \frac{1}{3}$ si n est pair et $\sqrt[n]{u_n} = \frac{\sqrt[4]{4}}{3}$ si n est impair.

Comme $\lim_{(n \rightarrow \infty)} \sqrt[n]{4} = 1$, la suite $(\sqrt[n]{u_n})$ converge vers $1/3 < 1$, ce qui permet d'affirmer par la règle de Cauchy que la série $\sum u_n$ converge. Pour tout $n \geq 1$:

$$v_n = \frac{u_{n+1}}{u_n} = \begin{cases} \frac{4}{3} & \text{si } n \text{ est pair,} \\ \frac{1}{12} & \text{si } n \text{ est impair.} \end{cases}$$

Par conséquent, la sous-suite (v_{2n}) de (v_n) admet pour limite $4/3 > 1$ dans $\mathbb{R}$ et la sous-suite (v_{2n+1}) de (v_n) admet pour limite $1/12 < 1$. En particulier la suite quotient (v_n) n'admet pas de limite dans $\mathbb{R}$: on ne peut pas conclure à l'aide de la règle de d'Alembert. $\square$

Nous remarquons que pour une série convergente $\sum u_n$, il peut y avoir une infinité d'entiers naturels n pour lesquels $u_{n+1}/u_n > 1$, comme le montre l'exemple que nous venons de voir. Par contre, si $\sqrt[n]{u_n} > 1$ pour une infinité d'entiers naturels n , la série $\sum u_n$ diverge ; en effet, son terme général ne tend pas vers 0.

Dans les cas où les suites (u_{n+1}/u_n) et $(\sqrt[n]{u_n})$ n'ont pas de limite dans $\overline{\mathbb{R}}$, on peut améliorer les règles de Cauchy et de d'Alembert grâce au théorème suivant, qui utilise les limites inférieures et supérieures des suites (u_{n+1}/u_n) et $(\sqrt[n]{u_n})$.

THÉORÈME 7.10. — Soit (u_n) une suite de réels strictement positifs.

a) Si $\limsup_{n \rightarrow \infty} \frac{u_{n+1}}{u_n} < 1$, la série $\sum u_n$ converge.

Si $\liminf_{n \rightarrow \infty} \frac{u_{n+1}}{u_n} > 1$, la série $\sum u_n$ diverge.

b) Si $\limsup_{n \rightarrow \infty} \sqrt[n]{u_n} < 1$, la série $\sum u_n$ converge.

Si $\liminf_{n \rightarrow \infty} \sqrt[n]{u_n} > 1$, la série $\sum u_n$ diverge.

Par les inégalités :

$$\liminf_{n \rightarrow \infty} \frac{u_{n+1}}{u_n} \leq \liminf_{n \rightarrow \infty} \sqrt[n]{u_n} \leq \limsup_{n \rightarrow \infty} \sqrt[n]{u_n} \leq \limsup_{n \rightarrow \infty} \frac{u_{n+1}}{u_n},$$

on voit que si l'on utilise le théorème 7.11, la règle avec la racine n -ième est plus efficace que celle utilisant la suite quotient. Dans l'exemple précédent 7.16, on a :

$$\limsup_{n \rightarrow \infty} \frac{u_{n+1}}{u_n} = \frac{4}{3} > 1 \text{ et } \liminf_{n \rightarrow \infty} \frac{u_{n+1}}{u_n} = \frac{1}{12} < 1,$$

donc le théorème 7.11 ne permet pas de déterminer la nature de la série $\sum u_n$.

On peut aussi améliorer la règle de d'Alembert en effectuant un développement limité à l'ordre 1 en $1/n$ du quotient u_{n+1}/u_n .

THÉORÈME 7.11. — Règle de Raabe-Duhamel⁶.

Si $u_n > 0$ pour tout n et s'il existe un nombre réel α tel que :

$$\frac{u_{n+1}}{u_n} = 1 - \frac{\alpha}{n} + o_{n \rightarrow \infty}\left(\frac{1}{n}\right),$$

la série $\sum u_n$ converge lorsque $\alpha > 1$ et diverge dans le cas où $\alpha < 1$.

7.17. Suite à termes strictement positifs pour laquelle la règle de Raabe-Duhamel permet de conclure, alors que la règle de d'Alembert ne le permet pas.

Soit α un réel, auquel nous associons la suite $(u_n)_{n \geq 1}$ de terme général $u_n = \frac{1}{n^\alpha}$.

On a, pour tout entier $n \geq 1$, $v_n = \frac{u_{n+1}}{u_n} = \frac{n^\alpha}{(n+1)^\alpha} = \left(\frac{n}{n+1}\right)^\alpha$.

La suite quotient (v_n) converge vers 1, donc la règle de d'Alembert ne permet pas de déterminer la nature de la série de Riemann $\sum u_n$. On a :

$$v_n = \frac{1}{\left(1 + \frac{1}{n}\right)^\alpha} = \frac{1}{1 + \frac{\alpha}{n} + o_{n \rightarrow \infty}\left(\frac{1}{n}\right)} = 1 - \frac{\alpha}{n} + o_{n \rightarrow \infty}\left(\frac{1}{n}\right)$$

donc on retrouve par la règle de Raabe-Duhamel que la série $\sum u_n$ converge pour $\alpha > 1$ et diverge pour $\alpha < 1$; cependant, on ne peut pas conclure pour $\alpha = 1$. $\square$

6. Du nom du mathématicien suisse Joseph Raabe (1801-1859) et du mathématicien français Jean-Marie Duhamel (1797-1872). Pour une démonstration, voir [ARN2], §IX.2.

THÉORÈME 7.12. — Critère de condensation de Cauchy⁷.

Si (u_n) est une suite décroissante de réels positifs, les séries $\sum u_n$ et $\sum (2^n u_{2^n})$ sont de même nature : elles convergent toutes deux ou divergent toutes deux.

L'hypothèse de décroissance de (u_n) est nécessaire. Dans les exemples qui suivent, nous imposerons que de plus cette suite converge vers 0.

7.18. Suite (u_n) de réels positifs telle que la série $\sum u_n$ converge, alors que la série $\sum (2^n u_{2^n})$ diverge.

Nous utilisons la suite $(u_n)_{n \geq 0}$, de terme général :

$$u_n = \begin{cases} \frac{1}{n} & \text{s'il existe } q \in \mathbb{N} \text{ tel que } n = 2^q, \\ 0 & \text{sinon.} \end{cases}$$

Comme $u_n \geq 0$ pour tout entier naturel n , la suite $(S_n)_{n \geq 0}$ des sommes partielles de la série $\sum u_n$ est croissante. Pour tout entier naturel n , on a $n < 2^n$ donc :

$$S_n \leq S_{2^n} = \sum_{p=0}^{2^n} u_p = \sum_{q=0}^n \frac{1}{2^q} = 2 - \frac{1}{2^n} < 2.$$

La suite (S_n) étant croissante et majorée dans $\mathbb{R}$, elle converge dans $\mathbb{R}$, donc la série $\sum u_n$ converge. Cependant, pour tout n :

$$T_n = \sum_{p=0}^n 2^p u_{2^p} = \sum_{p=0}^n 1 = n + 1$$

donc la suite (T_n) tend vers $+\infty$, ce qui montre que la série $\sum (2^n u_{2^n})$ diverge. $\square$

7.19. Suite (u_n) de réels positifs telle que la série $\sum u_n$ diverge, alors que la série $\sum (2^n u_{2^n})$ converge.

Nous considérons la suite $(u_n)_{n \geq 0}$ de terme général :

$$u_n = \begin{cases} 0 & \text{s'il existe } q \in \mathbb{N} \text{ tel que } n = 2^q, \\ \frac{1}{n} & \text{sinon} \end{cases}$$

et nous notons $(S_n)_{n \geq 0}$ la suite des sommes partielles de la série $\sum u_n$. On a, pour tout entier naturel n :

$$S_{2^n} = \sum_{p=0}^{2^n} \frac{1}{p} - \sum_{q=0}^n \frac{1}{2^q} = \sum_{p=0}^{2^n} \frac{1}{p} - 2 + \frac{1}{2^n}.$$

Comme la série harmonique diverge, la sous-suite (S_{2^n}) de (S_n) ne converge pas dans $\mathbb{R}$; il en est donc de même pour la suite (S_n) , ce qui montre que la série $\sum u_n$ diverge. Cependant, $2^n u_{2^n} = 0$ pour tout n , donc la série $\sum (2^n u_{2^n})$ converge. $\square$

7. Ce critère est énoncé et démontré par Cauchy en 1821 ([CAU1], chapitre VI, § 2^e, 3^e théorème, pages 135 et 136) ; voir aussi [WALT], § 5.10.

■ Comparaison série-intégrale

THÉORÈME 7.13. — Pour une fonction f de $\mathbb{R}$ dans $\mathbb{R}$ définie, positive et décroissante sur l'intervalle $[0, +\infty[$, il y a équivalence entre la convergence de l'intégrale⁸ :

$$\int_0^{+\infty} f(t) dt$$

— qui exprime l'intégrabilité de f sur $[0, +\infty[$ — et celle de la série $\sum f(n)$.

Le résultat devient faux si l'on ne suppose pas la fonction f décroissante.

7.20. Fonction f définie et positive sur $[0, +\infty[$, qui n'est pas intégrable sur $[0, +\infty[$ alors que $\sum f(n)$ converge.

Nous introduisons l'application, continue sur $\mathbb{R}$:

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R}_+ \\ x \longmapsto f(x) = |\sin(2\pi x)|. \end{array} \right.$$

Pour tout entier naturel n , $f(n) = |\sin(2\pi n)| = 0$, donc la série $\sum f(n)$ converge et sa somme est nulle. La fonction f étant de période 1, on a, pour tout $n \in \mathbb{N}$:

$$\begin{aligned} \int_n^{n+1} f(t) dt &= \int_0^1 f(t) dt = \int_0^{1/2} \underbrace{|\sin(2\pi t)|}_{=\sin(2\pi t)} dt + \int_{1/2}^1 \underbrace{|\sin(2\pi t)|}_{=-\sin(2\pi t)} dt \\ &= \frac{1}{2\pi} \left([-\cos(2\pi t)]_{t=0}^{t=1/2} + [\cos(2\pi t)]_{t=1/2}^{t=1} \right) = \frac{2}{\pi}. \end{aligned}$$

Pour tout nombre réel $x \geq 1$, on a, en notant p la partie entière de x :

$$\int_0^x f(t) dt \geq \int_0^p f(t) dt = \sum_{n=0}^{p-1} \int_n^{n+1} f(t) dt = \frac{2}{\pi}(p-1) \geq \frac{2}{\pi}(x-2)$$

donc l'intégrale $\int_0^{+\infty} f(t) dt$ diverge, ce qui, puisque $f \geq 0$ sur $[0, +\infty[$, montre que f n'est pas intégrable sur $[0, +\infty[$. $\square$

7.21. Fonction f définie, positive et intégrable sur l'intervalle $[0, +\infty[$ pour laquelle la série $\sum f(n)$ diverge.

Nous considérons la fonction f de l'exemple 11.12 (pages 216 et 217). Alors f est définie, positive et intégrable sur $[0, +\infty[$. Or $f(n) = 1$ pour tout entier $n \geq 2$, donc la série $\sum f(n)$ diverge. $\square$

■ Modification de l'ordre des termes

Modifier l'ordre des termes d'une suite (u_n) , c'est la remplacer par la suite (v_n) de terme général $v_n = u_{\varphi(n)}$ où φ est une bijection de $\mathbb{N}$ sur $\mathbb{N}$.

Pour une suite (u_n) à termes positifs, on ne change ni la nature (convergence ou divergence) ni la somme éventuelle de la série $\sum u_n$ en modifiant l'ordre des

8. Voir le chapitre 11, pages 213 à 219.

termes de cette suite. Pour une suite à termes quelconques, ce résultat reste vrai si la série est absolument convergente, mais il est faux dans le cas général. Donnons un exemple où la convergence est conservée alors que la somme de la série est modifiée, puis un autre où la nature de la série ne se conserve pas.

7.22. Modification de l'ordre des termes qui change la somme⁹.

Nous introduisons la suite $(u_n)_{n \geq 1}$ de terme général $u_n = \frac{(-1)^{n+1}}{n}$ et nous notons (S_n) la suite des sommes partielles de $\sum u_n$.

La série $\sum u_n$ converge par le critère des séries alternées. Nous notons $(H_n)_{n \geq 1}$ la suite des sommes partielles de la série harmonique. En désignant par γ la constante d'Euler, il existe une suite $(\varepsilon_n)_{n \geq 1}$ convergeant vers 0 telle que, pour tout entier $n \geq 1$, $H_n = \ln(n) + \gamma + \varepsilon_n$ — voir, dans ce qui suit la définition 7.3, la formule (2) de la page 111. On a, pour tout entier $n \geq 1$:

$$\begin{aligned} S_{2n} &= 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \cdots + \frac{1}{2n-1} - \frac{1}{2n} = H_{2n} - 2\left(\frac{1}{2} + \frac{1}{4} + \cdots + \frac{1}{2n}\right) \\ &= H_{2n} - H_n = \ln(2n) + \gamma + \varepsilon_{2n} - (\ln(n) + \gamma + \varepsilon_n) = \ln 2 + \varepsilon_{2n} - \varepsilon_n, \end{aligned}$$

donc la suite (S_{2n}) converge vers $\ln 2$. Comme la suite (S_n) converge dans $\mathbb{R}$, sa limite est $\ln 2$, donc la somme de la série $\sum u_n$ est égale à $\ln 2$.

Nous modifions l'ordre des termes de la suite $(u_n)_{n \geq 1}$ en construisant, à partir de (u_n) , la suite $(v_n)_{n \geq 1}$ ainsi : on prend alternativement et dans l'ordre, en partant de u_1 pour les termes positifs et de u_2 pour les termes négatifs, deux termes positifs consécutifs puis un terme négatif, et ainsi de suite, ce qui donne :

$$v_1 = 1, v_2 = \frac{1}{3}, v_3 = -\frac{1}{2}, v_4 = \frac{1}{5}, v_5 = \frac{1}{7}, v_6 = -\frac{1}{4}, v_7 = \frac{1}{9}, v_8 = \frac{1}{11}, v_9 = -\frac{1}{6}, \text{ etc.}$$

On a donc, pour tout $p \in \mathbb{N}$, $v_{3p+1} = \frac{1}{4p+1}$, $v_{3p+2} = \frac{1}{4p+3}$ et, si $p \geq 1$, $v_{3p} = -\frac{1}{2p}$.

Notons $(V_n)_{n \geq 1}$ la suite des sommes partielles de la série $\sum v_n$. En utilisant (1), on obtient, pour tout entier $n \geq 1$:

$$\begin{aligned} V_{3n} &= 1 + \frac{1}{3} + \frac{1}{5} + \cdots + \frac{1}{4n-1} - \left(\frac{1}{2} + \frac{1}{4} + \cdots + \frac{1}{2n}\right) \\ &= 1 + \frac{1}{2} + \frac{1}{3} + \cdots + \frac{1}{4n} - \left(\frac{1}{2} + \frac{1}{4} + \cdots + \frac{1}{4n}\right) - \left(\frac{1}{2} + \frac{1}{4} + \cdots + \frac{1}{2n}\right) \\ &= H_{4n} - \frac{1}{2}(H_{2n} + H_n) = \ln(4n) + \gamma + \varepsilon_{4n} - \frac{1}{2}(\ln(2n) + \gamma + \varepsilon_{2n} + \ln(n) + \gamma + \varepsilon_n) \\ &= \ln \frac{4n}{\sqrt{2n^2}} + \varepsilon_{4n} - \frac{1}{2}(\varepsilon_{2n} + \varepsilon_n) = \ln \frac{4}{\sqrt{2}} + \varepsilon_{4n} - \frac{1}{2}(\varepsilon_{2n} + \varepsilon_n). \end{aligned}$$

La suite (V_{3n}) converge donc dans $\mathbb{R}$ vers $\ell = \ln \frac{4}{\sqrt{2}} = \frac{3}{2} \ln 2$. Or, pour tout $n \geq 1$:

$$V_{3n+1} = V_{3n} + v_{3n+1} = V_{3n} + \frac{1}{4n+1} \quad \text{et} \quad V_{3n+2} = V_{3n+1} + v_{3n+2} = V_{3n+1} + \frac{1}{4n+3},$$

donc les trois sous-suites (V_{3n}) , (V_{3n+1}) et (V_{3n+2}) de (V_n) convergent dans $\mathbb{R}$ vers la même limite ℓ , ce qui montre qu'il en est de même de la suite (V_n) .

En conclusion, la série $\sum v_n$ converge et $\sum_{n=1}^{+\infty} v_n = \frac{3}{2} \ln 2 \neq \ln 2 = \sum_{n=1}^{+\infty} u_n$. $\square$

9. Emile Borel a démontré qu'en modifiant l'ordre des termes d'une série semi-convergente, on peut la faire converger vers toute valeur choisie à l'avance, ou bien la faire diverger vers $\pm\infty$.

7.23. Modification de l'ordre des termes qui transforme une série convergente en une série divergente.

Nous reprenons la suite $(u_n)_{n \geq 1}$ de l'exemple précédent 7.22 et nous modifions l'ordre des termes de (u_n) en construisant la suite $(r_n)_{n \geq 1}$ ainsi : on prend successivement, et dans l'ordre, le premier terme positif et le premier terme négatif de $(u_n)_{n \geq 1}$, les deux termes suivants positifs et le second terme négatif, les trois termes suivants positifs et le troisième terme négatif, et ainsi de suite, ce qui donne :

$$r_1 = 1, r_2 = -\frac{1}{2}, r_3 = \frac{1}{3}, r_4 = \frac{1}{5}, r_5 = -\frac{1}{4}, r_6 = \frac{1}{7}, r_7 = \frac{1}{9}, r_8 = \frac{1}{11}, r_9 = -\frac{1}{6}, \text{ etc.}$$

Nous posons, pour tout entier $n \geq 1$:

$$\varphi(n) = 1 + 2 + \dots + n = \frac{n(n+1)}{2} \quad \text{et} \quad \theta(n) = \varphi(n) + n = \varphi(n+1) - 1 = \frac{n(n+3)}{2}.$$

Pour tout entier $n \geq 1$, l'entier $\theta(n)$ est le rang du n -ième terme négatif de la suite $(r_k)_{k \geq 1}$ et $r_{\theta(n)} = u_{2n} = -1/(2n)$, et par une récurrence facile on voit que les $(n+1)$ termes positifs qui suivent $r_{\theta(n)}$ sont successivement :

$$r_{\varphi(n+1)} = u_{2\varphi(n)+1}, r_{\varphi(n+1)+1} = u_{2(\varphi(n)+1)+1}, \dots, r_{\varphi(n+1)+n} = u_{2(\varphi(n)+n)+1}.$$

Notons $(R_n)_{n \geq 1}$ la suite des sommes partielles de la série $\sum r_n$. Nous introduisons la sous-suite $(T_n)_{n \geq 1}$ de (R_n) de terme général $T_n = R_{\theta(n)}$. On a, pour tout $n \geq 1$:

$$T_{n+1} - T_n = \sum_{p=\theta(n)+1}^{\theta(n+1)} r_p = \sum_{p=\varphi(n+1)}^{\varphi(n+1)+n} r_p + r_{\theta(n+1)} = \sum_{k=0}^n u_{2(\varphi(n)+k)+1} - \frac{1}{2n+2},$$

et comme, pour tout $k \in \llbracket 0, n \rrbracket$:

$$u_{2(\varphi(n)+k)+1} \geq u_{2(\varphi(n)+n)+1} = u_{2\theta(n)+1} = \frac{1}{2\theta(n)+1} = \frac{1}{n(n+3)+1},$$

il vient :

$$\begin{aligned} T_{n+1} - T_n &\geq \frac{n+1}{n(n+3)+1} - \frac{1}{2(n+1)} = \frac{n^2+n+1}{2(n+1)(n(n+3)+1)} \\ &\geq \frac{n(n+1)}{2(n+1)(n(n+3)+n)} = \frac{1}{2(n+4)}. \end{aligned}$$

On a $T_1 = R_{\theta(1)} = R_2 = 1 - \frac{1}{2} = \frac{1}{2}$. Pour tout $n \geq 1$, $T_n - T_1 = \sum_{k=1}^{n-1} (T_{k+1} - T_k)$, donc :

$$T_n \geq T_1 + \sum_{k=1}^{n-1} \frac{1}{2(k+4)} > \sum_{k=1}^{n-1} \frac{1}{2(k+4)} = \frac{1}{2} \sum_{p=5}^{n+3} \frac{1}{p} = V_n.$$

La série harmonique diverge, donc (V_n) tend vers $+\infty$, ce qui montre que la sous-suite (T_n) de (R_n) diverge vers $+\infty$. Par conséquent, la suite (R_n) ne converge pas dans $\mathbb{R}$, donc la série $\sum r_n$ diverge, alors que la série $\sum u_n$ converge. $\square$

■ Séries doubles et produit de Cauchy

Etant donné une famille $(u_{n,p})_{n,p \in \mathbb{N}}$ de nombres réels indexée par $\mathbb{N}^2$, nous allons nous intéresser à l'existence éventuelle des sommes de séries :

$$\sum_{n=0}^{+\infty} \sum_{p=0}^{+\infty} u_{n,p} \quad \text{et} \quad \sum_{p=0}^{+\infty} \sum_{n=0}^{+\infty} u_{n,p}$$

et à la possibilité d'égalité de ces deux sommes.

Nous rappelons tout d'abord le résultat fondamental suivant ¹⁰.

THÉORÈME 7.14. — Soit $(u_{n,p})_{n,p \in \mathbb{N}}$ une famille de nombres réels indexée par $\mathbb{N}^2$.

On suppose que, pour tout entier naturel p , la série $\sum_{n \geq 0} u_{n,p}$ est absolument convergente et on pose :

$$S_p = \sum_{n=0}^{+\infty} u_{n,p}.$$

Si la série $\sum_{p \geq 0} S_p$ converge, alors, pour tout entier naturel n , la série $\sum_{p \geq 0} u_{n,p}$ est absolument convergente et en posant, pour tout $n \in \mathbb{N}$:

$$T_n = \sum_{p=0}^{+\infty} u_{n,p},$$

la série $\sum_{n \geq 0} T_n$ converge et :

$$\sum_{n=0}^{+\infty} \sum_{p=0}^{+\infty} u_{n,p} = \sum_{p=0}^{+\infty} \sum_{n=0}^{+\infty} u_{n,p}.$$

Cependant, si les séries sont seulement convergentes, les sommes peuvent différer.

7.24. Famille $(u_{n,p})_{n,p \in \mathbb{N}}$ de réels telle que la série $\sum_{n \geq 0} (\sum_{p=0}^{+\infty} u_{n,p})$ diverge alors que la série $\sum_{p \geq 0} (\sum_{n=0}^{+\infty} u_{n,p})$ converge.

Nous utilisons la famille $(u_{n,p})_{n,p \in \mathbb{N}}$ de réels définie par $u_{n,p} = \begin{cases} -\frac{1}{2^n} & \text{si } n \geq 1, \\ 1 & \text{si } n = 0. \end{cases}$

Soit p un entier naturel. La série $\sum_{n \geq 0} u_{n,p}$ converge et $S_p = \sum_{n=0}^{+\infty} u_{n,p} = 1 - \sum_{n=1}^{+\infty} \frac{1}{2^n} = 0$.

Par conséquent, la série $\sum_{p \geq 0} (\sum_{n=0}^{+\infty} u_{n,p})$ converge et sa somme est nulle.

Par contre, pour tout entier naturel n , $u_{n,p}$ est, pour tout $p \in \mathbb{N}$, la constante non nulle 1 ou $-1/2^n$, donc la série $\sum_{p \geq 0} u_{n,p}$ diverge. $\square$

7.25. Famille $(u_{n,p})_{n,p \in \mathbb{N}}$ de nombres réels telle que les séries :

$$\sum_{n \geq 0} \left(\sum_{p=0}^{+\infty} u_{n,p} \right) \text{ et } \sum_{p \geq 0} \left(\sum_{n=0}^{+\infty} u_{n,p} \right)$$

convergent mais n'ont pas la même somme.

Nous considérons la famille $(u_{n,p})_{n,p \in \mathbb{N}}$ définie par $u_{n,p} = \begin{cases} 1 & \text{si } n = p, \\ -\frac{1}{2^{p-n}} & \text{si } n < p, \\ 0 & \text{si } n > p. \end{cases}$

Soit $p \in \mathbb{N}$. Comme $u_{n,p} = 0$ pour tout entier $n > p$, la série $\sum_{n \geq 0} u_{n,p}$ converge et :

$$S_p = \sum_{n=0}^{+\infty} u_{n,p} = \sum_{n=0}^p u_{n,p} = -\sum_{n=0}^{p-1} \frac{1}{2^{p-n}} + 1 = 1 - \sum_{k=1}^p \frac{1}{2^k} = \frac{1}{2^p}.$$

10. Le théorème 7.14 qui suit découle de l'étude des familles sommables (voir [CHOQ], chapitre VII-III, §8, ou [ARN2], §IX.7) : si la famille $(u_{n,p})_{(n,p) \in \mathbb{N}^2}$ est sommable, les deux sommes considérées existent et sont égales.

Il en résulte que la série $\sum_{p \geq 0} S_p$ converge et que $\sum_{p=0}^{+\infty} \sum_{n=0}^{+\infty} u_{n,p} = 2$.

Soit $n \in \mathbb{N}$. Pour tout entier $p > n$, $u_{n,p} = -\frac{1}{2^{p-n}} = -2^n \frac{1}{2^p}$, donc la série $\sum_{p \geq 0} u_{n,p}$ converge et :

$$T_n = \sum_{p=0}^{+\infty} u_{n,p} = \sum_{p=n}^{+\infty} u_{n,p} = 1 - \sum_{p=n+1}^{+\infty} \frac{1}{2^{p-n}} = 1 - \sum_{k=1}^{+\infty} \frac{1}{2^k} = 0.$$

Par conséquent, la série $\sum_{n \geq 0} T_n$ converge et $\sum_{n=0}^{+\infty} \sum_{p=0}^{+\infty} u_{n,p} = 0 \neq 2 = \sum_{p=0}^{+\infty} \sum_{n=0}^{+\infty} u_{n,p}$. $\square$

DÉFINITION 7.5. — Le produit de Cauchy des suites (u_n) et (v_n) de nombres réels est la suite (w_n) , de terme général :

$$w_n = \sum_{k=0}^n u_k v_{n-k} = u_0 v_n + u_1 v_{n-1} + \cdots + u_k v_{n-k} + \cdots + u_{n-1} v_1 + u_n v_0,$$

et la série $\sum w_n$ est appelée la série produit des séries $\sum u_n$ et $\sum v_n$.

THÉORÈME 7.15. — Si l'une des deux séries $\sum u_n$ ou $\sum v_n$ converge absolument, leur série produit $\sum w_n$ converge et ¹¹ :

$$\sum_{n=0}^{+\infty} w_n = \left(\sum_{n=0}^{+\infty} u_n \right) \left(\sum_{n=0}^{+\infty} v_n \right).$$

7.26. Série produit divergente de deux séries convergentes¹².

Nous utilisons la suite $(u_n)_{n \geq 0}$ de terme général $u_n = \frac{(-1)^n}{\sqrt{n+1}}$.

Nous savons par le critère de Leibniz que la série $\sum u_n$ converge. Nous notons $(w_n)_{n \geq 0}$ le produit de Cauchy de la suite (u_n) par elle-même.

Soit n un entier naturel. On a :

$$w_n = \sum_{k=0}^n u_k u_{n-k} = (-1)^n \sum_{k=0}^n \frac{1}{\sqrt{k+1} \sqrt{n-k+1}} = (-1)^n \sum_{p=1}^{n+1} \frac{1}{\sqrt{p} \sqrt{n+2-p}}.$$

Or, pour tout $p \in [1, n+1]$:

$$p + (n+2-p) - 2\sqrt{p} \sqrt{n+2-p} = (\sqrt{p} - \sqrt{n+2-p})^2 \geq 0,$$

donc $0 < \sqrt{p} \sqrt{n+2-p} \leq \frac{n+2}{2}$, d'où l'on déduit que :

$$|w_n| \geq \sum_{p=1}^{n+1} \frac{2}{n+2} = 2 \frac{n+1}{n+2} = 2 \left(1 - \frac{1}{n+2} \right) \geq 2 \left(1 - \frac{1}{2} \right) = 1.$$

11. Ce résultat est démontré en 1821 par Cauchy, d'abord pour les séries à termes positifs ([CAU1], chapitre VI, §2^e, 6^e théorème, pages 141 et 142) puis pour les séries absolument convergentes ([CAU1], chapitre VI, §3^e, 6^e théorème, pages 147 à 149) et affiné en 1875 par le mathématicien autrichien Franz Mertens (1840-1927). Le résultat est faux dans le cas général mais cependant, si les trois séries convergent, on a l'égalité $\sum_{n=0}^{+\infty} w_n = \left(\sum_{n=0}^{+\infty} u_n \right) \left(\sum_{n=0}^{+\infty} v_n \right)$.

12. Ce contre-exemple est dû à Cauchy en 1821 ([CAU1], chapitre VI, §3^e, pages 149 et 150).

Il en résulte que la suite (w_n) ne converge pas vers 0, ce qui montre que la série produit $\sum w_n$ diverge. $\square$

7.27. Deux séries divergentes dont le produit de Cauchy converge.

Nous considérons les suites $(u_n)_{n \geq 0}$ et $(v_n)_{n \geq 0}$, de termes généraux :

$$u_n = \begin{cases} 2 & \text{si } n = 0, \\ 2^n & \text{si } n \geq 1 \end{cases} \quad \text{et} \quad v_n = \begin{cases} -1 & \text{si } n = 0, \\ 1 & \text{si } n \geq 1. \end{cases}$$

Les suites (u_n) et (v_n) ne convergent pas vers 0, donc les séries $\sum u_n$ et $\sum v_n$ divergent. Notons $(w_n)_{n \geq 0}$ le produit de Cauchy des suites (u_n) et (v_n) . Pour tout entier naturel n , $w_n = \sum_{k=0}^n u_k v_{n-k}$, donc $w_0 = -2$ et, pour $n \geq 1$:

$$w_n = u_0 v_n + \sum_{k=1}^{n-1} u_k v_{n-k} + u_n v_0 = 2 + \sum_{k=1}^{n-1} 2^k - 2^n = 0.$$

Ceci montre que la série produit $\sum w_n$ converge — et que sa somme vaut -2 . $\square$

■ Equivalence des suites des sommes partielles et équivalence des suites des restes

THÉORÈME 7.16. — Si $u_n \geq 0$ à partir d'un certain rang, si la série $\sum u_n$ diverge et si (v_n) est une suite équivalente à (u_n) , alors les suites (U_n) et (V_n) des sommes partielles respectives des séries $\sum u_n$ et $\sum v_n$ sont équivalentes.

Ceci devient faux si u_n n'est pas positif à partir d'un certain rang.

7.28. Suites (u_n) telle que la série $\sum u_n$ diverge et (v_n) équivalente à (u_n) , pour lesquelles les suites des sommes partielles des séries $\sum u_n$ et $\sum v_n$ ne sont pas équivalentes.

Nous considérons les suites $(u_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$, de termes généraux :

$$u_n = (-1)^n \text{ et } v_n = (-1)^n + \frac{1}{n}.$$

On a, pour tout n , $v_n = u_n \left(1 + \frac{(-1)^n}{n}\right)$. Or $\lim_{n \rightarrow \infty} \frac{(-1)^n}{n} = 0$, donc $v_n \underset{n \rightarrow \infty}{\sim} u_n$.

Nous introduisons les suites $(U_n)_{n \geq 1}$ et $(V_n)_{n \geq 1}$ des sommes partielles respectives des séries $\sum u_n$ et $\sum v_n$.

Pour tout entier $n \geq 1$, $U_n = 1$ si n est pair et $U_n = 0$ si n est impair, donc la suite (U_n) ne converge pas dans $\mathbb{R}$, ce qui prouve que la série $\sum u_n$ diverge.

La suite $(S_n)_{n \geq 1}$ des sommes partielles de la série harmonique tend vers $+\infty$ et, pour tout entier $n \geq 1$, $V_n = U_n + S_n \geq S_n$, donc la suite (V_n) tend vers $+\infty$; comme (U_n) est bornée, les suites $(U_n)_{n \geq 1}$ et $(V_n)_{n \geq 1}$ ne sont pas équivalentes. $\square$

DÉFINITION 7.6. — Si la série $\sum u_n$ converge, la suite des restes de la série $\sum u_n$ est la suite (R_n) de terme général :

$$R_n = \sum_{k=n+1}^{+\infty} u_k.$$

Nous complétons le théorème 7.13 (page 122) « comparaison série-intégrale ».

THÉORÈME 7.17. — Pour une fonction f de $\mathbb{R}$ dans $\mathbb{R}$ définie, positive et décroissante sur $[0, +\infty[$, il y a équivalence entre la convergence de l'intégrale :

$$\int_0^{+\infty} f(t) dt$$

et celle de la série $\sum f(n)$ et, si cette série converge, alors, en notant $(R_n)_{n \geq 0}$ la suite des restes de la série $\sum f(n)$, on a, pour tout n :

$$\int_{n+1}^{+\infty} f(t) dt \leq R_n \leq \int_n^{+\infty} f(t) dt.$$

Dans le texte de ce théorème, on peut bien sûr remplacer $[0, +\infty[$ par $[1, +\infty[$, la convergence éventuelle de l'intégrale $\int_0^{+\infty} f(t) dt$ par celle de $\int_1^{+\infty} f(t) dt$, les suites étant alors définies à partir du rang 1.

THÉORÈME 7.18. — Si $u_n \geq 0$ à partir d'un certain rang, si la série $\sum u_n$ converge et si la suite (v_n) est équivalente à la suite (u_n) , la série $\sum v_n$ converge et les suites des restes des séries $\sum u_n$ et $\sum v_n$ sont équivalentes.

Ceci devient faux si les termes de la suite (u_n) ne sont pas tous positifs à partir d'un certain rang. Dans l'exemple suivant nous imposons, pour que le problème ait un sens, à la série $\sum v_n$ de converger — voir l'exemple 7.6 (pages 114 et 115).

7.29. Suites (u_n) et (v_n) équivalentes telles que les séries $\sum u_n$ et $\sum v_n$ convergent, alors que les suites des restes de ces séries ne sont pas équivalentes.

Nous introduisons les suites $(u_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$, de termes généraux :

$$u_n = \frac{(-1)^n}{n} \text{ et } v_n = \frac{(-1)^n}{n} + \frac{1}{n\sqrt{n}} = u_n + \frac{1}{n\sqrt{n}}.$$

On a, pour tout n , $v_n = u_n \left(1 + \frac{(-1)^n}{\sqrt{n}}\right)$. Or $\lim_{n \rightarrow \infty} \frac{(-1)^n}{\sqrt{n}} = 0$, donc $v_n \underset{n \rightarrow \infty}{\sim} u_n$.

La série $\sum u_n$ converge par le critère des séries alternées et le critère de Riemann montre que la série :

$$\sum_{n \geq 1} \frac{1}{n\sqrt{n}}$$

converge, donc la série $\sum v_n$ converge. Ces convergences justifient l'introduction des suites respectives $(R_n)_{n \geq 1}$ et $(S_n)_{n \geq 1}$ des restes des séries $\sum u_n$ et $\sum v_n$.

Nous considérons alors les suites $(w_n)_{n \geq 1}$, $(z_n)_{n \geq 1}$ et $(t_n)_{n \geq 1}$, de termes généraux :

$$w_n = u_{2n} + u_{2n+1} = \frac{1}{2n} - \frac{1}{2n+1} = \frac{1}{2n(2n+1)}, \quad z_n = \frac{1}{4n^2} \text{ et } t_n = \frac{1}{n\sqrt{n}}.$$

Les séries $\sum z_n$ et $\sum t_n$ convergent par le critère de Riemann et $0 < w_n < z_n$ pour tout n , donc la série $\sum w_n$ converge ; de plus la suite (w_n) est équivalente à la suite (z_n) . Nous notons $(W_n)_{n \geq 1}$, $(Z_n)_{n \geq 1}$ et $(T_n)_{n \geq 1}$ les suites respectives des

restes des séries $\sum w_n$, $\sum z_n$ et $\sum t_n$. Les fonctions :

$$f : x \mapsto f(x) = \frac{1}{4x^2} \text{ et } g : x \mapsto g(x) = \frac{1}{x\sqrt{x}}$$

sont définies, continues, positives et décroissantes sur l'intervalle $[1, +\infty[$, donc on déduit du théorème 7.17 (page 128) que, pour tout entier $n \geq 1$:

$$\frac{1}{4(n+1)} = \int_{n+1}^{+\infty} \frac{1}{4t^2} dt \leq Z_n \leq \int_n^{+\infty} \frac{1}{4t^2} dt = \frac{1}{4n}$$

et :

$$\frac{2}{\sqrt{n+1}} = \int_{n+1}^{+\infty} \frac{1}{t\sqrt{t}} dt \leq T_n \leq \int_n^{+\infty} \frac{1}{t\sqrt{t}} dt = \frac{2}{\sqrt{n}}$$

ce qui montre que $Z_n \underset{n \rightarrow \infty}{\sim} \frac{1}{4n}$ et $T_n \underset{n \rightarrow \infty}{\sim} \frac{2}{\sqrt{n}}$.

Par le théorème 7.18, $W_n \underset{n \rightarrow \infty}{\sim} Z_n$ donc $W_n \underset{n \rightarrow \infty}{\sim} \frac{1}{4n}$. Pour tout entier $n \geq 1$:

$$R_{2n+1} = \sum_{p=2n+2}^{+\infty} u_p = \sum_{k=n+1}^{+\infty} w_k = W_n > 0 \text{ et } R_{2n} = u_{2n+1} + R_{2n+1} = -\frac{1}{2n+1} + R_{2n+1}$$

donc :

$$\frac{R_{2n}}{R_{2n+1}} = -\frac{1}{(2n+1)R_{2n+1}} + 1 = 1 - \frac{1}{(2n+1)W_n}.$$

Comme $(2n+1)W_n \underset{n \rightarrow \infty}{\sim} \frac{2n}{4n} = \frac{1}{2}$, $\lim_{n \rightarrow \infty} (2n+1)W_n = \frac{1}{2}$ donc $\lim_{n \rightarrow \infty} \frac{R_{2n}}{R_{2n+1}} = -1$.

Par suite $|R_{2n}| \underset{n \rightarrow \infty}{\sim} |R_{2n+1}| = R_{2n+1} \underset{n \rightarrow \infty}{\sim} \frac{1}{4n}$; on en déduit que $|R_n| \underset{n \rightarrow \infty}{\sim} \frac{1}{2n}$.

Pour tout $k \in \mathbb{N}^*$, $v_k = u_k + t_k$. Il en résulte que $S_n = R_n + T_n$ pour tout $n \geq 1$.

On a $|R_n| \underset{n \rightarrow \infty}{\sim} \frac{1}{2n}$, $\frac{1}{2n} = \underset{n \rightarrow \infty}{o}\left(\frac{2}{\sqrt{n}}\right)$ et $\frac{2}{\sqrt{n}} \underset{n \rightarrow \infty}{\sim} T_n$, donc $R_n = \underset{n \rightarrow \infty}{o}(T_n)$.

Par conséquent, $S_n \underset{n \rightarrow \infty}{\sim} T_n$ donc $R_n = \underset{n \rightarrow \infty}{o}(S_n)$.

En conclusion, les suites (R_n) et (S_n) ne sont pas équivalentes. $\square$

■ Autres types de convergence

Euler considère que la série de terme général $(-1)^n$ converge et que sa somme est égale à $1/2$. Plus près de nous, Emile Borel développe une théorie des séries divergentes¹³. Les mathématiciens ont introduit des notions de convergence des séries plus générales que la définition traditionnelle, permettant ainsi de donner une somme à des séries divergentes au sens habituel.

Nous avons vu, dans le chapitre 6, à la page 104, la notion de convergence d'une suite au sens de Cesàro.

13. *Leçons sur les séries divergentes*, Gauthier-Villard 1928 (pour la 2^e édition), réédition Jacques Gabay 1988.

DÉFINITION 7.7. — Une série $\sum u_n$ converge au sens de Cesàro si la suite (S_n) des sommes partielles de $\sum u_n$ converge au sens de Cesàro, ce qui signifie que la suite (T_n) , de terme général :

$$T_n = \frac{S_0 + S_1 + \cdots + S_n}{n},$$

converge dans $\mathbb{R}$ au sens habituel et, dans ce cas, la limite « ordinaire » de la suite (T_n) s'appelle la somme au sens de Cesàro de la série $\sum u_n$.

Si une suite converge vers un réel ℓ au sens habituel, elle converge vers ℓ au sens de Cesàro donc, si la série $\sum u_n$ converge au sens traditionnel, elle converge au sens de Cesàro et les sommes au sens habituel et au sens de Cesàro sont les mêmes.

7.30. Suite (u_n) telle que la série $\sum u_n$ converge au sens de Cesàro mais ne converge pas au sens habituel.

L'exemple traditionnel consiste à utiliser la suite (u_n) de terme général $u_n = (-1)^n$.

Nous proposons un cas plus général. Nous considérons une suite $(u_n)_{n \geq 0}$ non identiquement nulle admettant une période $p \in \mathbb{N}^*$; alors $u_{n+p} = u_n$ pour tout $n \in \mathbb{N}$ et il existe au moins un entier q appartenant à $[0, p-1]$ tel que $u_q \neq 0$.

Pour tout entier naturel n , $u_{q+np} = u_q \neq 0$, donc la sous-suite $(u_{q+np})_{n \geq 0}$ de (u_n) ne converge pas vers 0. Il en résulte que la suite (u_n) ne tend pas vers 0, ce qui montre que la série $\sum u_n$ ne converge pas au sens habituel.

On suppose de plus que $S = \sum_{k=0}^{p-1} u_k = 0$.

Alors la suite $(S_n)_{n \geq 0}$ des sommes partielles de la série $\sum u_n$ admet p pour période ; en effet, pour tout $n \in \mathbb{N}$:

$$S_{n+p} = \sum_{k=0}^{n+p} u_k = \sum_{k=0}^n u_k + \sum_{k=n+1}^{n+p} u_k = S_n + S = S_n.$$

La suite (S_n) est donc bornée, ce qui montre que (S_n/n) converge vers 0 : la série $\sum u_n$ converge au sens de Cesàro et sa somme au sens de Cesàro est égale à 0. $\square$

DÉFINITION 7.8. — On dit d'une série $\sum u_n$ qu'elle converge au sens d'Abel si le rayon de convergence de la série entière $\sum u_n x^n$ est au moins égal à 1 et si la fonction :

$$f : x \mapsto f(x) = \sum_{n=0}^{+\infty} u_n x^n,$$

définie sur l'intervalle $] -1, 1[$, admet une limite finie ℓ quand x tend vers 1 par valeurs inférieures et, si la série $\sum u_n$ converge au sens d'Abel, sa somme au sens d'Abel est le nombre réel $\ell = \lim_{x \rightarrow 1^-} f(x)$.

On démontre que si la série $\sum u_n$ converge au sens de Cesàro, elle converge au sens d'Abel et qu'alors les sommes sont les mêmes pour les deux types de convergence. Cependant la réciproque est fausse.

7.31. Suite (u_n) telle que la série $\sum u_n$ converge au sens d'Abel mais ne converge pas au sens de Cesàro.

Nous utilisons la suite $(u_n)_{n \geq 0}$ de terme général $u_n = (-1)^{n+1} n(n+1)$.

Alors $u_0 = 0$. Le rayon de convergence de la série entière $\sum u_n x^n$ est égal à 1, ce qui justifie l'existence de la fonction f définie sur l'intervalle $] -1, 1[$ par :

$$f(x) = \sum_{n=0}^{+\infty} u_n x^n = x \sum_{n=1}^{+\infty} (-1)^{n+1} n(n+1) x^{n-1}.$$

Comme on peut dériver terme à terme, on a, pour tout point x de $] -1, 1[$:

$$\begin{aligned} f(x) &= x \sum_{n=1}^{+\infty} (-1)^{n+1} \frac{d^2}{dx^2} (x^{n+1}) = x \times \frac{d^2}{dx^2} \left(\sum_{n=1}^{+\infty} (-1)^{n+1} x^{n+1} \right) \\ &= x \times \frac{d^2}{dx^2} \left(\sum_{n=2}^{+\infty} (-1)^n x^n \right) = x \times \frac{d^2}{dx^2} \left(1 - x + \sum_{n=2}^{+\infty} (-1)^n x^n \right) \quad \left[\frac{d^2}{dx^2} (1-x) = 0 \right] \\ &= x \times \frac{d^2}{dx^2} \left(\sum_{n=0}^{+\infty} (-1)^n x^n \right) = x \times \frac{d^2}{dx^2} \left(\frac{1}{1+x} \right) = \frac{2x}{(1+x)^3}, \end{aligned}$$

donc $f(x)$ admet pour limite $1/4$ quand x tend vers 1 par valeurs inférieures. Par suite, la série $\sum u_n$ converge au sens d'Abel et sa somme au sens d'Abel vaut $1/4$.

Nous notons $(S_n)_{n \geq 0}$ la suite des sommes partielles de la série $\sum u_n$ et nous introduisons la suite (T_n) de terme général :

$$T_n = \frac{S_0 + S_1 + \dots + S_n}{n}.$$

On a, pour tout $n \in \mathbb{N}$:

$$\begin{aligned} S_{2n+1} &= \sum_{p=0}^{2n+1} u_p = \sum_{k=0}^n (u_{2k} + u_{2k+1}) = \sum_{k=0}^n (-(2k)(2k+1) + (2k+1)(2k+2)) \\ &= \sum_{k=0}^n (4k+2) = 4 \sum_{k=0}^n k + 2 \sum_{k=0}^n 1 = 2n(n+1) + 2(n+1) = 2(n+1)^2. \end{aligned}$$

Par conséquent, pour tout entier naturel n :

$$\begin{cases} S_{2n} = S_{2n-1} + u_{2n} = 2n^2 - (2n)(2n+1) = 2n(n - (2n+1)) = -2n(n+1) \\ \text{et } S_{2n} + S_{2n+1} = -2n(n+1) + 2(n+1)^2 = 2(n+1)(-n + (n+1)) = 2(n+1). \end{cases}$$

Il en résulte que, pour tout $n \in \mathbb{N}$:

$$\sum_{p=0}^{2n+1} S_p = \sum_{k=0}^n (S_{2k} + S_{2k+1}) = \sum_{k=0}^n 2(k+1) = 2 \sum_{q=1}^{n+1} q = (n+1)(n+2).$$

Ainsi $T_{2n+1} = \frac{(n+1)(n+2)}{2n+1} \underset{n \rightarrow \infty}{\sim} \frac{n}{2}$, donc la suite (T_n) ne converge pas dans $\mathbb{R}$.

En conclusion, la série $\sum u_n$ ne converge pas au sens de Cesàro. $\square$

On doit aussi à Emile Borel un autre type de convergence, qui découle de la remarque suivante :

$$\int_0^{+\infty} \frac{1}{n!} x^n e^{-x} dx = 1.$$

DÉFINITION 7.9. — a) La série $\sum u_n$ converge au sens de Borel si, pour tout nombre réel $x \geq 0$, la série :

$$\sum_{n \geq 0} \frac{u_n}{n!} x^n e^{-x}$$

converge et si, pour la fonction S définie sur l'intervalle $[0, +\infty[$ par :

$$S(x) = \sum_{n=0}^{+\infty} \frac{u_n}{n!} x^n e^{-x},$$

l'intégrale $\int_0^{+\infty} S(x) dx$ est convergente.

b) Si la série $\sum u_n$ converge au sens de Borel, sa somme au sens de Borel est le nombre réel :

$$B = \int_0^{+\infty} \left(\sum_{n=0}^{+\infty} \frac{u_n}{n!} x^n e^{-x} \right) dx.$$

On démontre que si la série $\sum u_n$ converge au sens d'Abel, elle converge au sens de Borel et qu'alors les sommes sont les mêmes pour les deux types de convergence. Cependant la réciproque est fausse.

7.32. Suite (u_n) telle que la série $\sum u_n$ converge au sens de Borel mais ne converge pas au sens d'Abel.

Nous choisissons un nombre réel $a > 1$ et nous lui associons la suite $(u_n)_{n \geq 0}$ de terme général $u_n = (-1)^n a^n$. Le rayon de convergence de la série entière $\sum u_n x^n$ est $1/a$ et $1/a < 1$, donc la fonction :

$$f : x \mapsto f(x) = \sum_{n=0}^{+\infty} u_n x^n$$

n'est pas définie au voisinage de 1 à gauche. Il en résulte que la série $\sum (-1)^n a^n$ ne converge pas au sens d'Abel. Par ailleurs, pour tout nombre réel $x \geq 0$:

$$S(x) = \sum_{n=0}^{+\infty} \frac{u_n}{n!} x^n e^{-x} = \sum_{n=0}^{+\infty} \frac{(-ax)^n}{n!} e^{-x} = e^{-(1+a)x}$$

donc l'intégrale $\int_0^{+\infty} S(x) dx$ est convergente et :

$$B = \int_0^{+\infty} S(x) dx = \int_0^{+\infty} e^{-(1+a)x} dx = \left[\frac{-e^{-(1+a)x}}{1+a} \right]_{x=0}^{x=+\infty} = \frac{1}{1+a}.$$

La série $\sum (-1)^n a^n$ converge donc au sens de Borel. $\square$

La somme au sens de Borel de la série $\sum (-1)^n a^n$ de l'exemple précédent 7.32 est égale à $1/(1+a)$. Il est amusant de voir que cette somme au sens de Borel vérifie la même formule que la somme habituelle d'une série géométrique dont la raison appartient à $] -1, 1[$.

Signalons enfin que pour les séries à termes réels positifs, les quatre types de convergence coïncident.

Chapitre 8

Fonctions d'une variable réelle Continuité et limites

La notion de continuité a été évoquée dès le XVII^e siècle. Cependant il faut attendre Bolzano en 1817 et Cauchy en 1821 pour en obtenir une définition satisfaisante. On doit à Karl Weierstrass, vers 1860, la définition que nous utilisons de nos jours. Les structures multiples que l'on peut mettre sur $\mathbb{R}$ (ordre, topologie, structure de corps) interfèrent les unes avec les autres et donnent aux fonctions d'une variable réelle, à savoir les applications d'une partie de $\mathbb{R}$ dans $\mathbb{R}$, un certain nombre de propriétés qui ne peuvent en général que partiellement se généraliser aux fonctions définies sur une partie de $\mathbb{R}^n$ où $n \geq 2$, comme par exemple le théorème des valeurs intermédiaires.

Dans ce chapitre, les fonctions sont des fonctions de $\mathbb{R}$ dans $\mathbb{R}$, que l'on appelle aussi des fonctions réelles d'une variable réelle. Une telle fonction est une application d'une partie de $\mathbb{R}$, appelée son domaine, dans $\mathbb{R}$. Si f est une fonction et $\mathcal{A}$ une partie de $\mathbb{R}$, on dit que f est définie sur $\mathcal{A}$ si $\mathcal{A}$ est incluse dans le domaine de f .

■ Continuité

Dans les définitions et le théorème qui suivent, $\mathcal{A}$ est une partie de $\mathbb{R}$.

La continuité est une propriété des fonctions liée à la topologie ; elle se définit de la manière suivante.

DÉFINITION 8.1. — Une application f de $\mathcal{A}$ dans $\mathbb{R}$ est continue en un point a de $\mathcal{A}$ si, pour tout réel $\varepsilon > 0$, il existe un réel $\eta > 0$ tel que, pour tout point x de $\mathcal{A}$:

$$|x - a| < \eta \text{ entraîne } |f(x) - f(a)| < \varepsilon.$$

Les fanatiques du symbolisme de la logique des prédicats du premier ordre — pas de la logique, seulement des symboles ! — pourront écrire cette définition ainsi :

$$(\forall \varepsilon > 0)(\exists \eta > 0)(\forall x \in \mathcal{A})(|x - a| < \eta \implies |f(x) - f(a)| < \varepsilon).$$

On comprend facilement que les applications purement ensemblistes de $\mathbb{R}$ dans $\mathbb{R}$ n'ont que peu d'intérêt au regard de la continuité, car elles n'utilisent pas les propriétés qui font de $\mathbb{R}$ un espace topologique ordonné.

La topologie de $\mathbb{R}$ provient d'une métrique obtenue par la distance canonique $d : (x, y) \mapsto d(x, y) = |x - y|$ de $\mathbb{R}$, ce qui se traduit par le théorème qui suit¹.

THÉORÈME 8.1. — Une application f de $\mathcal{A}$ dans $\mathbb{R}$ est continue en un point a de $\mathcal{A}$ si, et seulement si, pour toute suite (x_n) de points de $\mathcal{A}$ qui converge vers a , la suite image $(f(x_n))$ converge vers $f(a)$.

DÉFINITION 8.2. — a) Une application de $\mathcal{A}$ dans $\mathbb{R}$ est continue sur $\mathcal{A}$ si elle est continue en tout point de $\mathcal{A}$.

b) Une application de $\mathcal{A}$ dans $\mathbb{R}$ est discontinue en un point a de $\mathcal{A}$ si elle n'est pas continue en a .

Une application de $\mathcal{A}$ dans $\mathbb{R}$ est discontinue sur $\mathcal{A}$ si elle n'est pas continue sur $\mathcal{A}$, ce qui équivaut à l'existence d'au moins un point de $\mathcal{A}$ en lequel elle est discontinue.

DÉFINITION 8.3. — Si f est une fonction définie sur $\mathcal{A}$, f est continue sur $\mathcal{A}$ si la restriction de f à $\mathcal{A}$ est une application de $\mathcal{A}$ dans $\mathbb{R}$ continue sur $\mathcal{A}$.

8.1. Fonction qui n'est continue en aucun point.

Nous considérons la fonction de Dirichlet², c'est-à-dire l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 1 & \text{si } x \text{ est rationnel,} \\ 0 & \text{si } x \text{ est irrationnel.} \end{cases} \end{array} \right.$$

Nous voulons démontrer que f est discontinue en tout point de $\mathbb{R}$. Soit a un réel. Nous devons prouver³ qu'il existe un réel $\varepsilon_0 > 0$ tel que, pour tout réel $\eta > 0$, il existe au moins un nombre réel x tel que $|x - a| < \eta$ et $|f(x) - f(a)| \geq \varepsilon_0$.

Rappelons tout d'abord qu'entre deux nombres réels distincts, il existe au moins un nombre rationnel et au moins un irrationnel — donc une infinité de rationnels et une infinité d'irrationnels.

Nous raisonnons dans le cas où a est rationnel (resp. irrationnel). Nous posons $\varepsilon_0 = 1$. Soit η un nombre réel > 0 . Le rappel ci-dessus nous autorise à choisir un irrationnel (resp. un rationnel) x tel que $a - \eta < x < a + \eta$. Alors $|x - a| < \eta$ et $|f(x) - f(a)| = |0 - 1| = 1$ (resp. $|f(x) - f(a)| = |1 - 0| = 1$), donc $|f(x) - f(a)| \geq \varepsilon_0$. En conclusion, f est discontinue en a . $\square$

L'exemple que nous venons de voir s'applique à une fonction qui ne prend que deux valeurs, ici 0 et 1. Nous allons constater que même une bijection d'une partie infinie $\mathcal{A}$ de $\mathbb{R}$ sur $\mathcal{A}$ peut n'être continue en aucun point.

1. Voir l'exemple 15.7, page 294.

2. Cette fonction a été introduite en 1829 par le mathématicien allemand Gustav Lejeune Dirichlet (1805-1859) ; voir aussi l'exemple 11.1, page 208.

3. Si l'on trouve que l'expression en français de la négation de l'assertion « f est continue en a » est trop difficile, on peut l'obtenir en niant, de façon automatique, le $(\forall \varepsilon > 0) \dots$, ce qui donne :

$$(\exists \varepsilon > 0)(\forall \eta > 0)(\exists x \in \mathbb{R})(|x - a| < \eta \text{ et } |f(x) - f(a)| \geq \varepsilon).$$

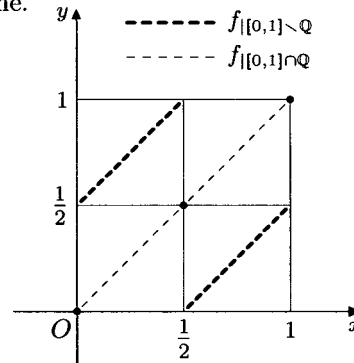
8.2. Bijection de $[0, 1]$ sur $[0, 1]$ qui n'est continue en aucun point du segment $[0, 1]$.

Nous notons E la fonction « partie entière » et nous introduisons l'application :

$$f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} x & \text{si } x \text{ est rationnel,} \\ x + \frac{1}{2} - E\left(x + \frac{1}{2}\right) & \text{si } x \text{ est irrationnel.} \end{cases}$$

Pour tout réel y , $E(y) \leq y < 1 + E(y)$ donc $0 \leq y - E(y) < 1$. Il en résulte que f est une application du segment $[0, 1]$ dans lui-même.

Nous ne pouvons bien sûr pas représenter le graphe de f , mais nous pouvons l'imaginer sur $[0, 1] \cap \mathbb{Q}$ et sur $[0, 1] \setminus \mathbb{Q}$ par le dessin ci-contre, sur lequel la représentation graphique de la restriction de f à $[0, 1] \setminus \mathbb{Q}$ est en pointillé gras et celle de la restriction de f à $[0, 1] \cap \mathbb{Q}$ en pointillé plus fin.



Pour prouver la discontinuité de f , nous allons cette fois-ci utiliser la caractérisation de la continuité par les suites (théorème 8.1).

Soit a un point du segment $[0, 1]$. Supposons

d'abord que a soit rationnel et que $0 \leq a < 1/2$. Le réel e est irrationnel. La suite $(x_n)_{n \geq 1}$ de terme général $x_n = a + (e/n)$ est à valeurs dans $\mathbb{R} \setminus \mathbb{Q}$ et converge vers a et on a, à partir d'un certain rang, $a < x_n < 1/2$ donc $f(x_n) = x_n + (1/2)$, ce qui montre que la suite image $(f(x_n))$ converge vers $a + (1/2) \neq a = f(a)$. On opère de même lorsque a est rationnel et appartient à $[1/2, 1[$ à l'aide de la même suite (x_n) , et enfin, pour $a = 1$, en utilisant la suite $(x_n)_{n \geq 1}$ de terme général $x_n = 1 - (e/n)$. Si a est un irrationnel appartenant à $[0, 1]$, la suite $(r_n)_{n \geq 0}$ des valeurs décimales approchées par défaut de a est à valeurs dans $[0, 1] \cap \mathbb{Q}$ et converge vers a , et $f(r_n) = r_n$ pour tout n , donc la suite $(f(r_n))$ converge vers a , alors que $f(a) = a + (1/2) \neq a$ si $a < 1/2$ et $f(a) = a - (1/2) \neq a$ si $a > 1/2$. En conclusion, dans tous les cas, la fonction f est discontinue en a .

Il reste à démontrer que f est une bijection de $[0, 1]$ sur $[0, 1]$. Soit x un point de $[0, 1]$. Si x est rationnel, $f(x) = x$ donc $f(x)$ est rationnel, ce qui montre que $f(f(x)) = f(x) = x$. Si x est irrationnel et si $0 < x < 1/2$, $f(x) = x + (1/2)$ est irrationnel et $1/2 < f(x) < 1$, donc $f(f(x)) = f(x) - (1/2) = x$; enfin, si x est irrationnel et si $1/2 < x < 1$, $f(x) = x - (1/2)$ est irrationnel et $0 < f(x) < 1/2$, donc $f(f(x)) = f(x) + (1/2) = x$. Finalement, la composée $f \circ f$ est l'identité de $[0, 1]$, donc f est une bijection de $[0, 1]$ sur lui-même⁴. $\square$

Les exemples 8.1 et 8.2 utilisent le fait que $\mathbb{Q}$ et son complémentaire sont tous les deux denses dans $\mathbb{R}$ et qu'une fonction définie sur un ensemble dense ne peut avoir qu'un seul prolongement continu.

Pour une fonction f définie au voisinage d'un réel a , la notion de continuité en ce point, bien qu'elle soit locale, c'est-à-dire qu'elle ne dépende que des valeurs

4. Une application $f : E \rightarrow E$ telle que $f \circ f = \text{Id}_E$ est une involution de E ; c'est alors une bijection de E sur E et on a $f^{-1} = f$.

prises par f sur un voisinage de a , fait entrer en ligne de compte la valeur de f sur les points « très proches » de a ; cependant, elle ne permet pas de conclure à la continuité de f sur un voisinage de a comme le montre l'exemple suivant.

8.3. Fonction continue en 0 qui n'est continue en aucun autre point.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} x & \text{si } x \text{ est rationnel,} \\ 0 & \text{si } x \text{ est irrationnel.} \end{cases} \end{array} \right.$$

Remarquons que $f(0) = 0$ et que, pour tout réel x , $|f(x)| \leq |x|$.

Soit ε un réel > 0 . Nous posons $\eta = \varepsilon$. Alors $\eta > 0$ et, pour tout nombre réel x , $|x| < \eta$ entraîne $|f(x)| \leq |x| < \eta = \varepsilon$. La fonction f est donc continue en 0.

Soit a un réel différent de zéro. Nous raisonnons dans le cas où a est rationnel (resp. irrationnel). Nous posons $\varepsilon_0 = |a|$; c'est un réel > 0 . Soit η un réel > 0 . Nous choisissons un irrationnel (resp. un rationnel) x tel que $a < x < a + \eta$ si $a > 0$ et $a - \eta < x < a$ si $a < 0$; par conséquent $|x| > |a|$. On a $|x - a| < \eta$ et $|f(x) - f(a)| = |0 - a| = |a|$ (resp. $|f(x) - f(a)| = |x - 0| = |x| > |a|$), donc $|f(x) - f(a)| \geq \varepsilon_0$. En conclusion, la fonction f est discontinue en a . $\square$

Dans l'exemple précédent 8.3, la fonction f est continue en 0 mais, dans tout voisinage de 0, il existe un élément en lequel f est discontinue. Cette propriété de f en 0 peut s'étendre à tout un ensemble dense.

8.4. Fonction continue en tout point de $\mathbb{R} \setminus \mathbb{Q}$ et discontinue sur les voisinages des nombres irrationnels.

Nous allons en fait construire une application de $\mathbb{R}$ dans $\mathbb{R}$, continue en tout point de $\mathbb{R} \setminus \mathbb{Q}$ et discontinue en chaque point de $\mathbb{Q}$.

Nous utilisons dans cet exemple la notion de limite — voir la suite du présent chapitre — et les propriétés des suites et des séries de fonctions, en particulier celles de la convergence uniforme, qui sont rappelées aux chapitres 12 et 13.

Nous notons E la fonction « partie entière » et nous introduisons l'application :

$$h : x \mapsto h(x) = x - E(x)$$

de $\mathbb{R}$ dans $\mathbb{R}$ et la suite $(u_n)_{n \geq 1}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} u_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_n(x) = \frac{h(nx)}{2^n}. \end{array} \right.$$

On a, pour tout réel y , $0 \leq h(y) < 1$, donc, pour tout réel x et tout entier $n \geq 1$, $0 \leq u_n(x) < 1/2^n$. Ainsi la série de fonctions $\sum_{n \geq 1} u_n$ converge normalement, donc uniformément, sur $\mathbb{R}$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \sum_{n=1}^{+\infty} u_n(x). \end{array} \right.$$

La fonction f est continue en tout réel a pour lequel, quel que soit l'entier $n \geq 1$, la fonction u_n est continue en a . Les points de discontinuité de la fonction h sont les entiers relatifs. Si a est un nombre réel irrationnel, alors, pour tout entier $n \geq 1$, na est irrationnel donc h est continue en na , ce qui montre que u_n est continue en a , d'où l'on déduit que f est continue en a . Ceci prouve la continuité de f en tous les points de $\mathbb{R} \setminus \mathbb{Q}$.

Soit a un nombre rationnel. Notons (p, q) le représentant irréductible de a de dénominateur positif. Les entiers p et q sont premiers entre eux, donc, pour tout entier $n \geq 1$, na appartient à $\mathbb{Z}$ si, et seulement si, q divise n , ce qui équivaut à l'existence de $k \in \mathbb{N}^*$ tel que $n = qk$. Pour tout entier $n \geq 1$ tel que na appartient à $\mathbb{Z}$, $h(na) = 0$ et $\lim_{(x \rightarrow a^-)} h(nx) = 1$, donc $u_n(a) = 0$ et $\lim_{(x \rightarrow a^-)} u_n(x) = 1/2^n$. La convergence uniforme de la série de fonctions nous permet d'affirmer que :

$$f(a^-) = \lim_{x \rightarrow a^-} f(x) = \sum_{n=1}^{+\infty} \lim_{x \rightarrow a^-} u_n(x)$$

On en déduit que :

$$\begin{aligned} f(a) - f(a^-) &= \sum_{n=1}^{+\infty} (u_n(a) - \lim_{x \rightarrow a^-} u_n(x)) = \sum_{\substack{n \in \mathbb{N}^* \\ na \in \mathbb{Z}}} \frac{1}{2^n} \\ &= \sum_{k=1}^{+\infty} \frac{1}{2^{qk}} = \sum_{k=1}^{+\infty} \frac{1}{(2^q)^k} = \frac{1}{2^q} \frac{1}{1 - \frac{1}{2^q}} = \frac{1}{2^q - 1} > 0. \end{aligned}$$

Par conséquent, la limite de f en a à gauche n'est pas égale à $f(a)$, donc f n'est pas continue en a . En conclusion, f est discontinue en tous les nombres rationnels. $\square$

■ Théorème des valeurs intermédiaires

THÉORÈME 8.2. — Théorème des valeurs intermédiaires⁵.

Si I est un intervalle et f une fonction de $\mathbb{R}$ dans $\mathbb{R}$ définie et continue sur I , alors, quels que soient les points a et b de I , le segment $[f(a), f(b)]$ est inclus dans $f([a, b])$ — ce qui signifie que⁶, pour tout réel y compris entre $f(a)$ et $f(b)$, il existe au moins un réel x compris entre a et b tel que $y = f(x)$.

Il en résulte que si I est un intervalle et f une fonction définie et continue sur I , l'image directe $f(I)$ de I par f est un intervalle.

On dit d'une fonction f définie sur un intervalle I qu'elle vérifie le théorème des valeurs intermédiaires sur I si, quels que soient les points a et b de I , le segment $[f(a), f(b)]$ est inclus dans $f([a, b])$. Il peut se faire, comme nous le verrons dans les exemples suivants, qu'une telle fonction ne soit pas continue sur cet intervalle⁷.

5. Ce théorème fut longtemps utilisé sans autre justification que graphique : on invoquait « l'évidence géométrique ». La première tentative de preuve est due à Bernhard Bolzano en 1817.

6. La propriété qui suit explique le nom du théorème 8.2.

7. Le mathématicien français Ossian Bonnet a démontré, en 1868, que la fonction dérivée d'une fonction définie et dérivable sur un intervalle vérifie sur cet intervalle le théorème des valeurs intermédiaires, même si elle n'est pas continue ; voir à la page 168 ce qui suit les exemples 9.10 et 9.11. Ce résultat est appelé le théorème de Darboux ; il est en effet souvent attribué à Gaston Darboux, qui a donné le premier exemple de fonction discontinue vérifiant le théorème des valeurs intermédiaires.

Pour *démontrer* qu'une fonction f définie sur un intervalle I vérifie le théorème des valeurs intermédiaires sur I , il *suffit* bien sûr de prouver que, quels que soient les points a et b de I tels que $a < b$, $[f(a), f(b)]$ est inclus dans $f([a, b])$.

8.5. Fonction discontinue vérifiant le théorème des valeurs intermédiaires.

Nous introduisons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} \cos \frac{1}{x} & \text{si } x \neq 0, \\ 1 & \text{si } x = 0. \end{cases} \end{array} \right.$$

Pour tout réel $a \neq 0$, la fonction $x \mapsto 1/x$ est continue en a , donc aussi la fonction f . La suite $(u_n)_{n \geq 1}$, de terme général $u_n = 1/(n\pi)$, converge vers 0 ; or, pour tout $n \in \mathbb{N}^*$, $f(u_n) = \cos(n\pi) = (-1)^n$, donc la suite $(f(u_n))$ ne converge pas dans $\mathbb{R}$. Il en résulte que la fonction f est discontinue en 0 (théorème 8.1, page 134). Ainsi, la fonction f est continue en tout point de $\mathbb{R} \setminus \{0\}$ et discontinue en 0.

Soit a et b des réels tels que $a < b$. Clairement, $f([a, b]) \subset [-1, 1]$. Si $0 \notin [a, b]$, alors $a > 0$ ou $b < 0$, donc $[a, b]$ est inclus dans $]0, +\infty[$ ou $]-\infty, 0[$, d'où l'on déduit que f est continue sur $[a, b]$, ce qui, en appliquant le théorème des valeurs intermédiaires à la restriction de f à l'intervalle $[a, b]$, montre que $[f(a), f(b)] \subset f([a, b])$. Il reste à examiner le cas où $0 \in [a, b]$, c'est-à-dire $a \leq 0 \leq b$. Supposons que $b \neq 0$. La suite $(u_n)_{n \geq 1}$, de terme général $u_n = 1/(n\pi)$, est strictement décroissante et converge vers 0. Comme $b > 0$, il existe des entiers $n \geq 1$ tels que $u_n < b$ donc $a \leq 0 < u_{n+1} < u_n < b$; nous choisissons un tel n . L'image de $[n\pi, (n+1)\pi]$ par cosinus est $[-1, 1]$, donc l'image du segment $[u_{n+1}, u_n]$ par f est $[-1, 1]$. Par suite $[-1, 1] = f([u_{n+1}, u_n]) \subset f([a, b]) \subset [-1, 1]$, donc $f([a, b]) = [-1, 1]$; or $f(a)$ et $f(b)$ appartiennent à $[-1, 1]$, donc $[f(a), f(b)] \subset [-1, 1] = f([a, b])$. Enfin, si $b = 0$, alors $a < 0$, et en remplaçant b par a et $(u_n)_{n \geq 1}$ par la suite opposée, on prouve comme dans ce qui précède que $[f(a), f(b)] \subset [-1, 1] = f([a, b])$.

Finalement, la fonction f vérifie le théorème des valeurs intermédiaires sur $\mathbb{R}$. $\square$

8.6. Fonction vérifiant le théorème des valeurs intermédiaires et discontinue en tous les points⁸.

Si $x \in [0, 1[$, le développement décimal propre de x est l'unique développement décimal illimité $x = 0, a_1 a_2 a_3 \dots$ pour lequel il n'existe aucun entier $p \geq 1$ tel que $a_n = 9$ pour tout entier $n \geq p$.

Soit $x \in [0, 1[$, de développement décimal propre $x = 0, a_1 a_2 a_3 \dots$. Si la suite $(a_{2n-1})_{n \geq 1}$ des décimales impaires de x n'est pas périodique à partir d'un certain rang, on pose $f(x) = 0$; sinon, p étant le plus petit des entiers naturels k tels que cette suite est périodique à partir du rang k , on pose $f(x) = 0, a_{2p} a_{2p+2} a_{2p+4} \dots$. En posant $f(1) = 0$, on obtient ainsi une application f de $[0, 1]$ dans $[0, 1]$.

Soit a et b des points du segment $[0, 1]$ tels que $a < b$. On a donc $0 \leq a < b \leq 1$.

8. Cet exemple est dû à Henri Lebesgue, en 1904.

La suite $(1/10^{2(n-1)})$ tend vers 0, ce qui justifie le choix d'un entier $n \geq 1$ tel que :

$$\frac{1}{10^{2(n-1)}} < \frac{b-a}{2}.$$

Par la propriété d'Archimède, il existe $k \in \mathbb{Z}$ tel que $a \leq \frac{k}{10^{2(n-1)}} < \frac{k+1}{10^{2(n-1)}} < b$.

On a $0 \leq k/10^{2(n-1)} < 1$, donc $k \in \llbracket 0, 10^{2(n-1)} - 1 \rrbracket$. Par suite il existe des chiffres $a_1, a_2, \dots, a_{2(n-1)} \in \llbracket 0, 9 \rrbracket$ tels qu'en base 10, $k = a_1 a_2 a_3 \dots a_{2(n-1)}$. On a alors :

$$\frac{k}{10^{2(n-1)}} = 0, a_1 a_2 a_3 \dots a_{2(n-1)}.$$

Soit y un point de $]0, 1[$, de développement décimal propre $y = 0, b_1 b_2 b_3 \dots$. Nous posons :

$$x = 0, a_1 a_2 a_3 \dots a_{2n-2} \alpha b_1 \beta b_2 \alpha b_3 \beta b_4 \alpha b_5 \beta b_6 \alpha b_7 \beta b_8 \dots = 0, c_1 c_2 c_3 \dots$$

où α et β sont choisis dans $\llbracket 0, 9 \rrbracket$ de telle sorte que $\alpha \neq a_{2n-3}$, $\beta \neq a_{2n-1}$ et $\alpha \neq \beta$.

La suite $(c_{2k-1})_{k \geq 1}$ des décimales impaires de x est périodique à partir du rang n — et pas avant —, de séquence $\alpha\beta$, donc $f(x) = 0, b_1 b_2 b_3 b_4 b_5 b_6 b_7 b_8 \dots = y$. De plus, les $2n - 2$ premières décimales de x sont celles de $k/10^{2(n-1)}$, donc :

$$a \leq \frac{k}{10^{2(n-1)}} \leq x < \frac{k+1}{10^{2(n-1)}} < b,$$

ce qui montre que y appartient à $f([a, b])$.

Par conséquent, $]0, 1[$ est inclus dans $f([a, b])$. Or $f([a, b]) \subset f([0, 1]) \subset [0, 1]$, donc $]0, 1[\subset f([a, b]) \subset [0, 1]$. Il en résulte que $J = f([a, b]) =]0, 1[, [0, 1[,]0, 1]$ ou $[0, 1]$, donc J est un intervalle et, comme $f(a)$ et $f(b)$ appartiennent à J , $[f(a), f(b)]$ est inclus dans $J = f([a, b])$.

Il en résulte que f vérifie le théorème des valeurs intermédiaires sur $[0, 1]$.

Soit c un point du segment $[0, 1]$. Nous posons $\varepsilon_0 = 1/4$. Soit η un réel > 0 . Nous posons $\alpha = \eta/2$, $a = \text{Max}(0, c - \alpha)$ et $b = \text{Min}(1, c + \alpha)$. Alors $0 \leq a < b \leq 1$ et, pour tout point t du segment $[a, b]$, $|t - c| \leq \alpha < \eta$. Si $f(c) \leq 1/2$ (resp. $f(c) > 1/2$), nous choisissons un point y de $]3/4, 1[$ (resp. de $]0, 1/4[$) et un antécédent x , appartenant à $[a, b]$, de y par f — un tel x existe, puisque $]0, 1[\subset f([a, b])$; voir ce qui précède. Comme $x \in [a, b]$, $|x - c| < \eta$ et, par le choix du point y , $|y - f(c)| > 1/4$ donc $|f(x) - f(c)| \geq \varepsilon_0$. Par suite, f est discontinue en c .

En conclusion, f n'est continue en aucun point de l'intervalle $[0, 1]$. $\square$

8.7. Deux fonctions qui vérifient le théorème des valeurs intermédiaires, mais pas leur somme.

Nous introduisons les applications :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} \cos \frac{1}{x} & \text{si } x \neq 0, \\ 1 & \text{si } x = 0 \end{cases} \end{array} \right.$$

et :

$$\left| \begin{array}{l} g : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} -\cos \frac{1}{x} & \text{si } x \neq 0, \\ 1 & \text{si } x = 0. \end{cases} \end{array} \right.$$

Nous avons vu dans l'exemple 8.5 (page 138) que f vérifie le théorème des valeurs intermédiaires sur $\mathbb{R}$, et on prouve de manière identique qu'il en est de même de g .

On a $(f + g)(0) = 2$ et, pour tout réel $x \neq 0$, $(f + g)(x) = 0$. Nous choisissons un réel a différent de zéro. Alors $(f + g)(a) = 0$, donc $[(f + g)(0), (f + g)(a)] = [0, 2]$, qui n'est pas inclus dans $(f + g)([0, a]) = \{0, 2\}$. En conclusion, la fonction somme $f + g$ ne vérifie pas sur $\mathbb{R}$ le théorème des valeurs intermédiaires. $\square$

■ Limites

DÉFINITION 8.4. — Si $\mathcal{A}$ est une partie de $\mathbb{R}$, a un réel adhérent à $\mathcal{A}$, f une fonction définie sur $\mathcal{A}$ et ℓ un nombre réel, $f(x)$ admet pour limite ℓ dans $\mathbb{R}$ quand x tend vers a suivant $\mathcal{A}$ si, pour tout réel $\varepsilon > 0$, il existe un nombre réel $\eta > 0$ tel que, pour tout réel x , $x \in \mathcal{A}$ et $|x - a| < \eta$ entraînent $|f(x) - \ell| < \varepsilon$.

Il y a bien sûr unicité de la limite éventuelle dans $\mathbb{R}$ de f en a suivant $\mathcal{A}$ — c'est-à-dire de la limite de $f(x)$ dans $\mathbb{R}$ quand x tend vers a suivant $\mathcal{A}$ — et cette limite se note $\lim_{\substack{x \rightarrow a \\ x \in \mathcal{A}}} f(x)$ ou simplement, s'il n'y a pas d'ambiguïté sur $\mathcal{A}$, $\lim_{x \rightarrow a} f(x)$.

On dit aussi que $f(x)$ tend vers ℓ dans $\mathbb{R}$ quand x tend vers a suivant $\mathcal{A}$ pour exprimer que $f(x)$ admet pour limite ℓ dans $\mathbb{R}$ quand x tend vers a suivant $\mathcal{A}$.

Si $\mathcal{A}$ est une partie de $\mathbb{R}$, a un point de $\mathcal{A}$ et f une fonction définie sur $\mathcal{A}$, $f(x)$ admet une limite ℓ dans $\mathbb{R}$ quand x tend vers a suivant $\mathcal{A}$ si f est continue en a , ce qui montre que $f(a) = \ell$; la continuité est donc un cas particulier de limite.

DÉFINITION 8.5. — Soit f une fonction et a un nombre réel.

a) Si le domaine $\mathcal{D}$ de f contient un intervalle $]a, b[$ où $b > a$, et si $f(x)$ admet une limite ℓ dans $\mathbb{R}$ quand x tend vers a suivant $\mathcal{D} \cap]a, +\infty[$, ℓ s'appelle la limite de f en a à droite, ou la limite de $f(x)$ quand x tend vers a à droite, et cette limite, qui est unique, se note $\lim_{a^+} f$, $\lim_{x \rightarrow a^+} f(x)$ ou $\lim_{\substack{x \rightarrow a \\ x > a}} f(x)$.

b) Si le domaine $\mathcal{D}$ de f contient un intervalle $]b, a[$ où $b < a$, et si $f(x)$ admet une limite ℓ dans $\mathbb{R}$ quand x tend vers a suivant $\mathcal{D} \cap]-\infty, a[$, ℓ s'appelle la limite de f en a à gauche, ou la limite de $f(x)$ quand x tend vers a à gauche, et cette limite, qui est unique, se note $\lim_{a^-} f$, $\lim_{x \rightarrow a^-} f(x)$ ou $\lim_{\substack{x \rightarrow a \\ x < a}} f(x)$.

Les hypothèses étant celles du a) (resp. du b)) de la définition 8.5, le réel ℓ est la limite de f en a à droite (resp. à gauche) si, pour tout réel $\varepsilon > 0$, il existe un réel $\eta > 0$ tel que, pour tout nombre réel x :

$$a < x < a + \eta \quad (\text{resp. } a - \eta < x < a) \quad \text{entraîne} \quad |f(x) - \ell| < \varepsilon.$$

Si a est un réel et f une fonction définie sur un voisinage de a , f est continue en a si, et seulement si, f admet le réel $f(a)$ pour limite en a à gauche et à droite.

On peut remplacer le réel a par $+\infty$ ou $-\infty$. On dit d'une fonction f qu'elle est définie au voisinage de $+\infty$ (resp. de $-\infty$) s'il existe un réel A_0 (resp. B_0) tel que f est définie sur $]A_0, +\infty[$ (resp. sur $] -\infty, B_0[$).

DÉFINITION 8.6. — Si f est une fonction définie au voisinage de $+\infty$ (resp. de $-\infty$) et ℓ un nombre réel, $f(x)$ admet pour limite ℓ dans $\mathbb{R}$ quand x tend vers $+\infty$ (resp. vers $-\infty$) si, pour tout réel $\varepsilon > 0$, il existe un nombre réel A (resp. B) tel que $|f(x) - \ell| < \varepsilon$ pour tout réel $x > A$ (resp. pour tout réel $x < B$).

8.8. Fonction admettant une limite à gauche et une limite à droite en un point, mais discontinue en ce point.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 1 & \text{si } x = 0, \\ 0 & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

Les limites de f en 0 à gauche et à droite existent et sont égales à 0, différent de $f(0)$. De ce fait, f est discontinue en 0. $\square$

8.9. Fonction n'admettant ni limite à droite, ni limite à gauche en 0.

Nous introduisons l'application :

$$\left| \begin{array}{l} f : \mathbb{R}^* \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \sin \frac{1}{x}. \end{array} \right.$$

Nous considérons les suites $(a_n)_{n \geq 0}$ et $(b_n)_{n \geq 0}$, de termes généraux :

$$a_n = \frac{2}{(4n+1)\pi} \text{ et } b_n = \frac{2}{(4n+3)\pi}.$$

On a, pour tout entier naturel n , $0 < b_n < a_n$, $f(a_n) = 1$ et $f(b_n) = -1$, donc $f(a_n) - f(b_n) = 2$; de plus, la suite (a_n) converge vers 0.

Soit ℓ un nombre réel. Nous posons $\varepsilon_0 = 1$. Soit η un réel > 0 . Nous choisissons un entier naturel N tel que $a_n < \eta$ pour tout entier $n \geq N$; en particulier, $a_N < \eta$, donc $0 < a_N < \eta$ et $0 < b_N < \eta$. En notant c le milieu du bipoint $(f(b_N), f(a_N))$, on a $f(b_N) < c < f(a_N)$, $c - f(b_N) = 1$ et $f(a_N) - c = 1$, donc, en raisonnant dans les cas $\ell \leq c$ et $c < \ell$, on voit que $|f(a_N) - \ell| \geq 1$ ou $|f(b_N) - \ell| \geq 1$. Nous avons trouvé un réel x tel que $0 < x < \eta$ et $|f(x) - \ell| \geq \varepsilon_0$. Par conséquent la fonction f n'admet pas pour limite ℓ en 0 à droite. En conclusion, f n'admet pas de limite en 0 à droite et, comme f est impaire, il en est de même pour «à gauche». $\square$

8.10. Fonction n'admettant ni limite à gauche, ni limite à droite en aucun point de $\mathbb{R}$.

Nous reprenons la fonction de Dirichlet f de l'exemple 8.1 (page 134) :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 1 & \text{si } x \text{ est rationnel,} \\ 0 & \text{si } x \text{ est irrationnel.} \end{cases} \end{array} \right.$$

Soit a un nombre réel. Soit ℓ un nombre réel. Nous posons $\varepsilon_0 = 1/2$. Soit η un nombre réel > 0 . Nous choisissons un rationnel r et un irrationnel α dans $]a, a + \eta[$. Alors $f(r) - f(\alpha) = 1$. En notant c le milieu du bipoint $(f(\alpha), f(r))$, on a $f(\alpha) < c < f(r)$, $c - f(\alpha) = 1/2$ et $f(r) - c = 1/2$, donc, en raisonnant dans les cas $\ell \leq c$ et $c < \ell$, on voit que $|f(r) - \ell| \geq 1/2$ ou $|f(\alpha) - \ell| \geq 1/2$. Nous avons trouvé un réel x tel que $a < x < a + \eta$ et $|f(x) - \ell| \geq \varepsilon_0$. Par suite la fonction f n'admet pas pour limite ℓ en a à droite. Ceci prouve que f n'admet pas de limite en a à droite. On opère de même, en remplaçant $]a, a + \eta[$ par $]a - \eta, a[$, pour établir que f n'admet pas de limite en a à gauche. $\square$

Si $\mathcal{A}$ est une partie de $\mathbb{R}$, a un réel adhérent à $\mathcal{A}$, f une fonction définie sur $\mathcal{A}$ et ℓ un nombre réel, et si $f(x)$ admet pour limite ℓ quand x tend vers a suivant $\mathcal{A}$, alors, pour toute suite (x_n) à valeurs dans $\mathcal{A}$ qui converge vers a , la suite $(f(x_n))$ converge vers ℓ . De la même façon, si f est une fonction et a un nombre réel, si le domaine de f contient un intervalle $]a, b[$ où $b > a$ (resp. un intervalle $]b, a[$ où $b < a$) et si f admet pour limite un réel ℓ en a à droite (resp. à gauche), alors, pour toute suite (x_n) à valeurs dans $]a, +\infty[$ (resp. $]-\infty, a[$) qui converge vers a , la suite $(f(x_n))$ converge vers ℓ . Enfin, si f est une fonction définie au voisinage de $+\infty$ (resp. de $-\infty$) et si f admet pour limite un réel ℓ en $+\infty$ (resp. en $-\infty$), alors, pour toute suite (x_n) à valeurs réelles qui tend vers $+\infty$ (resp. vers $-\infty$), la suite $(f(x_n))$ converge vers ℓ . En fait, toutes ces implications sont des équivalences ; il s'agit de la caractérisation des limites par les suites.

On en déduit par exemple que si f est une fonction définie au voisinage de $+\infty$ et si f admet pour limite un nombre réel ℓ en $+\infty$, alors, pour tout nombre réel $x > 0$, la suite $(f(nx))$ converge vers ℓ . La réciproque est fautive en général, comme le montre l'exemple suivant.

8.11. Fonction n'admettant pas pour limite 0 en $+\infty$ mais telle que, pour tout réel $x > 0$, $\lim_{n \rightarrow \infty} f(nx) = 0$.

Nous posons $\mathcal{A} = \{n + \pi \mid n \in \mathbb{N}\}$ et nous considérons l'application :

$$\left| \begin{array}{l} f :]0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 1 & \text{si } x \in \mathcal{A}, \\ 0 & \text{sinon.} \end{cases} \end{array} \right.$$

La suite $(n + \pi)$ tend vers $+\infty$ alors que la suite $(f(n + \pi))$ converge vers 1, donc f ne tend pas vers 0 en $+\infty$.

Soit x un réel > 0 . Supposons l'existence de deux entiers naturels distincts p et q tels que px et qx appartiennent tous deux à $\mathcal{A}$. Il existe alors des entiers naturels distincts m et n tels que $px = m + \pi$ et $qx = n + \pi$, donc $q(m + \pi) = p(n + \pi) (= pqx)$, d'où l'égalité $(p - q)\pi = qm - pn$, ou encore $\pi = (qm - pn)/(p - q)$, en contradiction avec l'irrationalité de π . Ainsi, $f(nx) = 0$ pour tout entier naturel n , sauf peut-être pour une seule valeur de n ; on en déduit que $\lim_{n \rightarrow \infty} f(nx) = 0$. $\square$

DÉFINITION 8.7. — Si a est un nombre réel et f une fonction définie sur un voisinage de a , f est continue à droite (resp. à gauche) en a si $f(x)$ admet pour limite $f(a)$ quand x tend vers a à droite (resp. à gauche).

8.12. Fonction n'admettant pas de limite à droite ni à gauche en chaque point de $\mathbb{Q}$, mais continue en tout point de $\mathbb{R} \setminus \mathbb{Q}$.

Nous utilisons dans cet exemple les propriétés de la convergence uniforme des suites et séries de fonctions, rappelées aux chapitres 12 et 13.

Nous introduisons l'application :

$$\left| \begin{array}{l} g : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} 0 & \text{si } x = 0, \\ \sin \frac{1}{x} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

Remarquons que cette fonction est continue en tout point de $\mathbb{R}^*$, mais qu'elle n'admet en 0 ni limite à droite, ni limite à gauche — en effet, les suites $(x_n)_{n \geq 0}$ et $(y_n)_{n \geq 0}$, de termes généraux $x_n = 1/((\pi/2) + n\pi)$ et $y_n = -x_n$, à valeurs dans $]0, +\infty[$ pour (x_n) et dans $]-\infty, 0[$ pour (y_n) , convergent vers 0 alors que, pour tout n , $g(x_n) = (-1)^n$ et $g(y_n) = (-1)^{n+1}$, termes généraux de suites divergentes. L'ensemble $\mathbb{Q}$ des nombres rationnels est dénombrable. Nous choisissons une bijection φ de $\mathbb{N}$ sur $\mathbb{Q}$, nous posons $c_n = \varphi(n)$ pour tout entier naturel n et nous considérons la suite $(u_n)_{n \geq 1}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} u_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_n(x) = \frac{g(x - c_n)}{2^n}. \end{array} \right.$$

On a, pour tout réel y , $|g(y)| \leq 1$, donc, pour tout nombre réel x et tout entier $n \geq 1$, $|u_n(x)| \leq 1/2^n$. Par conséquent la série de fonctions $\sum_{n \geq 1} u_n$ converge normalement, donc uniformément, sur $\mathbb{R}$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \sum_{n=1}^{+\infty} u_n(x). \end{array} \right.$$

Si a est un nombre réel irrationnel, alors, pour tout $n \in \mathbb{N}$, $a - c_n \neq 0$, donc u_n est continue en a , et la convergence uniforme de la série de fonctions nous assure la continuité de f en a .

Soit a un nombre rationnel. Il existe un entier naturel p et un seul tel que $a = c_p$. Comme ci-dessus, la fonction $x \mapsto f(x) - u_p(x)$ est continue en a ; or la fonction u_p n'admet en a ni limite à gauche, ni limite à droite, ce qui entraîne la même propriété pour la fonction f . $\square$

8.13. Fonction continue à gauche en tout réel qui n'admet pas de limite à droite en tous les points d'un ensemble dense.

Reprenons l'exemple précédent 8.12 en remplaçant la fonction g de cet exemple par :

$$\left| \begin{array}{l} g : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} 0 & \text{si } x \leq 0, \\ \sin \frac{1}{x} & \text{si } x > 0. \end{cases} \end{array} \right.$$

Le même raisonnement que dans l'exemple 8.12 montre que, pour chaque nombre rationnel a , f n'admet pas de limite en a à droite. Que le nombre réel a soit rationnel ou irrationnel, la fonction u_n est, pour tout n , continue à gauche en a , et la convergence uniforme de la série de fonctions nous assure qu'il en est de même pour f : la fonction f est continue à gauche en tout point de $\mathbb{R}$. $\square$

■ Continuité uniforme, continuité absolue et fonctions lipschitziennes

Pour la conservation de certaines propriétés, la continuité est insuffisante et l'on est amené à introduire d'autres conditions plus fortes, comme la continuité uniforme, l'absolue continuité et les fonctions lipschitziennes. La continuité

uniforme⁹ est utile pour conserver les suites de Cauchy et la continuité absolue¹⁰ est surtout intéressante pour la conservation de propriétés métriques ; son intérêt principal est qu'une fonction absolument continue est dérivable presque partout, que sa fonction dérivée est intégrable au sens de Lebesgue et que son intégrale indéfinie est égale à la fonction initiale à une constante additive près.

Nous rappelons trois définitions pour une fonction f définie sur une partie $\mathcal{A}$ de $\mathbb{R}$.

DÉFINITION 8.8. — La fonction f est uniformément continue sur $\mathcal{A}$ si, pour tout nombre réel $\varepsilon > 0$, il existe un réel $\eta > 0$ tel que, quels que soient les points x et y de $\mathcal{A}$, $|x - y| < \eta$ entraîne $|f(x) - f(y)| < \varepsilon$.

Pour les fanatiques du symbolisme logique, cette définition s'écrit :

$$(\forall \varepsilon > 0)(\exists \eta > 0)(\forall x \in \mathcal{A})(\forall y \in \mathcal{A})(|x - y| < \eta \implies |f(x) - f(y)| < \varepsilon).$$

La continuité uniforme sur $\mathcal{A}$ est plus forte que la continuité en tout point a de $\mathcal{A}$; en effet, dans le premier cas, η ne dépend pas de a , alors qu'il peut en dépendre dans le second. La continuité uniforme sur $\mathcal{A}$ entraîne donc la continuité sur $\mathcal{A}$.

DÉFINITION 8.9. — La fonction f est absolument continue sur $\mathcal{A}$ si, pour tout réel $\varepsilon > 0$, il existe un réel $\eta > 0$ tel que, pour tout entier $n \geq 1$ et quels que soient les points $x_1, \dots, x_n, y_1, \dots, y_n$ de $\mathcal{A}$ vérifiant $x_1 < y_1 \leq x_2 < y_2 \leq \dots \leq x_n < y_n$:

$$\sum_{i=1}^n (y_i - x_i) < \eta \text{ entraîne } \sum_{i=1}^n |f(y_i) - f(x_i)| < \varepsilon.$$

Clairement, la continuité absolue entraîne la continuité uniforme — il suffit d'appliquer celle-là à $n = 1$ pour obtenir celle-ci.

DÉFINITION 8.10. — Soit k un nombre réel positif. La fonction f est lipschitzienne¹¹ de rapport k sur $\mathcal{A}$ si, quels que soient les points x et y de $\mathcal{A}$:

$$|f(x) - f(y)| \leq k|x - y|.$$

On prouve facilement que si k est un réel positif et si f est lipschitzienne de rapport k sur $\mathcal{A}$, f est absolument continue sur $\mathcal{A}$. On a donc les implications suivantes :

- (1) Si f est lipschitzienne sur $\mathcal{A}$, f est absolument continue sur $\mathcal{A}$.
- (2) Si f est absolument continue sur $\mathcal{A}$, f est uniformément continue sur $\mathcal{A}$.
- (3) Si f est uniformément continue sur $\mathcal{A}$, f est continue sur $\mathcal{A}$.

Nous allons prouver que les trois implications réciproques sont en général fausses. Cependant, (3) est une équivalence pour un compact $\mathcal{A}$, en particulier un segment.

THÉORÈME 8.3. — Théorème de Heine¹².

Toute fonction définie et continue sur un compact $\mathcal{A}$ de $\mathbb{R}$ est uniformément continue sur $\mathcal{A}$.

9. Cette notion est introduite par le mathématicien allemand Eduard Heine (1821-1881), en 1872, pour démontrer le théorème 8.3 ci-après qui porte son nom.

10. C'est le mathématicien italien Giuseppe Vitali (1875-1932) qui définit cette notion.

11. Cette notion est introduite en 1864 par Rudolph Lipschitz (1832-1903) pour généraliser le théorème de Cauchy sur l'existence et l'unicité d'une solution d'une équation différentielle.

12. Pour une démonstration, voir [ARN2], §III.5.

8.14. Fonction continue qui n'est pas uniformément continue.

Nous considérons l'application $f : x \mapsto f(x) = x^2$ de $\mathbb{R}$ dans $\mathbb{R}$, continue sur $\mathbb{R}$. Pour prouver que f n'est pas uniformément continue sur $\mathbb{R}$, on doit trouver un réel $\varepsilon_0 > 0$ tel que, pour tout nombre réel $\eta > 0$, il existe des réels x et y pour lesquels $|x - y| < \eta$ et $|f(x) - f(y)| \geq \varepsilon_0$.

Nous posons $\varepsilon_0 = 1$. Soit η un réel > 0 . Si y est un réel positif et si $x = y + \frac{\eta}{2}$, alors $|x - y| = \eta/2 < \eta$ et :

$$|f(x) - f(y)| = f(x) - f(y) = \left(y + \frac{\eta}{2}\right)^2 - y^2 = \eta y + \frac{\eta^2}{4} \geq \eta y.$$

Nous choisissons un réel $y \geq 1/\eta$ et nous posons $x = y + (\eta/2)$. Alors $|x - y| < \eta$ et $|f(x) - f(y)| \geq \eta y \geq 1$, donc $|f(x) - f(y)| \geq \varepsilon_0$.

En conclusion, f est continue mais n'est pas uniformément continue sur $\mathbb{R}$. $\square$

8.15. Fonction uniformément continue qui n'est pas lipschitzienne.

La fonction racine carrée, de symbole $\sqrt{}$, est définie et continue sur $\mathbb{R}_+$.

Soit ε un réel > 0 . Le réel $\eta = \varepsilon^2$ est strictement positif. Soit x et y des points de $\mathbb{R}_+$ tels que $|x - y| < \eta$. Si $x < \varepsilon^2$ et $y < \varepsilon^2$, alors $0 \leq \sqrt{x} < \varepsilon$ et $0 \leq \sqrt{y} < \varepsilon$, donc $|\sqrt{x} - \sqrt{y}| < \varepsilon$; sinon, $x \geq \varepsilon^2$ ou $y \geq \varepsilon^2$, donc $0 < \varepsilon \leq \sqrt{x} + \sqrt{y}$, ce qui donne :

$$|\sqrt{x} - \sqrt{y}| = \left| \frac{x - y}{\sqrt{x} + \sqrt{y}} \right| = \frac{|x - y|}{\sqrt{x} + \sqrt{y}} \leq \frac{|x - y|}{\varepsilon} < \frac{\eta}{\varepsilon} = \varepsilon.$$

Par suite, dans tous les cas, $|\sqrt{x} - \sqrt{y}| \leq \varepsilon$. La fonction $\sqrt{}$ est donc uniformément continue sur $\mathbb{R}_+$.

Si a est un réel > 0 et si $x, y \in [a, +\infty[$, alors $\sqrt{x} + \sqrt{y} \geq 2\sqrt{a} > 0$, donc :

$$|\sqrt{x} - \sqrt{y}| = \left| \frac{(\sqrt{x})^2 - (\sqrt{y})^2}{\sqrt{x} + \sqrt{y}} \right| = \left| \frac{x - y}{\sqrt{x} + \sqrt{y}} \right| = \frac{|x - y|}{\sqrt{x} + \sqrt{y}} \leq \frac{1}{2\sqrt{a}} |x - y|.$$

Ainsi, pour tout réel $a > 0$, $\sqrt{}$ est lipschitzienne de rapport $\frac{1}{2\sqrt{a}}$ sur $[a, +\infty[$.

Supposons que $\sqrt{}$ soit lipschitzienne sur $\mathbb{R}_+$. Il existe alors un réel $k \geq 0$ tel que, quels que soient $x, y \in \mathbb{R}_+$, $|\sqrt{x} - \sqrt{y}| \leq k|x - y|$, d'où l'on déduit que, pour tout réel $x > 0$, $\sqrt{x} \leq kx$ donc $1 \leq k\sqrt{x}$, en contradiction avec $\lim_{x \rightarrow 0^+} \sqrt{x} = 0$. En conclusion, $\sqrt{}$ n'est pas lipschitzienne sur $\mathbb{R}_+$. $\square$

8.16. Fonction uniformément continue qui n'est pas absolument continue.

Nous considérons l'application :

$$\left| \begin{array}{l} f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ x \sin \frac{\pi}{x} & \text{si } x \in]0, 1]. \end{cases} \end{array} \right.$$

La fonction f est clairement continue sur $]0, 1]$. De plus, pour tout point x de l'intervalle $]0, 1]$, $|f(x)| = |x| |\sin(\pi/x)| \leq x$, donc $\lim_{(x \rightarrow 0^+)} f(x) = 0 = f(0)$, ce qui prouve la continuité de f en 0. Par conséquent la fonction f est continue sur le segment $[0, 1]$, donc uniformément continue sur $[0, 1]$ (théorème de Heine).

Nous introduisons les suites $(x_n)_{n \geq 0}$ et $(y_n)_{n \geq 1}$, de termes généraux :

$$x_n = \frac{2}{2n+1} \text{ et } y_n = \frac{2}{2n-1}.$$

Ces deux suites sont strictement décroissantes, (y_n) converge vers 0 et on a, pour tout entier $n \geq 2$, $0 < x_n < x_{n-1} = y_n \leq 1$.

Nous posons $\varepsilon_0 = 1$. Soit η un réel > 0 . Nous choisissons un entier $p \geq 2$ tel que $y_n < \eta$ pour tout entier $n \geq p$. Nous considérons la famille :

$$(x_{2p}, x_{2p-1}, \dots, x_p, y_{2p}, y_{2p-1}, \dots, y_p)$$

de points de $[0, 1]$. Alors $0 < x_{2p} < x_{2p-1} = y_{2p} < \dots < y_{p+1} = x_p < y_p$ et :

$$\sum_{k=0}^p (y_{p+k} - x_{p+k}) = y_p - x_{2p} < y_p < \eta.$$

On a, pour tout $n \in \mathbb{N}$, $f(x_n) = \frac{2}{2n+1} \sin \frac{(2n+1)\pi}{2} = \frac{2}{2n+1} \sin \left(n\pi + \frac{\pi}{2} \right) = \frac{2(-1)^n}{2n+1}$.

Par conséquent, pour tout entier $n \geq 1$:

$$\begin{aligned} |f(y_n) - f(x_n)| &= |f(x_{n-1}) - f(x_n)| = \frac{2}{2n-1} + \frac{2}{2n+1} \\ &= \frac{2(2n-1+2n+1)}{(2n+1)(2n-1)} = \frac{8n}{4n^2-1} > \frac{8n}{4n^2} = \frac{2}{n} \end{aligned}$$

donc :

$$\sum_{k=0}^p |f(y_{p+k}) - f(x_{p+k})| \geq \sum_{k=0}^p \frac{2}{p+k} \geq \frac{2(p+1)}{2p} > 1 = \varepsilon_0.$$

On conclut de cette étude que f n'est pas absolument continue sur $[0, 1]$. $\square$

8.17. Fonction absolument continue qui n'est pas lipschitzienne.

Nous avons vu dans l'exemple 8.15 que la fonction racine carrée est uniformément continue sur $\mathbb{R}_+$ mais qu'elle n'est pas lipschitzienne sur $\mathbb{R}_+$. Nous démontrons qu'en fait $\sqrt{\cdot}$ est absolument continue sur $\mathbb{R}_+$.

Soit ε un réel > 0 . Le réel $\eta = \varepsilon^2/8$ est strictement positif et $\sqrt{\eta} = \varepsilon/(2\sqrt{2})$. La croissance stricte de la fonction $\sqrt{\cdot}$ nous assure que si y est un nombre réel et si $0 \leq y < 2\eta$, alors $0 \leq \sqrt{y} < \sqrt{2}\sqrt{\eta} = \varepsilon/2$.

Soit n un entier ≥ 1 et $x_1, x_2, \dots, x_n, y_1, y_2, \dots, y_n$ des nombres réels tels que $0 \leq x_1 < y_1 \leq x_2 < y_2 \leq \dots \leq x_n < y_n$ et :

$$\sum_{i=1}^n (y_i - x_i) < \eta.$$

Nous posons $k = 0$ si $\eta \leq x_1$, $k = n$ si $x_n < \eta$ et, si $x_1 < \eta \leq x_n$, nous notons k l'unique élément de $\llbracket 1, n-1 \rrbracket$ tel que $x_k < \eta \leq x_{k+1}$.

On a $y_k - x_k \leq \sum_{i=1}^n (y_i - x_i) < \eta$ donc $y_k < 2\eta$, et $\sqrt{x_1} < \sqrt{y_1} \leq \dots \leq \sqrt{x_k} < \sqrt{y_k}$ donc :

$$\sum_{i=1}^k |\sqrt{y_i} - \sqrt{x_i}| = \sum_{i=1}^k (\sqrt{y_i} - \sqrt{x_i}) \leq \sqrt{y_k} - \sqrt{x_1} \leq \sqrt{y_k} < \frac{\varepsilon}{2}.$$

Nous avons vu dans l'exemple 8.15 que $\sqrt{\cdot}$ est lipschitzienne de rapport $\frac{1}{2\sqrt{\eta}} = \frac{\sqrt{2}}{\varepsilon}$

sur $[\eta, +\infty[$. On a, pour tout $i \in \llbracket k+1, n \rrbracket$, $y_i > x_i \geq \eta$ donc :

$$\begin{aligned} \sum_{i=k+1}^n |\sqrt{y_i} - \sqrt{x_i}| &\leq \sum_{i=k+1}^n \frac{\sqrt{2}}{\varepsilon} |y_i - x_i| = \frac{\sqrt{2}}{\varepsilon} \sum_{i=k+1}^n |y_i - x_i| \\ &\leq \frac{\sqrt{2}}{\varepsilon} \sum_{i=1}^n |y_i - x_i| = \frac{\sqrt{2}}{\varepsilon} \sum_{i=1}^n (y_i - x_i) < \frac{\sqrt{2}}{\varepsilon} \eta = \frac{\sqrt{2}}{\varepsilon} \frac{\varepsilon^2}{8} = \frac{\varepsilon}{4\sqrt{2}}, \end{aligned}$$

d'où finalement :

$$\sum_{i=1}^n |\sqrt{y_i} - \sqrt{x_i}| < \frac{\varepsilon}{2} + \frac{\varepsilon}{4\sqrt{2}} < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon.$$

En conclusion, $\sqrt{\cdot}$ est absolument continue sur $\mathbb{R}_+$. $\square$

La composée de deux applications continues est encore continue. Ce résultat reste valable si l'on remplace *continue* par *uniformément continue* ou *lipschitzienne*. Par contre il devient faux avec *absolument continue*.

8.18. Deux fonctions absolument continues dont la composée ne l'est pas.

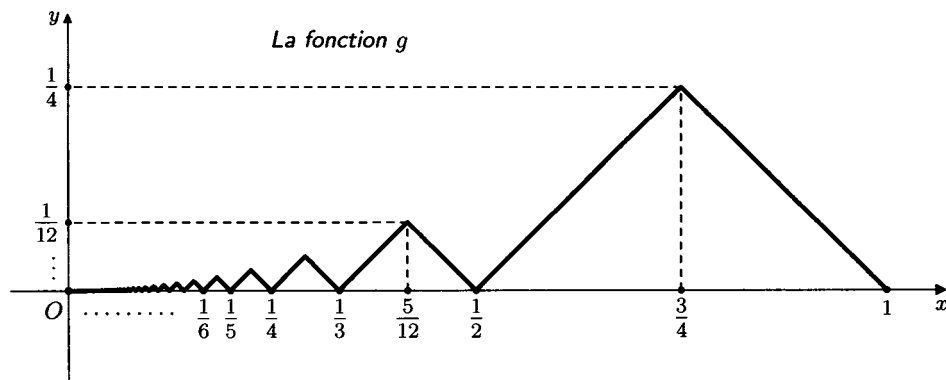
Il résulte de l'exemple 8.17 que la fonction $\sqrt{\cdot}$ est absolument continue sur $\mathbb{R}_+$, donc sur $[0, 1]$. Nous posons :

$$S = \{0\} \cup \left\{ \frac{1}{n} \mid n \in \mathbb{N}^* \right\} = \left\{ 0, 1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \frac{1}{5}, \frac{1}{6}, \dots \right\},$$

et nous introduisons l'application « distance à S » :

$$\left| \begin{array}{l} g : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = d(x, S) = \inf_{t \in S} |x - t| \end{array} \right.$$

représentée graphiquement ci-dessous.



Soit x et y des points du segment $[0, 1]$. On a, pour tout $t \in S$:

$$d(x, S) \leq |x - t| \leq |x - y| + |y - t|$$

donc $d(x, S) - |x - y| \leq |y - t|$. Par suite, $d(x, S) - |x - y|$ minore l'ensemble des $|y - t|$ pour t parcourant S , donc $d(x, S) - |x - y| \leq d(y, S)$, ce qui montre que $d(x, S) - d(y, S) \leq |x - y|$. En intervertissant les rôles de x et de y , on obtient

$d(y, S) - d(x, S) \leq |y - x| = |x - y|$. Il en résulte que $|d(x, S) - d(y, S)| \leq |x - y|$, ce qui s'écrit :

$$|g(x) - g(y)| \leq |x - y|.$$

Cette étude montre que g est lipschitzienne donc absolument continue sur $[0, 1]$.

La fonction g est à valeurs dans $\mathbb{R}_+$, ce qui justifie l'existence de la composée $f = \sqrt{} \circ g$, application de $[0, 1]$ dans $\mathbb{R}$. Nous prouvons que f n'est pas absolument continue sur $[0, 1]$. Pour ceci, nous introduisons les suites $(x_n)_{n \geq 1}$ et $(y_n)_{n \geq 1}$, de termes généraux :

$$x_n = \frac{1}{2} \left(\frac{1}{n} + \frac{1}{n+1} \right) \text{ et } y_n = \frac{1}{n},$$

à valeurs dans le segment $[0, 1]$. Remarquons que les x_n sont les points où g admet un maximum relatif. On a, pour tout entier $n \geq 1$:

$$g(x_n) = \frac{1}{2} \left(\frac{1}{n} - \frac{1}{n+1} \right) = \frac{1}{2n(n+1)}$$

et, comme y_n appartient à S , $g(y_n) = 0$.

Posons $\varepsilon_0 = 1$. Soit η un réel > 0 . La suite de terme général $(1/n)$ converge vers 0, ce qui justifie le choix d'un entier $N \geq 1$ tel que $1/n < \eta$ pour tout entier $n \geq N$. Soit p un entier $> N$. Considérons la famille $(x_p, x_{p-1}, \dots, x_N, y_p, y_{p-1}, \dots, y_N)$, qui vérifie les inégalités $x_p < y_p < x_{p-1} < y_{p-1} < \dots < x_N < y_N$. On a :

$$\sum_{k=N}^p (y_k - x_k) \leq y_N - x_p < y_N = \frac{1}{N} < \eta$$

et :

$$\sum_{k=N}^p |f(y_k) - f(x_k)| = \sum_{k=N}^p \frac{1}{\sqrt{2k(k+1)}} \geq \frac{1}{\sqrt{2}} \sum_{k=N}^p \frac{1}{k+1}.$$

La série harmonique étant divergente, il existe un entier $p > N$ tel que, pour tout entier $n \geq p$, $\sum_{k=N}^n 1/(k+1) > \sqrt{2}$. En choisissant un tel entier p et en définissant la famille $(x_p, x_{p-1}, \dots, x_N, y_p, y_{p-1}, \dots, y_N)$ comme ci-dessus, on a :

$$\sum_{k=N}^p |f(y_k) - f(x_k)| \geq 1 = \varepsilon_0.$$

Par suite la fonction composée $\sqrt{} \circ g$ n'est pas absolument continue sur $[0, 1]$. $\square$

Le produit de deux fonctions définies et continues sur une partie $\mathcal{A}$ de $\mathbb{R}$ est continue sur $\mathcal{A}$. Ceci devient faux si l'on remplace *continue* par *uniformément continue* ; pour obtenir ceci, il faut bien sûr que $\mathcal{A}$ ne soit pas un compact.

8.19. Produit de deux fonctions lipschitziennes qui n'est pas uniformément continue.

La fonction $f : x \mapsto f(x) = x$ est définie et lipschitzienne de rapport 1 sur $\mathbb{R}$, alors que la fonction produit $g = f \times f : x \mapsto g(x) = x^2$ n'est pas uniformément continue sur $\mathbb{R}$ comme le montre l'exemple 8.14, page 145. $\square$

DÉFINITION 8.11. — Si $\mathcal{A}$ est une partie de $\mathbb{R}$ et f une fonction définie sur $\mathcal{A}$, f est contractante sur $\mathcal{A}$ si f est lipschitzienne de rapport $k < 1$ sur $\mathcal{A}$.

Si $\mathcal{A}$ est un fermé de $\mathbb{R}$, c'est un espace métrique complet, ce qui permet de prouver que si f est une application contractante de $\mathcal{A}$ dans $\mathcal{A}$, l'équation $f(x) = x$ admet dans $\mathcal{A}$ une solution unique. Ceci devient faux si l'on suppose seulement que $|f(x) - f(y)| < |x - y|$ quels que soient les points distincts x et y de $\mathcal{A}$.

8.20. Application de $\mathbb{R}$ dans $\mathbb{R}$ qui réduit strictement les distances mais qui n'admet pas de point fixe.

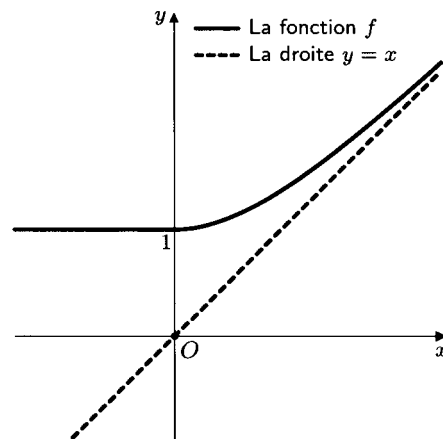
Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} x + e^{-x} & \text{si } x \geq 0, \\ 1 & \text{si } x < 0. \end{cases} \end{array} \right.$$

Comme $f(0) = 1$, l'application f est continue sur $\mathbb{R}$. La fonction f est dérivable sur $]-\infty, 0[$ et sur $]0, +\infty[$, $f'(x) = 1 - e^{-x}$ pour tout réel $x > 0$ et $f'(x) = 0$ pour tout réel $x < 0$; par suite, $f'(x)$ admet pour limite 0 quand x tend vers 0, donc f est dérivable en 0 et $f'(0) = 0$. De plus $0 \leq f'(x) < 1$ pour tout réel x .

Si x et y sont des réels et si $x < y$, le théorème des accroissements finis appliqué à f sur $[x, y]$ nous affirme l'existence d'un point z de $]x, y[$ tel que $f(y) - f(x) = (y - x)f'(z)$, donc $|f(x) - f(y)| = |y - x|f'(z) < |x - y|$.

Si $x \geq 0$, $f(x) = x + e^{-x} > x$ et, si $x < 0$, $f(x) = 1 > x$. Ceci montre que l'équation $f(x) = x$ n'a pas de solution dans $\mathbb{R}$, bien que $\mathbb{R}$ soit un fermé de $\mathbb{R}$. $\square$



■ Continuité de l'application réciproque d'une fonction continue

THÉORÈME 8.4. — Si f est une fonction définie, continue et injective sur un intervalle I d'intérieur non vide, $J = f(I)$ est un intervalle d'intérieur non vide, $f|_I$ est une bijection de I sur J et $g = (f|_I)^{-1}$ est continue sur J .

Ce théorème appelle deux remarques : d'une part, il impose que l'ensemble de départ soit un intervalle et d'autre part, il traite d'un problème global, c'est-à-dire qu'il impose la continuité sur tout l'intervalle de départ. Nous prouvons dans les exemples suivants que ces conditions sont nécessaires à l'établissement d'un théorème sur la continuité de l'application réciproque.

8.21. Bijection continue f d'une partie $\mathcal{A}$ et $\mathbb{R}$ sur une partie $\mathcal{B}$ de $\mathbb{R}$ pour laquelle f^{-1} n'est continue en aucun point de $\mathcal{B}$.

Les ensembles $\mathbb{Z}$ et $\mathbb{Q}$ sont tous deux dénombrables, donc ils sont équipotents. Nous choisissons une bijection f de $\mathbb{Z}$ sur $\mathbb{Q}$.

Soit a un entier relatif. Soit ε un réel > 0 . Le réel $\eta = 1$ est strictement positif et, pour tout point x de $\mathbb{Z}$, $|x - a| < \eta$ entraîne $x = a$ donc $|f(x) - f(a)| = 0 < \varepsilon$. Il en résulte que f est continue en a .

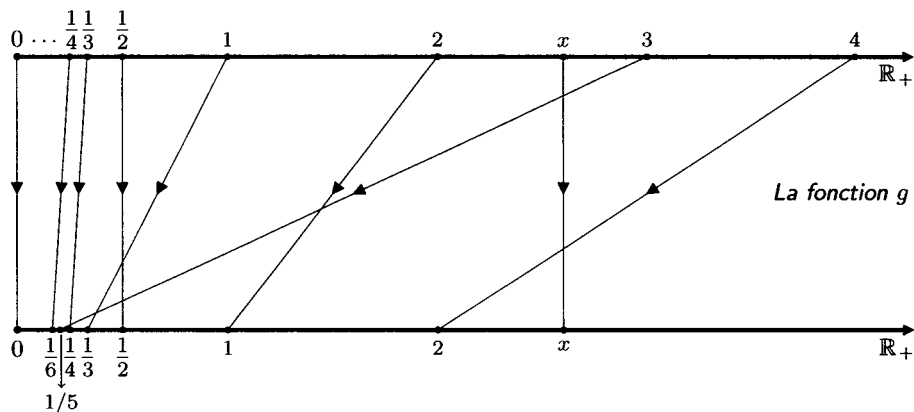
Soit b un nombre rationnel. Nous posons $\varepsilon_0 = 1$. Soit η un réel > 0 . Nous choisissons un rationnel c tel que $b < c < b + \eta$. Alors $c \in \mathbb{Q}$, $|c - b| < \eta$ et, comme l'application réciproque f^{-1} de f est injective et à valeurs dans $\mathbb{Z}$, $f^{-1}(c)$ et $f^{-1}(b)$ sont des entiers relatifs distincts donc $|f^{-1}(c) - f^{-1}(b)| \geq 1 = \varepsilon_0$. Par conséquent f^{-1} est discontinue en b .

En conclusion, f est une bijection de $\mathbb{Z}$ sur $\mathbb{Q}$, continue en tous les points de $\mathbb{Z}$ et dont l'application réciproque est discontinue en tous les points de $\mathbb{Q}$. $\square$

8.22. Bijection f de $\mathbb{R}$ sur $\mathbb{R}$, continue en 0, mais dont l'application réciproque est discontinue en $f(0)$.

Nous posons $S = \left\{ \frac{1}{n} \mid n \in \mathbb{N} \text{ et } n \geq 2 \right\}$ et nous considérons les applications :

$$g : \mathbb{R}_+ \longrightarrow \mathbb{R}_+ \quad \left| \quad \begin{aligned} x \longmapsto g(x) = \begin{cases} \frac{n}{2} & \text{si } x = n \text{ où } n \text{ est un entier naturel pair,} \\ \frac{1}{n+2} & \text{si } x = n \text{ où } n \text{ est un entier naturel impair,} \\ \frac{1}{2(n-1)} & \text{si } x = \frac{1}{n} \text{ où } n \text{ est un entier et } n \geq 2, \\ x & \text{dans tous les autres cas} \end{cases} \end{aligned} \right.$$



et :

$$f : \mathbb{R} \longrightarrow \mathbb{R} \quad \left| \quad \begin{aligned} x \longmapsto f(x) = \begin{cases} g(x) & \text{si } x \geq 0, \\ -g(-x) & \text{si } x < 0. \end{cases} \end{aligned} \right.$$

Soit y un point de $\mathbb{R}_+$. Si $y \notin S \cup \mathbb{N}^*$, y est l'unique antécédent de y par g ; si $y = n$ où $n \in \mathbb{N}^*$, $2n$ est l'unique antécédent de y par g ; si $y \in S$, $y = 1/n$ où $n \in \mathbb{N}$ et $n \geq 2$, et alors, pour n pair, $1/(1 + (n/2))$ est l'unique antécédent de y par g et pour n impair, $n - 2$ est l'unique antécédent de y par g . Il en résulte que g est une bijection de $\mathbb{R}_+$ sur $\mathbb{R}_+$. Comme g est la restriction de f à $\mathbb{R}_+$ et que la fonction f est impaire, on en déduit que f est une bijection de $\mathbb{R}$ sur $\mathbb{R}$.

Nous remarquons que $|f(x)| \leq |x|$ pour tout réel x ; il en résulte que f est continue en 0. Or $f(2n - 1) = 1/(2n + 1)$ pour tout entier $n \geq 1$, la suite de terme général $1/(2n + 1)$ tend vers 0 et la suite de terme général $f^{-1}(1/(2n + 1)) = 2n - 1$ ne converge pas vers 0, ce qui montre que f^{-1} n'est pas continue en $0 = f(0)$. $\square$

■ Continuité et topologie

Nous rappelons que si $\mathcal{A}$ est une partie de $\mathbb{R}$, son adhérence se note $\overline{\mathcal{A}}$.

On peut caractériser la continuité d'une application f de $\mathbb{R}$ dans $\mathbb{R}$ de la manière suivante : pour toute partie $\mathcal{A}$ de $\mathbb{R}$, $f(\overline{\mathcal{A}}) \subset \overline{f(\mathcal{A})}$.

8.23. Application continue f de $\mathbb{R}$ dans $\mathbb{R}$ et partie $\mathcal{A}$ de $\mathbb{R}$ telles que $f(\overline{\mathcal{A}}) \neq \overline{f(\mathcal{A})}$.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \frac{x}{1+|x|} \end{array} \right.$$

La fonction f est continue et strictement croissante sur l'intervalle $\mathcal{A} = \mathbb{R}$ et f admet -1 et 1 pour limites respectives en $-\infty$ et $+\infty$, donc $f(\overline{\mathcal{A}}) = f(\mathbb{R}) =]-1, 1[$, alors que l'adhérence de $f(\mathcal{A})$ est $\overline{f(\mathcal{A})} = \overline{f(\mathbb{R})} = [-1, 1]$. $\square$

DÉFINITION 8.12. — Le graphe d'une application f d'une partie $\mathcal{A}$ de $\mathbb{R}$ dans $\mathbb{R}$ est le sous-ensemble $\mathcal{G}(f) = \{(x, f(x)) \mid x \in \mathcal{A}\}$ de $\mathbb{R}^2$.

Si $\mathcal{A}$ est un fermé de $\mathbb{R}$ et f une application continue de $\mathcal{A}$ dans $\mathbb{R}$, le graphe de f est un fermé de $\mathbb{R}^2$ — muni de sa topologie canonique. Ceci ne caractérise pas les applications continues, comme nous le montre l'exemple suivant.

8.24. Application discontinue dont le graphe est fermé.

Nous introduisons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} \frac{1}{x} & \text{si } x > 0, \\ 0 & \text{si } x \leq 0. \end{cases} \end{array} \right.$$

Comme $f(x)$ tend vers $+\infty$ quand x tend vers 0 à droite, f est discontinue en 0 , donc f n'est pas une application continue de $\mathbb{R}$ dans $\mathbb{R}$.

Le graphe de f est la réunion des ensembles :

$$\mathcal{L} = \{(x, 0) \mid x \leq 0\} \text{ et } \mathcal{M} = \{(x, 1/x) \mid x > 0\}.$$

Les projections $p : z = (x, y) \mapsto p(z) = x$ et $q : z = (x, y) \mapsto q(z) = y$ sont des applications continues de $\mathbb{R}^2$ dans $\mathbb{R}$, $\mathbb{R}_-$ et $\{0\}$ sont des fermés de $\mathbb{R}$ et $\mathcal{L} = p^{-1}(\mathbb{R}_-) \cap q^{-1}(\{0\})$, donc $\mathcal{L}$ est un fermé de $\mathbb{R}^2$ comme intersection de deux fermés. De plus $g : z = (x, y) \mapsto g(z) = xy$ est une application continue de $\mathbb{R}^2$ dans $\mathbb{R}$, $\mathbb{R}_+$ et $\{1\}$ sont des fermés de $\mathbb{R}$ et :

$$\mathcal{M} = \{(x, y) \mid x \geq 0 \text{ et } xy = 1\} = p^{-1}(\mathbb{R}_+) \cap g^{-1}(\{1\})$$

donc $\mathcal{M}$ est un fermé de $\mathbb{R}^2$. Le graphe de f est donc un fermé de $\mathbb{R}^2$. $\square$

DÉFINITION 8.13. — Si E et F sont des espaces topologiques¹³ et f une application de E dans F , f est ouverte (resp. fermée) si l'image directe de tout ouvert (resp. de tout fermé) de E par f est un ouvert (resp. un fermé) de F .

¹³. Voir le chapitre 15.

Nous démontrons que, dans l'espace topologique $\mathbb{R}$, les trois notions d'application ouverte, fermée et continue sont distinctes.

8.25. Application ouverte qui n'est pas fermée.

La fonction arc tangente, de symbole $\arctan$, est l'application réciproque de la restriction T de la fonction tangente à l'intervalle ouvert $I =]-\pi/2, \pi/2[$. La fonction T est une bijection strictement croissante et continue de I sur $\mathbb{R}$ et la fonction $\arctan$ une bijection strictement croissante et continue de $\mathbb{R}$ sur I .

L'image réciproque par T de tout ouvert de $\mathbb{R}$ est un ouvert de I , muni de la topologie induite par la topologie canonique de $\mathbb{R}$, et comme I est un ouvert de $\mathbb{R}$, l'image réciproque par T de tout ouvert de $\mathbb{R}$ est un ouvert de $\mathbb{R}$, donc l'image directe par $\arctan$ de tout ouvert de $\mathbb{R}$ est un ouvert de $\mathbb{R}$. Il en résulte que $\arctan$ est une application ouverte de $\mathbb{R}$ dans $\mathbb{R}$.

Par ailleurs $\mathbb{R}$ est un fermé de $\mathbb{R}$, mais son image directe par $\arctan$ est $]-\pi/2, \pi/2[$, qui n'est pas un fermé de $\mathbb{R}$. Par conséquent $\arctan$ n'est pas une application fermée de $\mathbb{R}$ dans $\mathbb{R}$. $\square$

8.26. Application fermée qui n'est pas ouverte.

Nous considérons la fonction « valeur absolue » :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = |x|. \end{array} \right.$$

Soit $\mathcal{F}$ un fermé de $\mathbb{R}$. Les ensembles $\mathcal{F}_1 = \mathcal{F} \cap [0, +\infty[$ et $\mathcal{F}_2 = \mathcal{F} \cap]-\infty, 0]$ sont des fermés comme intersections de fermés. L'application $\sigma : x \mapsto \sigma(x) = -x$ est une bijection continue de $\mathbb{R}$ sur $\mathbb{R}$ et $\sigma^{-1} = \sigma$, donc l'image par σ de tout fermé de $\mathbb{R}$ est un fermé de $\mathbb{R}$. Par suite, $\mathcal{G}_2 = \{-x \mid x \in \mathcal{F}_2\}$ est un fermé de $\mathbb{R}$; or $f(\mathcal{F}) = \mathcal{F}_1 \cup \mathcal{G}_2$ donc $f(\mathcal{F})$ est un fermé de $\mathbb{R}$. Il en résulte que f est une application fermée de $\mathbb{R}$ dans $\mathbb{R}$.

Cependant l'image directe de l'intervalle ouvert $]-1, 1[$ par f est l'intervalle $[0, 1[$, qui n'est pas un ouvert, donc f n'est pas une application ouverte de $\mathbb{R}$ dans $\mathbb{R}$. $\square$

8.27. Application continue qui n'est ni ouverte, ni fermée.

Nous introduisons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} \sin x & \text{si } x \leq 0, \\ \arctan x & \text{si } x > 0. \end{cases} \end{array} \right.$$

On a $\arctan 0 = 0 = f(0)$, donc $f(x) = \arctan x$ pour tout $x \geq 0$. Par suite, f est continue sur les intervalles fermés $]-\infty, 0]$ et $[0, +\infty[$, donc f est une application continue de $\mathbb{R}$ dans $\mathbb{R}$.

Le segment $[-2\pi, 0]$ étant inclus dans $]-\infty, 0]$, on a $f([-2\pi, 0]) = [-1, 1]$. D'autre part, $f([0, +\infty[) = \arctan([0, +\infty[) = [0, \pi/2[$, donc $f(\mathbb{R}) = [-1, \pi/2[$, ce qui montre que $f(\mathbb{R})$ n'est ni un ouvert, ni un fermé de $\mathbb{R}$, alors que $\mathbb{R}$ est à la fois ouvert et fermé. $\square$

8.28. Application fermée et discontinue.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x < 0, \\ 1 & \text{si } x \geq 0. \end{cases} \end{array} \right.$$

Comme $f(x)$ admet pour limite 0 quand x tend vers 0 à gauche et pour limite 1 quand x tend vers 0 à droite, la fonction f est discontinue en 0. Si $\mathcal{F}$ est un fermé de $\mathbb{R}$, son image directe $f(\mathcal{F})$ est incluse dans $\{0, 1\}$, donc finie, ce qui montre que $f(\mathcal{F})$ est un fermé de $\mathbb{R}$. L'application f est donc fermée et discontinue. $\square$

8.29. Application ouverte et fermée, mais discontinue.

Nous munissons l'intervalle $I = [0, 2\pi[$ de la topologie induite par la topologie canonique de $\mathbb{R}$ et l'ensemble $\mathbb{U} = \{z \in \mathbb{C} \mid |z| = 1\}$ des nombres complexes de module 1 de la topologie induite par la topologie canonique de $\mathbb{C}$ et nous introduisons l'application :

$$\left| \begin{array}{l} f : I \longrightarrow \mathbb{U} \\ x \longmapsto f(x) = e^{ix}. \end{array} \right.$$

L'application f est une bijection continue de I sur $\mathbb{U}$, donc l'image réciproque de tout ouvert (resp. de tout fermé) de $\mathbb{U}$ par f est un ouvert (resp. un fermé) de I . Ceci montre que l'application réciproque $g = f^{-1}$, bijection de $\mathbb{U}$ sur I , est une application ouverte et fermée de $\mathbb{U}$ dans I . Cependant, g est discontinue en 1; en effet, $g(1) = 0$ alors que tout voisinage de 1 dans $\mathbb{U}$ contient des nombres complexes de module 1 et de partie imaginaire strictement négative, donc d'argument appartenant à l'intervalle $] \pi, 2\pi[$. Par suite, g n'est pas une application continue de $\mathbb{U}$ dans I . $\square$

Dans l'exemple précédent 8.29, les ouverts et fermés le sont pour la topologie induite. Par exemple $g(\mathbb{U}) = I$ est un ouvert de I mais ce n'est pas un ouvert de $\mathbb{R}$.

8.30. Application fermée discontinue en tout point.

Nous reprenons la fonction de Dirichlet de l'exemple 8.1 (page 134) :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 1 & \text{si } x \text{ est rationnel,} \\ 0 & \text{si } x \text{ est irrationnel.} \end{cases} \end{array} \right.$$

Nous avons vu dans cet exemple que f n'est continue en aucun point de $\mathbb{R}$. Comme dans l'exemple 8.28, l'image directe par f d'une partie quelconque, en particulier celle d'un fermé, est finie donc fermée, ce qui montre que f est une application fermée de $\mathbb{R}$ dans $\mathbb{R}$. $\square$

8.31. Application ouverte discontinue en tout point.

Nous nous inspirons de la fonction définie à l'exemple 8.6 (pages 138 et 139) et nous construisons une application de période 1 qui prend toutes les valeurs réelles sur n'importe quel intervalle ouvert. Nous notons E la fonction « partie entière ». Soit x un nombre réel. Le nombre réel $x - E(x)$ appartient à l'intervalle $[0, 1[$. Nous notons $0, a_1 a_2 a_3 a_4 \dots$ le développement décimal illimité propre de $x - E(x)$

et nous considérons la suite $(a_{3i-2})_{i \in \mathbb{N}^*} = (a_1, a_4, a_7, \dots, a_{3i-2}, \dots)$ des décimales de $x - E(x)$ de rang congru à 1 modulo 3. Si cette suite n'est pas périodique à partir d'un certain rang, nous posons $\psi(x) = 0$ et, si elle est périodique à partir de a_{3n-2} — et pas avant — et de période $p \in \mathbb{N}^*$, nous notons $\psi(x)$ le réel d'écriture décimale $\varepsilon a_{3(n+p-1)-1} \dots a_{3n+2} a_{3n-1}, a_{3n} a_{3n+3} a_{3n+6} \dots$ où $\varepsilon = +1$ si a_{3n-2} est pair et $\varepsilon = -1$ si a_{3n-2} est impair ; remarquons que nous utilisons les décimales de rang congrus à 2 modulo 3 pour la partie entière et les décimales de rang divisible par 3 pour la partie décimale. Nous avons construit ainsi une application $\psi : x \mapsto \psi(x)$ de $\mathbb{R}$ dans $\mathbb{R}$, de période 1 puisque la définition de $\psi(x)$ ne dépend que des décimales de x . On prouve comme dans l'exemple 8.6 (pages 138 et 139) que ψ n'est continue en aucun point de $[0, 1[$, donc en aucun point de $\mathbb{R}$.

Soit a et b des réels tels que $0 \leq a < b \leq 1$. Soit y un nombre réel. Nous cherchons un point x de $]a, b[$ tel que $y = \psi(x)$. Nous choisissons un entier $n \geq 2$ tel que :

$$\frac{1}{10^{3n-3}} < \frac{b-a}{2}.$$

La propriété d'Archimède nous fournit $q \in \mathbb{N}$ tel que $a \leq \frac{q}{10^{3n-3}} < \frac{q+1}{10^{3n-3}} < b$.

Comme $b \leq 1$, on a $q < 10^{3n-3}$ donc l'écriture décimale de q s'écrit $e_1 e_2 \dots e_m$ où $m \leq 3n-3$, d'où, en posant $u_k = 0$ pour tout $k \in \llbracket 1, 3n-3-m \rrbracket$ et $u_{3n-3-m+i} = e_i$ pour tout $i \in \llbracket 1, m \rrbracket$, le développement décimal :

$$\frac{q}{10^{3n-3}} = 0, u_1 u_2 \dots u_{3n-3}.$$

Nous traitons d'abord le cas où $y = 0$. Nous choisissons une suite non périodique $(v_i)_{i \geq 1}$ de chiffres — par exemple la suite des décimales de e — et nous posons :

$$x = 0, a_1 a_2 a_3 \dots = 0, u_1 \dots u_{3n-3} v_1 00 v_2 00 \dots v_{i-1} 00 v_i 00 \dots$$

Les v_i n'étant pas tous nuls, le réel x appartient à $]a, b[$, et on a $\psi(x) = 0 = y$. Supposons que $y \neq 0$. Nous explicitons le développement décimal du nombre réel y , qui s'écrit $y = \varepsilon b_p b_{p-1} \dots b_1, c_1 c_2 c_3 \dots$ où $\varepsilon = \pm 1$, et nous posons :

$$x = 0, a_1 a_2 a_3 \dots = 0, u_1 \dots u_{3n-3} v_1 b_1 c_1 v_2 b_2 c_2 \dots v_p b_p c_p v_1 b_1 c_{p+1} v_2 b_2 \dots$$

où la suite finie $(v_1, v_2, \dots, v_p)$ de chiffres possède les propriétés suivantes : si $p \geq 2$, elle n'est pas périodique, v_1 est pair si $\varepsilon = +1$ et impair si $\varepsilon = -1$, et la suite :

$$(a_1, a_4, \dots, a_{3q+1}, \dots) = (u_1, u_4, \dots, u_{3n-5}, v_1, v_2, \dots, v_p, \dots, v_1, v_2, \dots, v_p, \dots)$$

des décimales de rang congru à 1 modulo 3 n'est périodique de période p qu'à partir de $a_{3n-2} = v_1$. Dans la suite $(b_1, \dots, b_p, c_1, c_2, c_3, \dots)$, il y a au moins un chiffre différent de zéro, donc $a < x < b$, et la définition de ψ montre que $\psi(x) = y$. Ainsi ψ prend toutes les valeurs réelles sur tout intervalle ouvert $]a, b[$ où a et b sont des réels tels que $0 \leq a < b \leq 1$.

Soit a et b des réels tels que $a < b$. On note q la partie entière de a . Alors $q \in \mathbb{Z}$ et $q \leq a < q+1$. On pose $c = \text{Min}(b, q+1)$. Comme $1 \leq a - q < c - q \leq 1$, ψ prend toutes les valeurs réelles sur $]a - q, c - q[$ et, q étant une période de ψ , ψ prend toutes les valeurs réelles sur $]a, c[$, donc sur $]a, b[$, ce qui montre que $\psi([a, b]) = \mathbb{R}$.

Tout ouvert non vide O de $\mathbb{R}$ contient un intervalle ouvert $]a, b[$ où a et b sont des réels tels que $a < b$, donc $\psi(O) = \mathbb{R}$. De plus $\mathbb{R}$ est un ouvert de $\mathbb{R}$, donc l'image directe par ψ de tout ouvert de $\mathbb{R}$ est un ouvert de $\mathbb{R}$, ce qui montre que ψ est une application ouverte de $\mathbb{R}$ dans $\mathbb{R}$. $\square$

Chapitre 9

Fonctions d'une variable réelle Dérivabilité

La notion de dérivée apparaît à la fin du XVII^e siècle, en même temps que le calcul intégral, sous l'impulsion de Newton et de Leibniz. Avant eux, Descartes s'était intéressé au problème des tangentes à une courbe et Pierre de Fermat avait introduit une notion très proche de la dérivée lors d'une recherche d'extremums. Dans la deuxième moitié du XIX^e siècle, on s'intéresse aux propriétés des fonctions dérivées. C'est alors que sont introduits de nombreux contre-exemples qui défient l'intuition généralement admise.

Dans tout le chapitre, les fonctions sont des fonctions de $\mathbb{R}$ dans $\mathbb{R}$.

■ Dérivabilité locale et globale

Nous associons à toute fonction f et à tout point a du domaine de f la fonction quotient :

$$Qf(a, \bullet) : x \mapsto Qf(a, x) = \frac{f(x) - f(a)}{x - a}.$$

DÉFINITION 9.1. — Si une fonction f est définie sur un voisinage d'un nombre réel a , f est dérivable en a si $Qf(a, x)$ admet une limite dans $\mathbb{R}$ quand x tend vers a et, si la fonction f est dérivable en a , la dérivée de f en a est le nombre réel :

$$f'(a) = \lim_{x \rightarrow a} Qf(a, x) = \lim_{x \rightarrow a} \frac{f(x) - f(a)}{x - a}.$$

DÉFINITION 9.2. — Si a est un nombre réel et f une fonction définie sur un intervalle $[a, a + r[$ (resp. $]a - r, a]$) où r est un réel strictement positif, f est dérivable à droite (resp. à gauche) en a si $Qf(a, x)$ admet une limite dans $\mathbb{R}$ quand x tend vers a à droite (resp. à gauche) et, si f est dérivable à droite (resp. à gauche) en a , la dérivée à droite (resp. à gauche) de f en a est le nombre réel :

$$f'_d(a) = \lim_{x \rightarrow a^+} Qf(a, x) \quad (\text{resp. } f'_g(a) = \lim_{x \rightarrow a^-} Qf(a, x)).$$

Si Ω est un ouvert de $\mathbb{R}$ et f une fonction définie sur Ω , f est dérivable sur Ω si f est dérivable en tout point de Ω , et si f est dérivable sur Ω , la fonction dérivée de f sur Ω est l'application $f' : x \mapsto f'(x)$ de Ω dans $\mathbb{R}$. Si f est une fonction définie sur un intervalle I d'intérieur non vide et si I n'est pas un intervalle ouvert, f est dérivable sur I si f est dérivable en tout point de l'intérieur de I , dérivable à droite en $\alpha = \inf I$ si $\alpha \in I$ et dérivable à gauche en $\beta = \sup I$ dans le cas où $\beta \in I$, et si f est dérivable sur I , la fonction dérivée de f sur I est l'application $f' : x \mapsto f'(x)$ de I dans $\mathbb{R}$ où $f'(\alpha) = f'_d(\alpha)$ si $\alpha = \inf I$ appartient à I et $f'(\beta) = f'_g(\beta)$ si $\beta = \sup I$ appartient à I . Avec ces hypothèses, la fonction f est dite de classe $\mathcal{C}^1$ sur Ω (sur I) si f est dérivable sur Ω (sur I) et f' continue sur Ω (sur I).

La notion de dérivabilité en un point a est un problème local ; elle est liée au comportement de f au voisinage de a , et on peut se demander si elle entraîne des propriétés pour f sur tout un voisinage de a .

9.1. Fonction dérivable en 0 mais discontinue en tout réel $x \neq 0$.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} x^2 & \text{si } x \text{ est rationnel,} \\ 0 & \text{si } x \text{ est irrationnel.} \end{cases} \end{array} \right.$$

Pour tout nombre réel x , $|f(x)| \leq x^2$, donc, pour tout réel $x \neq 0$:

$$|Qf(0, x)| = \left| \frac{f(x)}{x} \right| \leq |x|.$$

Par conséquent $Qf(0, \bullet)$ admet pour limite 0 en 0, donc f est dérivable en 0 et $f'(0) = 0$. Soit a un nombre réel différent de 0. Notons $(r_n)_{n \geq 0}$ la suite des valeurs décimales approchées par défaut de a et introduisons la suite $(\alpha_n)_{n \geq 0}$, de terme général $\alpha_n = r_n - (e/10^n)$; ces deux suites convergent vers a . Pour tout entier naturel n , r_n est rationnel et α_n irrationnel, donc $f(r_n) = (r_n)^2$ et $f(\alpha_n) = 0$, ce qui montre que la suite $(f(r_n))$ converge vers $a^2 > 0$ alors que la suite $(f(\alpha_n))$ converge vers 0. Il en résulte que la fonction f est discontinue en a . $\square$

Interrogeons-nous maintenant sur un problème similaire. Une fonction peut-elle être dérivable en tous les nombres réels sauf un ? Si l'on n'impose pas la continuité de f sur $\mathbb{R}$, le problème est très simple.

9.2. Fonction continue sur $\mathbb{R}$ et dérivable en tout nombre réel non nul mais pas en 0.

Nous utilisons la fonction « valeur absolue » $f : x \mapsto f(x) = |x|$.

Si a est un réel et si $a \neq 0$, alors $r = |a| > 0$, et $Qf(a, x) = (x - a)/(x - a) = 1$ pour tout $x \in]a - r, a + r[$ si $a > 0$ et $Qf(a, x) = (-x + a)/(x - a) = -1$ pour tout point x de $]a - r, a + r[$ si $a < 0$, donc f est dérivable en a , et $f'(a) = 1$ si $a > 0$ et $f'(a) = -1$ si $a < 0$.

De plus $Qf(0, x)$ vaut $x/x = 1$ pour tout $x > 0$ et $(-x)/x = -1$ pour tout $x < 0$, donc f est dérivable à droite et à gauche en 0 et $f'_d(0) = 1 \neq -1 = f'_g(0)$, ce qui montre que f n'est pas dérivable en 0. $\square$

Si l'on n'autorise pas l'existence de dérivées à droite et à gauche au point où la fonction n'est pas dérivable, le problème est plus délicat.

9.3. Fonction continue sur $\mathbb{R}$ qui est dérivable en tout réel non nul mais qui n'est dérivable ni à droite ni à gauche en 0.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} x \sin \frac{1}{x} & \text{si } x \neq 0, \\ 0 & \text{si } x = 0. \end{cases} \end{array} \right.$$

Pour tout réel $x \neq 0$, $|f(x)| = |x| |\sin(1/x)| \leq |x|$, inégalité vraie aussi pour $x = 0$, donc f admet pour limite $0 = f(0)$ en 0, ce qui prouve que f est continue en 0. De plus, si a est un réel et si $a \neq 0$, $x \mapsto (1/x)$ est dérivable en a , donc f est dérivable en a et par suite continue en a . La fonction f est donc continue sur $\mathbb{R}$ et dérivable en tout nombre réel non nul. Le calcul donne, pour tout réel $x \neq 0$:

$$f'(x) = \sin \frac{1}{x} - \frac{1}{x} \cos \frac{1}{x}.$$

$$\text{On a, pour tout réel } x \neq 0, \quad Qf(0, x) = \frac{f(x) - f(0)}{x - 0} = \frac{f(x)}{x} = \sin \frac{1}{x}.$$

La suite $(a_n)_{n \geq 0}$, de terme général $a_n = 1/((\pi/2) + n\pi)$, est à valeurs dans $]0, +\infty[$ et converge vers 0, et la suite $(b_n)_{n \geq 0}$, de terme général $b_n = -a_n$, est à valeurs dans $] -\infty, 0[$ et converge vers 0. Or, pour tout $n \in \mathbb{N}$, $Qf(0, a_n) = (-1)^n$ et $Qf(0, b_n) = (-1)^{n+1}$, termes généraux de suites divergentes, donc f n'est dérivable ni à droite ni à gauche en 0. $\square$

Remarquons que dans l'exemple précédent 9.3, la fonction f' n'est pas bornée au voisinage de 0, ni à droite ni à gauche. En effet, la suite $(c_n)_{n \geq 0}$, de terme général $c_n = 1/(2n\pi)$, est à valeurs dans $]0, +\infty[$ et converge vers 0, et la suite $(d_n)_{n \geq 0}$, de terme général $d_n = -c_n$, est à valeurs dans $] -\infty, 0[$ et converge vers 0, alors que, pour tout $n \in \mathbb{N}$, $f'(c_n) = -2n\pi$ et $f'(d_n) = 2n\pi$, ce qui montre que la suite $(f'(c_n))$ tend vers $-\infty$ et la suite $(f'(d_n))$ vers $+\infty$.

Voici un exemple du même type, mais dans lequel la fonction dérivée est bornée.

9.4. Fonction continue sur $\mathbb{R}$, dérivable en tout nombre réel non nul et dont la fonction dérivée est bornée, mais qui n'est pas dérivable à droite ni à gauche en 0.

Si n est un entier naturel non nul, nous posons, pour tout point x de $\left[\frac{1}{4^n}, \frac{1}{4^{n-1}}\right[$:

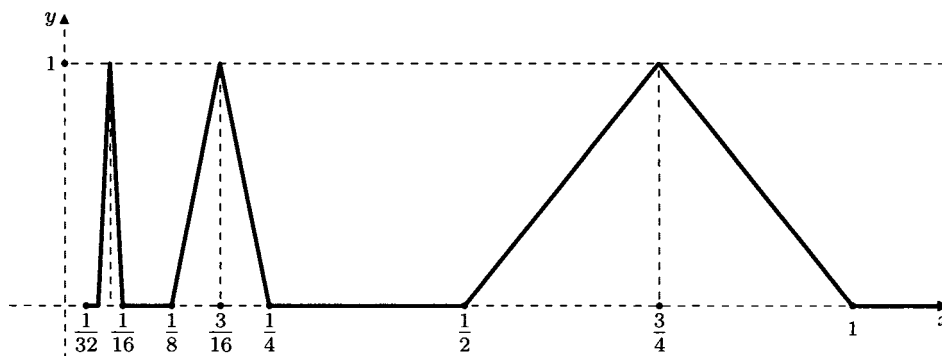
$$\varphi(x) = \begin{cases} 0 & \text{si } \frac{1}{4^n} \leq x < \frac{2}{4^n}, \\ 4^n x - 2 & \text{si } \frac{2}{4^n} \leq x < \frac{3}{4^n}, \\ -4^n x + 4 & \text{si } \frac{3}{4^n} \leq x < \frac{4}{4^n}. \end{cases}$$

La famille d'intervalles :

$$\left(\left[\frac{1}{4^n}, \frac{1}{4^{n-1}} \right[\right)_{n \in \mathbb{N}^*}$$

étant une partition de l'intervalle $]0, 1[$, ceci définit une application $\varphi : x \mapsto \varphi(x)$ de $]0, 1[$ dans $\mathbb{R}$. Nous construisons alors l'application g de $]0, +\infty[$ dans $\mathbb{R}$ en posant $g(x) = \varphi(x)$ si $x \in]0, 1[$ et $g(x) = 0$ si $x \geq 1$, et nous prolongeons enfin g en une application de $\mathbb{R}$ dans $\mathbb{R}$ en posant $g(0) = 0$ et $g(x) = g(-x)$ pour tout réel $x < 0$.

Nous représentons ci-dessous la restriction de g à l'intervalle $\left[\frac{1}{32}, +\infty\right]$, ce qui permet d'imaginer la représentation de sa restriction à $\mathbb{R}_+$.



La fonction g est clairement continue sur $]0, +\infty[$, donc aussi sur $]-\infty, 0[$, et, pour tout entier $n \geq 1$:

$$A_n = \int_{\frac{1}{4^n}}^{\frac{1}{4^{n-1}}} g(x) dx = \int_{\frac{2}{4^n}}^{\frac{1}{4^{n-1}}} g(x) dx = \frac{1}{4^n}$$

— c'est l'aire d'un triangle isocèle de base $2/4^n$ et de hauteur 1. La relation de Chasles donne, pour tout entier $n \geq 1$:

$$\int_{\frac{1}{4^n}}^1 g(x) dx = \sum_{k=1}^n A_k = \sum_{k=1}^n \frac{1}{4^k} = \frac{1}{4} \frac{1 - \frac{1}{4^n}}{1 - \frac{1}{4}} = \frac{1}{3} \left(1 - \frac{1}{4^n}\right).$$

La fonction g étant à valeurs dans $\mathbb{R}_+$, g est intégrable sur $[0, 1]$ et $\int_0^1 g(x) dx = \frac{1}{3}$.

La fonction g est, pour tout réel $x \geq 0$, intégrable sur $[0, x]$ et, comme elle est paire, elle est intégrable sur $[x, 0]$ pour tout réel $x \leq 0$, ce qui justifie l'introduction de l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \int_0^x g(t) dt. \end{array} \right.$$

La continuité de g sur $]0, +\infty[$ et sur $]-\infty, 0[$ nous assure la dérivabilité de f en tout réel $x \neq 0$ et montre que, pour tout réel x non nul, $f'(x) = g(x)$. La fonction f' est donc définie sur $\mathbb{R} \setminus \{0\}$ et, pour tout réel $x \neq 0$, $|f'(x)| = |g(x)| \leq 1$, donc la fonction dérivée f' de f est bornée sur $\mathbb{R} \setminus \{0\}$.

Comme $f(0) = 0$, on a, pour tout nombre réel x différent de 0, $Qf(0, x) = \frac{f(x)}{x}$. Pour tout $n \in \mathbb{N}^*$, g est nulle sur le segment $[1/4^n, 2/4^n]$, ce qui donne :

$$f\left(\frac{1}{4^n}\right) = f\left(\frac{2}{4^n}\right) = \sum_{k=n+1}^{+\infty} \frac{1}{4^k} = \frac{1}{3} \times \frac{1}{4^n}, \text{ donc } Qf\left(0, \frac{1}{4^n}\right) = \frac{1}{3} \text{ et } Qf\left(0, \frac{2}{4^n}\right) = \frac{1}{6}.$$

Or les suites $(1/4^n)$ et $(2/4^n)$ sont à valeurs dans $]0, +\infty[$ et convergent vers 0, donc $Qf(0, x)$ n'admet pas de limite quand x tend vers 0 à droite ; f étant impaire, $Qf(0, x)$ n'admet pas non plus de limite quand x tend vers 0 à gauche. Finalement, f n'est dérivable ni à droite ni à gauche en 0. $\square$

Posons-nous le même type de problème, mais au niveau global, en cherchant une fonction définie et continue sur $\mathbb{R}$ dont l'ensemble des points en lesquels elle est dérivable, ainsi que celui en lesquels elle n'est pas dérivable, sont denses dans $\mathbb{R}$.

Dans le chapitre 10 figure un autre exemple de ce type (exemple 10.9, page 201), dans lequel, en tout point où elle n'est pas dérivable, la fonction est dérivable à gauche et à droite.

9.5. Fonction continue sur $\mathbb{R}$, dérivable en tout point de $\mathbb{R} \setminus \mathbb{Q}$, mais qui n'est dérivable ni à droite ni à gauche en chaque point de $\mathbb{Q}$.

Nous utilisons dans cet exemple les propriétés de la convergence uniforme des suites et séries de fonctions, rappelées aux chapitres 12 et 13.

L'ensemble $\mathbb{Q}$ étant dénombrable, $\mathbb{Q} = \{r_n \mid n \in \mathbb{N}\} = \{r_0, r_1, r_2, \dots, r_n, \dots\}$ où $(r_n)_{n \geq 0}$ est une suite de nombres rationnels deux à deux distincts.

Nous utilisons les fonctions g et f définies dans l'exemple précédent 9.4 et nous en déduisons la série de fonctions $\sum_{n \geq 0} u_n$, de terme général :

$$\left| \begin{array}{l} u_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_n(x) = \frac{f(x - r_n)}{2^n}. \end{array} \right.$$

Pour tout nombre réel t , $|f(t)| \leq 1/3 < 1$, donc cette série de fonction converge uniformément sur $\mathbb{R}$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} h : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto h(x) = \sum_{n=0}^{+\infty} u_n(x). \end{array} \right.$$

La fonction f étant continue sur $\mathbb{R}$, il en est de même, pour tout $n \in \mathbb{N}$, de la fonction u_n , ce qui montre que¹ la fonction h est continue sur $\mathbb{R}$. Nous considérons, pour tout entier naturel n , l'application :

$$\left| \begin{array}{l} h_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto h_n(x) = \sum_{k=0}^n u_k(x). \end{array} \right.$$

Soit a un point de $\mathbb{R} \setminus \mathbb{Q}$. Pour tout $n \in \mathbb{N}$, $a - r_n$ est différent de 0, donc u_n est dérivable en a et :

$$u_n'(a) = \frac{f'(a - r_n)}{2^n} = \frac{g(a - r_n)}{2^n}$$

et, comme $|g(t)| \leq 1$ pour tout réel t , la série $\sum_{n \geq 0} u_n'(a)$ converge. Pour tout entier naturel n , la fonction h_n est dérivable en a et $h_n'(a) = \sum_{k=0}^n u_k'(a)$, donc la suite $(h_n'(a))_{n \geq 0}$ converge dans $\mathbb{R}$.

Nous posons $D_n(a) = h_n'(a)$ pour tout $n \in \mathbb{N}$ et $D(a) = \lim_{n \rightarrow \infty} h_n'(a) = \lim_{n \rightarrow \infty} D_n(a)$, et, pour tout nombre réel $x \neq a$:

$$D(x) = \frac{h(x) - h(a)}{x - a} \quad \text{et} \quad D_n(x) = \frac{h_n(x) - h_n(a)}{x - a} \quad \text{pour tout } n \in \mathbb{N}.$$

La fonction f est continue sur $[0, +\infty[$ et dérivable sur son intérieur $]0, +\infty[$, et

1. La méthode consistant à transporter la singularité en 0 d'une fonction f à un ensemble dense, ici $\{r_n \mid n \in \mathbb{N}\}$, en sommant la série de fonctions $x \mapsto c_n f(x - r_n)$ où $\sum_{n \geq 0} c_n$ est une série convergente de réels, ce qui assure la convergence uniforme, fut introduite par Weierstrass vers 1860 et exposée par Cantor en 1882. On l'appelle la méthode de condensation des singularités.

$|f'(x)| = |g(x)| \leq 1$ pour tout $x \in]0, +\infty[$, donc f est lipschitzienne de rapport 1 sur $[0, +\infty[$; de même f est lipschitzienne de rapport 1 sur $] -\infty, 0]$, donc f est lipschitzienne de rapport 1 sur $\mathbb{R}$. Si $p, q \in \mathbb{N}$ et si $p < q$, on a, pour tout $x \neq a$:

$$\begin{aligned} |(h_q(x) - h_q(a)) - (h_p(x) - h_p(a))| &\leq \sum_{n=p+1}^q \frac{|f(x - r_n) - f(a - r_n)|}{2^n} \\ &\leq \sum_{n=p+1}^q \frac{|x - a|}{2^n} \leq \frac{1}{2^p} |x - a| \end{aligned}$$

donc, en divisant par $|x - a|$, $|D_q(x) - D_p(x)| \leq 1/2^p$. Si $p \in \mathbb{N}$, le passage à la limite quand q tend vers l'infini dans l'inégalité précédente donne, pour tout $x \neq a$:

$$|D(x) - D_p(x)| \leq \frac{1}{2^p}.$$

Soit un réel $\varepsilon > 0$. Nous choisissons $p \in \mathbb{N}$ tel que $\frac{1}{2^p} \leq \frac{\varepsilon}{3}$ et $|D_p(a) - D(a)| \leq \frac{\varepsilon}{3}$.

La fonction h_p est dérivable en a et $h_p'(a) = D_p(a)$, donc il existe un réel $\alpha > 0$ tel que, pour tout réel $x \neq a$, $|x - a| < \alpha$ entraîne $|D_p(x) - D_p(a)| < \varepsilon/3$. On a donc, pour tout réel $x \neq a$ tel que $|x - a| < \alpha$:

$$|D(x) - D(a)| \leq \underbrace{|D(x) - D_p(x)|}_{\leq \varepsilon/3} + \underbrace{|D_p(x) - D_p(a)|}_{< \varepsilon/3} + \underbrace{|D_p(a) - D(a)|}_{\leq \varepsilon/3} < \varepsilon.$$

Il en résulte que h est dérivable en a et que $h'(a) = \lim_{n \rightarrow \infty} h_n'(a)$.

Soit $a \in \mathbb{Q}$. Alors $a = r_p$ où $p \in \mathbb{N}$. Par un raisonnement analogue au précédent, la fonction $\ell : x \mapsto \ell(x) = h(x) - u_p(x)$ est dérivable en a . Comme f n'est dérivable ni à droite ni à gauche en 0, la fonction u_p n'est dérivable ni à droite ni à gauche en a . Or $h = \ell + u_p$, donc h n'est dérivable ni à droite ni à gauche en a . $\square$

La dérivabilité en un point entraîne la continuité en ce point. La réciproque est fautive comme le montre l'exemple 9.2 (page 156). On peut même construire des fonctions continues qui ne sont dérivables en aucun point².

2. Jusqu'à la moitié du XIX^e siècle, on pensait généralement qu'une fonction continue était dérivable sauf peut-être en quelques points. Ampère prétendit même l'avoir démontré en 1806. Bolzano donne vers 1830 un exemple de fonction continue, mais dérivable nulle part ; cependant ses écrits restent méconnus. En 1854, Bernhard Riemann, propose sans preuve, la fonction :

$$R : x \mapsto R(x) = \sum_{n=1}^{+\infty} \frac{\sin(n^2 x)}{n^2}.$$

Karl Weierstrass se déclare incapable de le démontrer. Il faut attendre 1971 pour savoir que R n'est pas dérivable sauf en certains points. En 1872, Karl Weierstrass démontre que si a et b sont des réels tels que $a > 0$, $b > 0$ et $ab > 1 + (3\pi)/2$, la fonction :

$$f : x \mapsto f(x) = \sum_{n=1}^{+\infty} b^n \cos(a^n x)$$

est continue sur $\mathbb{R}$ et n'est dérivable en aucun point de $\mathbb{R}$. En 1903, le mathématicien japonais Teiji Takagi (1875-1960) propose les fonctions :

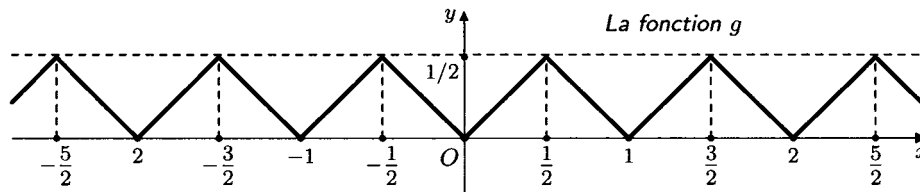
$$f : x \mapsto f(x) = \sum_{n=0}^{+\infty} b^n g(a^n x)$$

où g est la fonction de l'exemple 9.6 qui suit et a et b des réels tels que $0 < b < 1$ et $a \geq 4$; l'exemple 9.6 en est un cas particulier. Lorsque de plus $ab > 2$, la fonction f a des dérivées supérieures égales à $+\infty$ et inférieures égales à $-\infty$ en tout point de $\mathbb{R}$ — voir la note numéro 5 (page 181) et l'exemple 9.30, pages 182 et 183.

9.6. Fonction continue sur $\mathbb{R}$ qui n'est dérivable en aucun point.

Nous utilisons dans cet exemple les propriétés de la convergence uniforme des suites et séries de fonctions, rappelées aux chapitres 12 et 13.

Introduisons l'application $g : x \mapsto g(x) = d(x, \mathbb{Z})$ (distance de x à $\mathbb{Z}$) de $\mathbb{R}$ dans $\mathbb{R}$. Si x est un nombre réel, alors $|(x+1) - n| = |x - (n-1)|$ pour tout $n \in \mathbb{Z}$, donc $d(x+1, \mathbb{Z}) \leq d(x, \mathbb{Z})$, et $|x - n| = |(x+1) - (n+1)|$ pour tout $n \in \mathbb{Z}$, donc $d(x, \mathbb{Z}) \leq d(x+1, \mathbb{Z})$, ce qui montre que $g(x+1) = g(x)$. La fonction g admet donc 1 pour période. Pour tout point x de $[-1/2, 1/2]$, on a clairement $g(x) = |x|$, donc $0 \leq g(x) \leq 1/2$. Par conséquent $0 \leq g(x) \leq 1/2$ pour tout réel x , et la même démonstration que dans l'exemple 8.18 (pages 147 et 148) — obtenue en remplaçant S par $\mathbb{Z}$ — prouve que g est lipschitzienne de rapport 1 sur $\mathbb{R}$.



Nous associons à tout entier naturel n l'application :

$$\left| \begin{array}{l} u_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_n(x) = \frac{g(4^n x)}{4^n} \end{array} \right.$$

Pour tout n , u_n est lipschitzienne de rapport 1 sur $\mathbb{R}$ et admet $1/4^n$ pour période. On a, pour tout $n \in \mathbb{N}$ et tout nombre réel x , $|u_n(x)| \leq 1/(2 \times 4^n)$, donc la série de fonctions $\sum_n u_n$ converge uniformément sur $\mathbb{R}$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \sum_{n=0}^{+\infty} u_n(x) \end{array} \right.$$

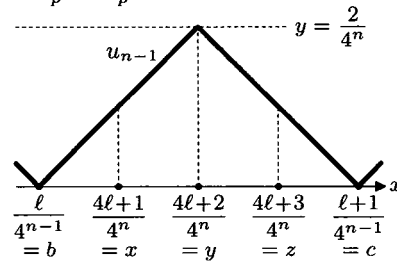
et prouve que la fonction f est continue sur $\mathbb{R}$. Posons $h_n = \frac{1}{4^n}$ pour tout $n \in \mathbb{N}$.

Soit a un nombre réel. Soit $n \in \mathbb{N}^*$. Nous posons, pour tout entier naturel p , $D_p = u_p(a + h_n) - u_p(a)$ et $C_p = u_p(a - h_n) - u_p(a)$. Si $p \in \mathbb{N}$ et si $p \geq n$, h_n et $-h_n$ sont des périodes de la fonction u_p , donc $D_p = C_p = 0$. Par suite :

$$Qf(a, h_n) = \frac{f(a + h_n) - f(a)}{h_n} = \sum_{p=0}^{n-1} \frac{D_p}{h_n}$$

et de même en remplaçant h_n par $-h_n$ et D_p par C_p pour tout $p \in \llbracket 0, n-1 \rrbracket$. Notons ℓ la partie entière de $4^{n-1}a$. Alors $\ell \in \mathbb{Z}$ et :

$$b = \frac{\ell}{4^{n-1}} \leq a < c = \frac{\ell+1}{4^{n-1}}.$$



Nous posons de plus (voir le dessin ci-dessus à droite) :

$$x = \frac{4\ell+1}{4^n}, y = \frac{4\ell+2}{4^n} \text{ et } z = \frac{4\ell+3}{4^n}.$$

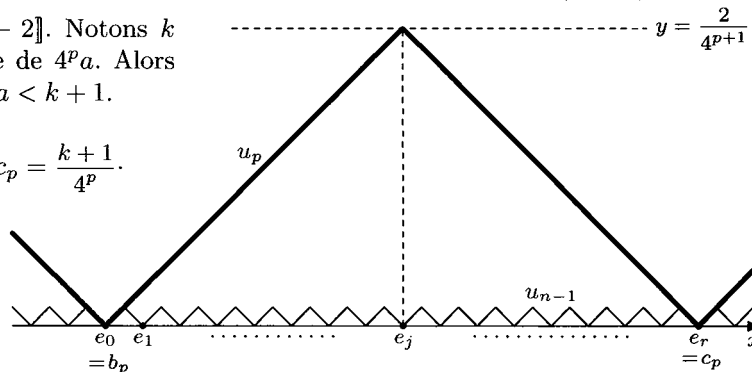
Alors y est le milieu de (b, c) et : si $b \leq a < x$, $b \leq a < a + h_n < y$; si $x \leq a < y$, $b \leq a - h_n < a < y$; si $y \leq a < z$, $y \leq a < a + h_n < c$; si $z \leq a < c$, $y \leq a - h_n < a < c$. Par suite, si $b \leq a < x$ ou $y \leq a < z$ (premier cas),

$D_{n-1}/h_n = \pm 1$, et si $x \leq a < y$ ou $z \leq a < c$ (second cas), $C_{n-1}/(-h_n) = \pm 1$.

Soit $p \in \llbracket 0, n-2 \rrbracket$. Notons k la partie entière de $4^p a$. Alors $k \in \mathbb{Z}$ et $k \leq 4^p a < k+1$.

Posons :

$$b_p = \frac{k}{4^p} \text{ et } c_p = \frac{k+1}{4^p}.$$



On a $4^{n-p-1}k \leq 4^{n-1}a < 4^{n-p-1}(k+1)$, donc $4^{n-p-1}k \leq \ell < 4^{n-p-1}(k+1)$. Alors $\ell+1 \leq 4^{n-p-1}(k+1)$, donc $4\ell+4 \leq 4^{n-p}(k+1)$. Posons $m_0 = 4^{n-p-1}k$, $r = 4^{n-p-1}$, $m_i = m_0 + i$ pour tout $i \in \llbracket 0, r \rrbracket$ et $j = r/2$. Comme $n-p-1 \geq 1$, $j \in \mathbb{N}$.

De plus $j \in \llbracket 1, r-1 \rrbracket$, $e_j = \frac{m_j}{4^{n-1}}$ est le milieu de $[b_p, c_p]$, on a :

$$b_p = e_0 = \frac{m_0}{4^{n-1}} < e_1 = \frac{m_1}{4^{n-1}} < \dots < e_j = \frac{m_j}{4^{n-1}} < \dots < e_{r-1} = \frac{m_{r-1}}{4^{n-1}} < \frac{m_r}{4^{n-1}} = e_r = c_p$$

et $\ell \in \{m_0, m_1, \dots, m_j, \dots, m_{r-1}\}$, donc $[b, c]$ est l'un des $[e_i, e_{i+1}]$. Le « coefficient directeur » de u_p sur $[b, c]$ est ± 1 (voir le dessin ci-dessus) ; c'est $+1$ si $i < j$ et -1 si $i \geq j$, donc $D_p/h_n = \pm 1$ dans le premier cas et $C_p/(-h_n) = \pm 1$ dans le second. Par conséquent, dans le premier cas (resp. dans le second), $Qf(a, h_n)$ (resp. $Qf(a, -h_n)$) est la somme de n termes égaux à ± 1 , donc $Qf(a, h_n)$ (resp. $Qf(a, -h_n)$) est un nombre entier de même parité que n . Or $\lim_{(n \rightarrow \infty)} h_n = 0$, donc $Qf(a, h_n)$ (premier cas) ou $Qf(a, -h_n)$ (second cas) n'admet pas de limite quand n tend vers l'infini, ce qui prouve que f n'est pas dérivable en a . En conclusion, f n'est dérivable en aucun point de $\mathbb{R}$. $\square$

9.7. Fonction continue sur le segment $[0, 1]$ qui n'est dérivable en aucun point.

On définit, par récurrence sur $n \in \mathbb{N}$, la suite $(f_n)_{n \geq 0}$ d'applications de $[0, 1]$ dans $\mathbb{R}$ de la manière suivante : $f_0(x) = x$ pour tout $x \in [0, 1]$ et, pour tout entier naturel n , f_{n+1} est l'unique application de $[0, 1]$ dans $\mathbb{R}$ telle que :

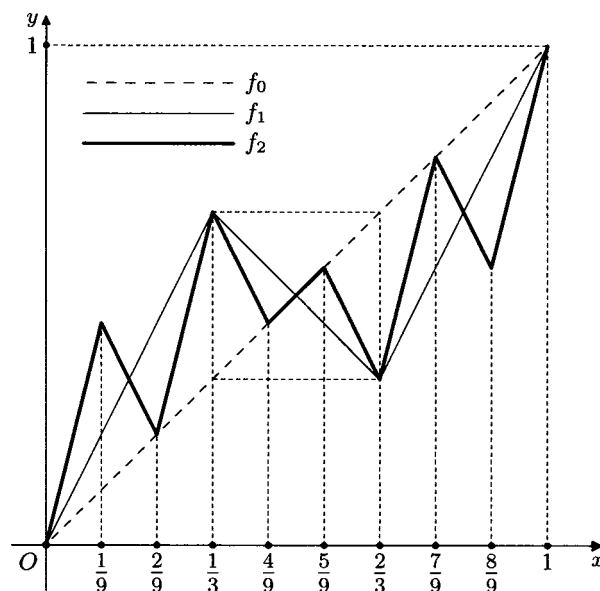
Pour tout $i \in \llbracket 0, 3^{n+1}-1 \rrbracket$, la restriction de f_{n+1} à $[i/3^{n+1}, (i+1)/3^{n+1}]$ est la restriction à ce segment d'une application affine, pour tout $k \in \llbracket 0, 3^n \rrbracket$:

$$f_{n+1}\left(\frac{3k}{3^{n+1}}\right) = f_n\left(\frac{k}{3^n}\right)$$

et, pour tout $k \in \llbracket 0, 3^n-1 \rrbracket$:

$$f_{n+1}\left(\frac{3k+1}{3^{n+1}}\right) = f_n\left(\frac{3k+2}{3^{n+1}}\right) \text{ et } f_{n+1}\left(\frac{3k+2}{3^{n+1}}\right) = f_n\left(\frac{3k+1}{3^{n+1}}\right).$$

Pour tout entier naturel n , f_n est une application de $[0, 1]$ dans $\mathbb{R}$, affine par morceaux continue sur le segment $[0, 1]$ — voir le dessin de la page suivante. Pour tout $n \in \mathbb{N}$ et tout $i \in \llbracket 0, 3^n-1 \rrbracket$, on note $\alpha(n, i)$ et $\beta(n, i)$ les réels tels que, pour tout point x du segment $[i/3^n, (i+1)/3^n]$, $f_n(x) = \alpha(n, i)x + \beta(n, i)$.



Soit $n \in \mathbb{N}$ et $i \in [0, 3^n - 1]$. On a, par définition de f_{n+1} :

$$f_{n+1}\left(\frac{3i+1}{3^{n+1}}\right) = f_n\left(\frac{3i+2}{3^{n+1}}\right) = f_n\left(\frac{i}{3^n} + \frac{2}{3^{n+1}}\right) = f_n\left(\frac{i}{3^n}\right) + \frac{2}{3^{n+1}}\alpha(n, i)$$

donc, comme $f_{n+1}\left(\frac{3i}{3^{n+1}}\right) = f_n\left(\frac{i}{3^n}\right)$, $f_{n+1}\left(\frac{3i+1}{3^{n+1}}\right) - f_{n+1}\left(\frac{3i}{3^{n+1}}\right) = \frac{2}{3^{n+1}}\alpha(n, i)$.

Par suite, en divisant par $1/3^{n+1}$, on obtient que $\alpha(n+1, 3i) = 2\alpha(n, i)$. Par des calculs analogues, il vient : $\alpha(n+1, 3i+1) = -\alpha(n, i)$ et $\alpha(n+1, 3i+2) = 2\alpha(n, i)$.

On a $\alpha(0, 0) = 1$, $\alpha(n+1, i) = 2\alpha(n, k)$ ou $-\alpha(n, k)$ pour un certain k , donc $|\alpha(n+1, i)| \leq 2|\alpha(n, k)|$, $3i$ est pair si, et seulement si, i est pair, $3i+1$ est pair si, et seulement si, i est impair et $3i+2$ est pair si, et seulement si, i est pair. Un raisonnement facile par récurrence montre alors que, pour tout $n \in \mathbb{N}$ et tout $i \in [0, 3^n - 1]$, $\alpha(n, i) \in \mathbb{Z}$, $|\alpha(n, i)| \leq 2^n$, $\alpha(n, i) > 0$ si i est pair et $\alpha(n, i) < 0$ si i est impair. En particulier, $\alpha(n, i) \neq 0$ pour tout $n \in \mathbb{N}$ et tout $i \in [0, 3^n - 1]$.

Nous majorons, pour tout $n \in \mathbb{N}$ et tout $x \in [0, 1]$, le nombre réel $|f_{n+1}(x) - f_n(x)|$. Soit $n \in \mathbb{N}$ et $x \in [0, 1]$. Il existe $i \in [0, 3^n - 1]$ tel que $x \in [i/3^n, (i+1)/3^n]$.

On a, dans le cas où $\frac{i}{3^n} \leq x \leq \frac{3i+1}{3^{n+1}}$, $f_n(x) = f_n\left(\frac{i}{3^n}\right) + \left(x - \frac{i}{3^n}\right)\alpha(n, i)$ et :

$$f_{n+1}(x) = f_{n+1}\left(\frac{3i}{3^{n+1}}\right) + \left(x - \frac{i}{3^n}\right)\alpha(n+1, 3i) = f_n\left(\frac{i}{3^n}\right) + 2\left(x - \frac{i}{3^n}\right)\alpha(n, i),$$

donc $f_{n+1}(x) - f_n(x) = \left(x - \frac{i}{3^n}\right)\alpha(n, i)$, et on déduit des inégalités $\left|x - \frac{i}{3^n}\right| \leq \frac{1}{3^{n+1}}$ et $|\alpha(n, i)| \leq 2^n$ que :

$$|f_{n+1}(x) - f_n(x)| \leq \frac{2^n}{3^{n+1}}.$$

Lorsque $(3i+1)/3^{n+1} < x < (3i+2)/3^{n+1}$, on écrit que le nombre réel $f_n(x)$ est la somme $f_n(3i+1)/3^{n+1} + (x - (3i+1)/3^{n+1})\alpha(n, i)$ et que $f_{n+1}(x)$ est la somme $f_{n+1}((3i+1)/3^{n+1}) + (x - (3i+1)/3^{n+1})\alpha(n+1, 3i+1)$, donc on a $f_{n+1}(x) = f_n((3i+2)/3^{n+1}) - (x - (3i+1)/3^{n+1})\alpha(n, i)$, et on prouve comme dans le premier cas que $|f_{n+1}(x) - f_n(x)| \leq 2^n/3^{n+1}$.

Enfin, si $(3i+2)/3^{n+1} \leq x \leq (i+1)/3^n$, on écrit que le nombre réel $f_n(x)$ est la somme $f_n((i+1)/3^n) + (x - (i+1)/3^n)\alpha(n, i)$ et $f_{n+1}(x)$ la somme des réels $f_{n+1}((3i+3)/3^{n+1})$ et $(x - (3i+3)/3^{n+1})\alpha(n+1, 3i+2)$, d'où l'on déduit l'égalité $f_{n+1}(x) = f_n((i+1)/3^n) + 2(x - (i+1)/3^n)\alpha(n, i)$ qui permet de conclure comme dans le premier cas que $|f_{n+1}(x) - f_n(x)| \leq 2^n/3^{n+1}$.

Ainsi, pour tout $n \in \mathbb{N}$ et tout point x de $[0, 1]$, $|f_{n+1}(x) - f_n(x)| \leq \frac{2^n}{3^{n+1}} = \frac{1}{3} \left(\frac{2}{3}\right)^n$.

Soit p et q des entiers naturels tels que $p < q$. On a, pour tout point x de $[0, 1]$:

$$\begin{aligned} |f_q(x) - f_p(x)| &= \left| \sum_{n=p}^{q-1} (f_{n+1}(x) - f_n(x)) \right| \leq \sum_{n=p}^{q-1} |f_{n+1}(x) - f_n(x)| \\ &\leq \frac{1}{3} \sum_{n=p}^{q-1} \left(\frac{2}{3}\right)^n = \frac{1}{3} \left(\frac{2}{3}\right)^p \sum_{i=0}^{q-p-1} \left(\frac{2}{3}\right)^i \leq \frac{1}{3} \left(\frac{2}{3}\right)^p \frac{1 - \left(\frac{2}{3}\right)^{q-p}}{1 - \frac{2}{3}} \leq \left(\frac{2}{3}\right)^p. \end{aligned}$$

La suite $((2/3)^n)$ converge vers 0, donc le critère de Cauchy uniforme montre que la suite (f_n) converge uniformément sur $[0, 1]$. Nous notons f la limite uniforme sur le segment $[0, 1]$ de la suite de fonctions (f_n) . Pour tout entier naturel n , f_n est continue sur $[0, 1]$, donc **f est continue sur $[0, 1]$** .

Il nous reste à démontrer que **f n'est dérivable en aucun point de $[0, 1]$** . Nous distinguons deux cas, les nombres triadiques et les autres.

Si x est un nombre triadique appartenant à $[0, 1]$, alors $x = k/3^n$ où $n \in \mathbb{N}$ et $k \in \llbracket 0, 3^n \rrbracket$, donc :

$$f_{n+1}(x) = f_{n+1}\left(\frac{k}{3^n}\right) = f_{n+1}\left(\frac{3k}{3^{n+1}}\right) = f_n\left(\frac{k}{3^n}\right) = f_n(x),$$

puis :

$$f_{n+2}(x) = f_{n+2}\left(\frac{k}{3^n}\right) = f_{n+2}\left(\frac{3 \times (3k)}{3^{n+2}}\right) = f_{n+1}\left(\frac{3k}{3^{n+1}}\right) = f_{n+1}(x) = f_n(x),$$

et on obtient ainsi, de proche en proche, que pour tout $p \in \mathbb{N}$, $f_{n+p}(x) = f_n(x)$, ce qui, en passant à la limite quand p tend vers l'infini, montre que $f(x) = f_n(x)$.

Soit a un nombre triadique tel que $0 \leq a < 1$. Alors $a = h/3^p$ où $p \in \mathbb{N}$ et où $h \in \llbracket 0, 3^p - 1 \rrbracket$. On introduit la suite $(x_n)_{n \geq 0}$, de terme général $x_n = a + (1/3^{p+n})$. Pour tout $n \in \mathbb{N}$, $x_n \in [0, 1]$ et x_n est le nombre triadique $((3^n h) + 1)/3^{p+n}$, donc $f(x_n) = f_{p+n}(x_n)$. Pour tout entier naturel n , $f_{p+n}(a) = f_p(a) = f(a)$, donc, en posant $i = 3^n h$, on a $f(x_n) - f(a) = f_{p+n}(x_n) - f_{p+n}(a)$ donc :

$$f(x_n) - f(a) = f_{p+n}\left(\frac{i+1}{3^{p+n}}\right) - f_{p+n}\left(\frac{i}{3^{p+n}}\right) = \frac{1}{3^{p+n}} \alpha(p+n, i) = \frac{1}{3^{p+n}} \alpha(p+n, 3^n h)$$

et, comme $x_n - a = 1/3^{p+n}$, il vient : $Qf(a, x_n) = \frac{f(x_n) - f(a)}{x_n - a} = \alpha(p+n, 3^n h)$.

Pour tout $n \in \mathbb{N}$, $\alpha(p+(n+1), 3^{n+1}h) = \alpha((p+n)+1, 3 \times (3^n h)) = 2\alpha(p+n, 3^n h)$. On en déduit de proche en proche que, pour tout $n \in \mathbb{N}$, $\alpha(p+n, 3^n h) = 2^n \alpha(p, h)$.



donc $|\alpha(p+n, 3^n h)| = 2^n |\alpha(p, h)|$ et, comme $|\alpha(p, h)| \geq 1$, on obtient la minoration $|\alpha(p+n, 3^n h)| \geq 2^n$ donc $|\mathbb{Q}f(a, x_n)| \geq 2^n$. Or $\lim_{(n \rightarrow \infty)} x_n = a^+$ donc, si f était dérivable à droite en a , la suite $(\mathbb{Q}f(a, x_n))$ convergerait dans $\mathbb{R}$ vers $f'_d(a)$, alors qu'elle tend vers $+\infty$. Par suite f n'est pas dérivable à droite en a .

Si $a = h/3^p$ est un nombre triadique tel que $0 < a \leq 1$, la suite (x_n) de terme général $x_n = a - (1/3^{p+n})$ permet de prouver que f n'est pas dérivable à gauche en a .

Soit a un point non triadique du segment $[0, 1]$. Alors $0 < a < 1$. On note $(x_n)_{n \geq 1}$ (resp. $(y_n)_{n \geq 1}$) la suite des valeurs triadiques approchées par défaut (resp. par excès) de a . Pour tout $n \in \mathbb{N}$, $x_n = i_n/3^n$ et $y_n = (1 + i_n)/3^n$ où i_n est la partie entière de $3^n a$, on a $x_n \leq a < y_n$ et a n'est pas triadique, donc $x_n < a < y_n$, ce qui montre que $i_n \in \llbracket 0, 3^n - 1 \rrbracket$ et que $0 \leq x_n < a < y_n \leq 1$. Ceci justifie l'existence de la suite $(u_n)_{n \geq 1}$ de terme général :

$$u_n = \frac{f(y_n) - f(x_n)}{y_n - x_n}.$$

Soit $n \in \mathbb{N}$. Comme $x_n = i_n/3^n$ et que $y_n = (1 + i_n)/3^n$, on a $f(y_n) = f_n(y_n)$ et $f(x_n) = f_n(x_n)$ donc $u_n = \alpha(n, i_n)$. De $x_n \leq x_{n+1} < a < y_{n+1} \leq y_n$ on déduit, en multipliant par $3^{n+1} > 0$, que $3i_n \leq i_{n+1} < 3i_n + 3$, donc que $i_{n+1} = 3i_n$, $3i_n + 1$ ou $3i_n + 2$, ce qui montre que $\alpha(n+1, i_{n+1}) = 2\alpha(n, i_n)$ ou $-\alpha(n, i_n)$. Par conséquent $u_{n+1} = 2u_n$ ou $-u_n$, donc $u_{n+1} - u_n = u_n$ ou $-2u_n$, d'où l'on déduit que $|u_{n+1} - u_n| = |u_n|$ ou $2|u_n|$. Or $|u_n| = |\alpha(n, i_n)| \geq 1$ donc $|u_{n+1} - u_n| \geq 1$. Pour tout $n \in \mathbb{N}$, $x_n < a < y_n$ et les suites (x_n) et (y_n) convergent vers a . Si f était dérivable en a , la suite (u_n) convergerait dans $\mathbb{R}$ vers le réel³ $f'(a)$, contredisant le fait que $|u_{n+1} - u_n| \geq 1$ pour tout n . La fonction f n'est donc pas dérivable en a . En conclusion, f est continue sur $[0, 1]$ et n'est dérivable en aucun point de $[0, 1]$. $\square$

■ Discontinuité de la fonction dérivée

Toute fonction f définie et continue sur un intervalle I est la fonction dérivée sur I d'une fonction définie et dérivable sur I , ce qui signifie qu'elle possède des primitives sur I ; en effet, en choisissant un point a de l'intervalle I , l'application $F : x \mapsto F(x) = \int_a^x f(t) dt$ de I dans $\mathbb{R}$ est dérivable sur I et $F'(x) = f(x)$ pour tout $x \in I$. Ceci ne se généralise pas aux fonctions discontinues.

9.8. Application de $\mathbb{R}$ dans $\mathbb{R}$ discontinue en un seul point n'ayant pas de primitive sur $\mathbb{R}$.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 1 & \text{si } x = 0, \\ 0 & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

La fonction f est continue en tout nombre réel non nul et discontinue en 0; en effet, f est constante au voisinage de tout réel différent de 0 et, comme la suite $(1/n)_{n \geq 1}$

3. On pose $\lambda = f'(a)$ et on part de ce que, pour tout n , on a : $y_n - a > 0$, $x_n - a < 0$ et :
 $f(y_n) - f(x_n) - \lambda(y_n - x_n) = (f(y_n) - f(a) - \lambda(y_n - a)) - (f(x_n) - f(a) - \lambda(x_n - a)).$

converge vers 0 et que $f(1/n) = 0$ pour tout entier $n \geq 1$, la suite $(f(1/n))$ converge vers 0 $\neq f(0)$.

Supposons que f admette une primitive F sur $\mathbb{R}$. Alors F est une application de $\mathbb{R}$ dans $\mathbb{R}$, dérivable sur $\mathbb{R}$, telle que $F'(x) = f(x)$ pour tout nombre réel x . Pour tout point x de $]0, +\infty[$, l'égalité des accroissements finis appliquée à F sur $[0, x]$ fournit $c \in]0, x[$ tel que $F(x) - F(0) = (x - 0)F'(c) = xf(c) = 0$, donc :

$$QF(0, x) = \frac{F(x) - F(0)}{x} = 0.$$

Il en résulte que $QF(0, x)$ admet pour limite 0 quand x tend vers 0 à droite, donc que $F'(0) = 0$, en contradiction avec $f(0) = 1$. En conclusion, la fonction f n'admet pas de primitive sur $\mathbb{R}$. $\square$

On prouve, en utilisant comme dans l'exemple précédent l'égalité des accroissements finis, que si a est un réel et r un réel strictement positif, si f est une fonction définie et continue sur $[a, a+r]$ (resp. sur $[a-r, a]$) et dérivable sur $]a, a+r[$ (resp. sur $]a-r, a[$) et si $f'(x)$ admet une limite ℓ dans $\mathbb{R}$ quand x tend vers a à droite (resp. à gauche), la fonction f est dérivable à droite (resp. à gauche) en a et la dérivée à droite (resp. à gauche) de f en a est égale à ℓ . Par suite, si a est un nombre réel et r un réel strictement positif, si f est une fonction définie et continue sur $[a-r, a+r]$ et dérivable sur $]a-r, a[$ et sur $]a, a+r[$ et si $f'(x)$ admet une limite ℓ dans $\mathbb{R}$ quand x tend vers a , f est dérivable en a et $f'(a) = \ell$.

On en déduit que si une fonction f est définie et dérivable sur un voisinage ouvert d'un réel a , la fonction dérivée f' est continue en a ou bien elle n'admet pas de limite en a . Si une fonction f est définie et dérivable sur un voisinage ouvert d'un réel a , la fonction dérivée f' n'a aucune raison d'admettre une limite en a — si ceci était toujours vrai, toutes les fonctions dérivées seraient continues.

9.9. Fonction dérivable sur $\mathbb{R}$ dont la fonction dérivée n'admet pas de limite à droite ni à gauche en 0.

Nous introduisons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ x^2 \sin \frac{1}{x} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

La fonction f est dérivable en tout point de $\mathbb{R} \setminus \{0\}$ et on a, pour tout réel $x \neq 0$:

$$(1) \quad f'(x) = 2x \sin \frac{1}{x} - \cos \frac{1}{x}.$$

Pour tout réel $x \neq 0$, $|f(x)| \leq x^2$ et $Qf(0, x) = f(x)/x$, donc $|Qf(0, x)| \leq |x|$, ce qui montre que $Qf(0, \bullet)$ admet pour limite 0 en 0 : f est dérivable en 0 et $f'(0) = 0$. Ainsi f est dérivable sur $\mathbb{R}$.

On déduit de (1) que, pour tout nombre réel $x \neq 0$, $\cos \frac{1}{x} = 2x \sin \frac{1}{x} - f'(x)$.

Pour tout $x \neq 0$, $|x \sin(1/x)| \leq |x|$, donc $x \sin(1/x)$ admet pour limite 0 quand x tend vers 0. Il en résulte que si f' admettait une limite ℓ en 0 à droite (resp. à gauche), $\cos(1/x)$ admettrait pour limite $-\ell$ quand x tend vers 0 à droite (resp. à gauche). Or la suite $(a_n)_{n \geq 1}$ (resp. $(b_n)_{n \geq 1}$), de terme général $a_n = 1/(n\pi)$ (resp. $b_n = -1/(n\pi)$), est à valeurs dans $]0, +\infty[$ (resp. dans $] -\infty, 0[$)

et converge vers 0, alors que la suite de terme général $\cos(1/a_n) = (-1)^n$ (resp. $\cos(1/b_n) = (-1)^n$) diverge. En conclusion, la fonction dérivée f' n'admet pas de limite à droite ni à gauche en 0. $\square$

Dans l'exemple que nous venons de traiter, la fonction dérivée f' de f est bornée sur tout voisinage de 0. Voici un exemple où la fonction dérivée n'est bornée sur aucun voisinage de 0.

9.10. Fonction dérivable sur $\mathbb{R}$ dont la fonction dérivée n'est bornée sur aucun voisinage de 0.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ x^2 \sin \frac{1}{|x|^{4/3}} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

Pour tout réel $x \neq 0$, $|f(x)| \leq x^2$ et $Qf(0, x) = \frac{f(x)}{x}$, donc $|Qf(0, x)| \leq |x|$. Par suite $Qf(0, \bullet)$ admet pour limite 0 en 0, donc f est dérivable en 0 et $f'(0) = 0$. La fonction f est dérivable en tout point de $\mathbb{R} \setminus \{0\}$ et, pour tout réel $x \neq 0$:

$$f'(x) = \underbrace{2x \sin \frac{1}{|x|^{4/3}}}_{= g(x)} \pm \underbrace{\frac{4}{3\sqrt[3]{|x|}} \cos \frac{1}{|x|^{4/3}}}_{= h(x)} \quad (- \text{ pour } x > 0, + \text{ pour } x < 0).$$

Il en résulte que la fonction f est dérivable sur $\mathbb{R}$. On a $|g(x)| \leq 2|x|$ pour tout réel $x \neq 0$, donc $g(x)$ admet pour limite 0 quand x tend vers 0. Nous introduisons les suites $(a_n)_{n \geq 1}$, $(u_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$, de termes généraux :

$$a_n = (2n\pi)^{3/4}, \quad u_n = \frac{1}{a_n} \quad \text{et} \quad v_n = -\frac{1}{a_n}.$$

La suite (a_n) tend vers $+\infty$, (u_n) est à valeurs dans $]0, +\infty[$ et converge vers 0, (v_n) est à valeurs dans $]-\infty, 0[$ et converge vers 0 et, pour tout entier $n \geq 1$:

$$h(u_n) = h(v_n) = \frac{4}{3} \sqrt[3]{a_n} \cos((a_n)^{4/3}) = \frac{4}{3} (2n\pi)^{1/4} \cos(2n\pi) = \frac{4}{3} (2n\pi)^{1/4}$$

donc la suite $(f'(u_n))$ tend vers $-\infty$ et la suite $(f'(v_n))$ vers $+\infty$. On en déduit que la fonction dérivée f' n'est bornée sur aucun voisinage de 0. $\square$

9.11. Fonction dérivable sur $] -1, 1[$ dont la fonction dérivée n'est bornée sur aucun voisinage de 0.

Nous introduisons l'application :

$$\left| \begin{array}{l} f :] -1, 1[\longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ \frac{x}{\ln|x|} \cos \frac{1}{x} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

On a, pour tout réel $x \neq 0$:

$$|Qf(0, x)| = \left| \frac{f(x)}{x} \right| = \frac{\left| \cos \frac{1}{x} \right|}{|\ln|x||} \leq \frac{1}{|\ln|x||},$$

donc $Qf(0, \bullet)$ admet pour limite 0 en 0. Par suite f est dérivable en 0 et $f'(0) = 0$.

La fonction f est dérivable en tout point de $] -1, 1[\setminus \{0\}$ et le calcul donne, pour tout $x \in] -1, 1[\setminus \{0\}$:

$$(1) \quad f'(x) = \frac{\cos \frac{1}{x}}{\ln|x|} - \frac{\cos \frac{1}{x}}{\ln^2|x|} + \frac{\sin \frac{1}{x}}{x \ln|x|}.$$

Nous considérons les suites $(u_n)_{n \geq 0}$ et $(v_n)_{n \geq 0}$, à valeurs dans l'intervalle $]0, 1[$, de termes généraux :

$$u_n = \frac{1}{\frac{\pi}{2} + (2n+1)\pi} \quad \text{et} \quad v_n = \frac{1}{\frac{\pi}{2} + 2n\pi}.$$

Pour tout $n \in \mathbb{N}$, $\sin(1/u_n) = \sin(\pi/2 + (2n+1)\pi) = (-1)^{2n+1} \sin(\pi/2) = -1$, $\sin(1/v_n) = \sin(\pi/2 + 2n\pi) = 1$ et $\cos(1/u_n) = \cos(1/v_n) = 0$, donc (1) donne :

$$f'(u_n) = -\frac{1}{u_n \ln u_n} \quad \text{et} \quad f'(v_n) = \frac{1}{v_n \ln v_n}.$$

Les suites (u_n) et (v_n) convergent vers 0, on a $0 < u_n < 1$ et $0 < v_n < 1$ pour tout entier naturel n et $\lim_{(x \rightarrow 0^+)} (x \ln x) = 0^-$, donc la suite $(f'(u_n))_{n \geq 0}$ tend vers $+\infty$ et la suite $(f'(v_n))_{n \geq 0}$ vers $-\infty$.

Il en résulte que la fonction dérivée f' n'est bornée sur aucun voisinage de 0. $\square$

Nous savons par le théorème de Darboux que la fonction dérivée d'une fonction définie et dérivable sur un intervalle vérifie sur cet intervalle le théorème des valeurs intermédiaires⁴. Ainsi dans les exemples précédents 9.10 et 9.11, la fonction dérivée f' prend une infinité de fois toutes les valeurs réelles sur tout intervalle $] -a, a[$ où $a > 0$ dans l'exemple 9.10 et $0 < a < 1$ dans l'exemple 9.11 ; en particulier $f'(\mathbb{R}) = \mathbb{R}$ dans l'exemple 9.10 et $f'([-1, 1]) = \mathbb{R}$ dans l'exemple 9.11.

Un résultat, assez difficile puisqu'il fait appel aux propriétés des ensembles de Baire, permet d'affirmer que si la fonction f est dérivable sur un intervalle I , l'ensemble des points en lesquels la fonction dérivée f' est continue est dense dans I . Il existe des fonctions dérivables où l'ensemble des points où f' est discontinue est, lui aussi, dense dans I .

9.12. Fonction dérivable sur $] -1, 1[$ de fonction dérivée continue en un point si, et seulement si, ce point est irrationnel.

Nous utilisons dans cet exemple les propriétés de la convergence uniforme des suites et séries de fonctions, rappelées aux chapitres 12 et 13.

Nous reprenons l'application f définie à l'exemple 9.9 (pages 166 et 167) :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ x^2 \sin \frac{1}{x} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

La fonction f est dérivable sur $\mathbb{R}$ et, pour $x \neq 0$, $f'(x) = 2x \sin(1/x) - \cos(1/x)$, ce qui montre que la fonction dérivée f' est continue en tout réel différent de zéro.

L'inégalité $|(\sin y)/y| \leq 1$, valable pour tout réel $y \neq 0$, montre que, pour tout réel $x \neq 0$, $|2x \sin(1/x)| \leq 2$, donc que $|f'(x)| \leq 3$ pour tout $x \in \mathbb{R}$; pour la même raison, on a $|f(x)| \leq |x|$ pour tout réel x .

4. Voir, dans le chapitre 8, la note 7 page 137 ; voir aussi [GOUR], chapitre 1, §4.4, exercice 8.

L'ensemble $Q = \mathbb{Q} \cap]-1, 1[$ est dénombrable donc $Q = \{r_0, r_1, r_2, \dots, r_n, \dots\}$ où $(r_n)_{n \geq 0}$ est une suite d'éléments deux à deux distincts de Q .

Nous associons à tout entier naturel n l'application :

$$\left| \begin{array}{l} u_n :]-1, 1[\longrightarrow \mathbb{R} \\ x \longmapsto u_n(x) = \frac{f(x-r_n)}{2^n}. \end{array} \right.$$

Pour tout n , u_n est dérivable sur $] -1, 1[$ et $u_n'(x) = \frac{f'(x-r_n)}{2^n}$ pour tout $x \in]-1, 1[$.

Pour tout entier naturel n et tout $x \in]-1, 1[$, $|x - r_n| \leq 2$ donc $|u_n(x)| \leq 2/2^n$ et $|u_n'(x)| \leq 3/2^n$. Par conséquent les séries de fonctions $\sum_{n \geq 0} u_n$ et $\sum_{n \geq 0} u_n'$ convergent uniformément sur l'intervalle $] -1, 1[$. En particulier l'application :

$$\left| \begin{array}{l} F :]-1, 1[\longrightarrow \mathbb{R} \\ x \longmapsto F(x) = \sum_{n=0}^{+\infty} u_n(x) \end{array} \right.$$

est continue sur $] -1, 1[$. On dispose de plus de l'application :

$$\left| \begin{array}{l} G :]-1, 1[\longrightarrow \mathbb{R} \\ x \longmapsto G(x) = \sum_{n=0}^{+\infty} u_n'(x). \end{array} \right.$$

Soit $a \in]-1, 1[$ et supposons a irrationnel. Pour tout $n \in \mathbb{N}$, $a - r_n \neq 0$ donc u_n' est continue en a , d'où l'on déduit que F est dérivable en a et que $F'(a) = G(a)$.

Soit q un entier naturel. On a $F = (F - u_q) + u_q$. La fonction f étant dérivable en 0, u_q est dérivable en r_q . Le même raisonnement que pour a irrationnel montre que $F - u_q$ est dérivable en r_q et que $(F - u_q)'(r_q) = G(r_q) - u_q'(r_q)$, ce qui montre que F est dérivable en r_q et que $F'(r_q) = (F' - u_q')(r_q) + u_q'(r_q) = G'(r_q)$.

Finalement, F est dérivable sur $] -1, 1[$ et, pour tout $x \in]-1, 1[$, $F'(x) = G(x)$.

Si $a \in]-1, 1[$ et si a est irrationnel, alors, pour tout $n \in \mathbb{N}$, $a - r_n \neq 0$ donc u_n' est continue en a , ce qui montre que $F' = G$ est continue en a .

Soit $a \in Q$. Il existe $q \in \mathbb{N}$ tel que $a = r_q$. Pour tout $n \neq q$, u_n' est continue en a , donc $G - u_q'$ est continue en a . Comme $a - r_q = 0$ et que f' est discontinue en 0, u_q' est discontinue en a . Or $F' = G = (G - u_q') + u_q'$, donc F' est discontinue en a .

En conclusion, F est dérivable sur $] -1, 1[$ et sa fonction dérivée F' est continue en un point a de $] -1, 1[$ si, et seulement si, a est irrationnel. $\square$

■ Sens de variation d'une fonction dérivable

L'une des utilisations les plus fréquentes de la fonction dérivée d'une fonction définie et dérivable sur un intervalle, est l'étude de ses variations sur cet intervalle.

THÉORÈME 9.1. — Une fonction f , définie et continue sur un intervalle I et dérivable sur l'intérieur I_0 de I , est croissante sur I si, et seulement si, sa fonction dérivée f' est positive ou nulle sur I_0 .

On peut, dans le texte de ce théorème, supposer la fonction f dérivable sur I et remplacer I_0 par I .

9.13. Fonction ayant une fonction dérivée positive ou nulle mais qui n'est pas croissante.

L'application :

$$\left| \begin{array}{l} f : \mathbb{R} \setminus \{0\} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = -\frac{1}{x} \end{array} \right.$$

est définie et dérivable sur l'ouvert $\Omega = \mathbb{R} \setminus \{0\}$ et, pour tout $x \in \Omega$, $f'(x) = \frac{1}{x^2} \geq 0$.

Cependant $1 > -1$ et $f(1) = -1 < f(-1) = 1$, donc f n'est pas croissante sur Ω . $\square$

L'exemple précédent 9.13 montre l'importance du mot *intervalle* dans le texte du théorème 9.1 (page 169).

On peut rechercher une condition nécessaire et suffisante sur sa fonction dérivée pour qu'une fonction, définie et continue sur un intervalle I et dérivable sur l'intérieur de I , soit strictement croissante sur I .

A l'aide de l'égalité des accroissements finis, on prouve que si une fonction est définie et continue sur un intervalle I et dérivable sur l'intérieur de I , et si sa fonction dérivée est strictement positive sur l'intérieur de I , alors elle est strictement croissante sur I . Cette condition n'est pas nécessaire, comme le montre l'exemple suivant.

9.14. Fonction dérivable et strictement croissante sur $\mathbb{R}$ telle que $f'(0) = 0$.

Nous considérons l'application $f : x \mapsto f(x) = x^3$ de $\mathbb{R}$ dans $\mathbb{R}$. La fonction f est dérivable sur $\mathbb{R}$ et, pour tout réel x , $f'(x) = 3x^2$; en particulier, $f'(0) = 0$ et 0 est le seul point en lequel f' s'annule. Pour tout nombre réel t :

$$t^2 + t + 1 = \left(t + \frac{1}{2}\right)^2 + \frac{3}{4} > 0$$

donc, si $x, y \in \mathbb{R}$ et $x \neq y$, la factorisation par x^2 si $x \neq 0$ et par y^2 si $y \neq 0$ donne l'inégalité stricte $y^2 + xy + x^2 > 0$. Quels que soient les réels x et y tels que $x < y$, $y^3 - x^3 = (y - x)(y^2 + xy + x^2)$, $y - x > 0$ et $y^2 + xy + x^2 > 0$, donc $y^3 - x^3 > 0$, ce qui montre que $f(x) < f(y)$. Ainsi f est strictement croissante sur $\mathbb{R}$. $\square$

Dans l'exemple précédent 9.14, la fonction dérivée ne s'annule qu'en un point isolé. L'exemple qui suit montre que la fonction dérivée d'une fonction définie, dérivable et croissante peut s'annuler sur un ensemble contenant des points non isolés.

9.15. Fonction dérivable et strictement croissante sur $\mathbb{R}$ telle que l'ensemble des points où la fonction dérivée s'annule contient un point non isolé.

Nous considérons l'application :

$$\left| \begin{array}{l} \varphi : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto \varphi(x) = \begin{cases} 0 & \text{si } x = 0, \\ \left|x \sin \frac{1}{x}\right| & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

Pour tout réel x , $0 \leq \varphi(x) \leq |x|$, donc φ admet pour limite $0 = \varphi(0)$ en 0 ; par suite φ est continue en 0. Or φ est clairement continue en tout réel différent de

zéro, donc φ est continue sur $\mathbb{R}$. Par conséquent l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \int_0^x \varphi(t) dt \end{array} \right.$$

est dérivable sur $\mathbb{R}$ et, pour tout réel x , $f'(x) = \varphi(x) \geq 0$; en particulier, f est croissante sur $\mathbb{R}$. Nous introduisons la suite $(u_n)_{n \geq 1}$, de terme général $u_n = 1/(n\pi)$. On a $f'(0) = \varphi(0) = 0$, la suite (u_n) converge vers 0 et, pour tout entier $n \geq 1$, $u_n > 0$ et $f'(u_n) = \varphi(u_n) = 0$, donc 0 appartient à l'ensemble $\mathcal{A}$ des points où f' s'annule, mais 0 n'est pas isolé dans $\mathcal{A}$.

Soit a et b des réels tels que $0 \leq a < b$. Comme $\lim_{n \rightarrow \infty} u_n = 0$, l'ensemble des entiers $k \geq 1$ tels que $u_k < b$ n'est pas vide; nous notons p le plus petit de ces entiers. Si $u_p \leq a$, alors $f'(x) > 0$ pour tout point x de l'intérieur $]a, b[$ de l'intervalle $[a, b]$, donc f est strictement croissante sur $[a, b]$, d'où l'on déduit que $f(a) < f(b)$. Si $a < u_p$, alors $f'(x) > 0$ pour tout point x de l'intérieur $]u_p, b[$ de l'intervalle $[u_p, b]$, donc f est strictement croissante sur $[u_p, b]$, ce qui montre que $f(a) \leq f(u_p) < f(b)$. Il découle de cette étude que f est strictement croissante sur $[0, +\infty[$; or f est impaire et $f(0) = 0$, donc f est strictement croissante sur $\mathbb{R}$. $\square$

Une équivalence entre la croissance stricte sur un intervalle I d'une fonction, définie et continue sur I et dérivable sur l'intérieur de I , et une propriété de sa fonction dérivée est donnée par le théorème suivant.

THÉORÈME 9.2. — Si f est une fonction, définie et continue sur un intervalle I et dérivable sur l'intérieur I_0 de I , f est strictement croissante sur I si, et seulement si, la fonction dérivée f' est positive ou nulle sur I_0 et l'ensemble des points de I_0 en lesquels f' s'annule est d'intérieur vide.

Si a est un réel et f une fonction définie et dérivable sur un voisinage de 0, si f' est continue en 0 et si $f'(0) > 0$, il existe un réel $r > 0$ tel que f' est strictement positive sur l'intervalle ouvert $]a - r, a + r[$, donc f est strictement croissante sur le voisinage ouvert $V =]a - r, a + r[$ de a . Sans la continuité de f' en 0 on ne pourrait pas conclure.

9.16. Fonction f dérivable sur $\mathbb{R}$ telle que $f'(0) > 0$ alors que f n'est croissante sur aucun voisinage de 0.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ x^2 \sin \frac{1}{x} + \frac{x}{4} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

On déduit des résultats de l'exemple 9.9 (pages 166 et 167) que la fonction f est dérivable sur $\mathbb{R}$, que $f'(0) = 1/4$ et que, pour tout nombre réel $x \neq 0$:

$$(1) \quad f'(x) = \left(2x \sin \frac{1}{x} + \frac{1}{4} \right) - \cos \frac{1}{x}.$$

Nous introduisons les suites $(u_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$, de termes généraux :

$$u_n = \frac{1}{n\pi + \frac{\pi}{3}} \quad \text{et} \quad v_n = \frac{1}{n\pi}.$$

Les suites (u_n) et (v_n) convergent vers 0 et $0 < u_n < v_n$ pour tout entier $n \geq 1$. Nous choisissons un entier $N \geq 1$ tel que $v_n < 1/16$ pour tout entier $n \geq N$. Soit n un entier $\geq N$. Soit x un point de $]u_n, v_n[$. On a $n\pi < 1/x < n\pi + (\pi/3)$ donc $0 < (1/x) - n\pi < \pi/3$. Cosinus étant strictement décroissante sur $[0, \pi]$, on en déduit que $1/2 < \cos((1/x) - n\pi) < 1$, donc que $(-1)^n \cos(1/x) > 1/2$, ce qui montre que $\cos(1/x) > 0$ si n est pair, $\cos(1/x) < 0$ si n est impair et $|\cos(1/x)| > 1/2$. De plus $0 < x < v_n < 1/16$, donc :

$$\left| 2x \sin \frac{1}{x} + \frac{1}{4} \right| \leq 2|x| + \frac{1}{4} = 2x + \frac{1}{4} < \frac{1}{8} + \frac{1}{4} < \frac{1}{2}$$

et on déduit de (1) que $f'(x)$ et $\cos \frac{1}{x}$ sont différents de zéro et de signes opposés. Nous avons ainsi établi que, pour tout entier $n \geq N$, f est strictement croissante sur le segment $[u_n, v_n]$ si n est impair et strictement décroissante sur $[u_n, v_n]$ si n est pair. Ainsi $f'(0) > 0$ mais f n'est croissante sur aucun voisinage de 0. $\square$

■ Dérivées et limites

Si a est un réel et f une fonction définie, dérivable et uniformément continue sur l'intervalle $]a, +\infty[$, et si f admet une limite ℓ dans $\mathbb{R}$ en $+\infty$, alors f' tend vers 0 en $+\infty$. Ceci devient faux si f n'est pas uniformément continue.

9.17. Fonction dérivable sur $]0, +\infty[$ qui tend vers 0 en $+\infty$ mais dont la fonction dérivée n'admet pas de limite en $+\infty$.

Nous introduisons l'application :

$$\left| \begin{array}{l} f :]0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \frac{\sin(x^2)}{x}. \end{array} \right.$$

Pour tout réel $x > 0$, $|f(x)| \leq 1/x$, donc f admet pour limite 0 dans $\mathbb{R}$ en $+\infty$. Par ailleurs f est dérivable sur l'intervalle $]0, +\infty[$ et, pour tout réel $x > 0$:

$$f'(x) = 2 \cos(x^2) - \underbrace{\frac{\sin(x^2)}{x^2}}_{= g(x)}.$$

Pour tout $x > 0$, $|g(x)| \leq 1/x^2$, donc $g(x)$ tend vers 0 quand x tend vers $+\infty$. La suite $(u_n)_{n \geq 1}$, de terme général $u_n = \sqrt{n\pi}$, tend vers $+\infty$ et, pour tout $n \in \mathbb{N}^*$, $\cos((u_n)^2) = \cos(n\pi) = (-1)^n$, terme général d'une suite divergente, donc la suite $(f'(u_n))$ ne converge pas. La fonction f' n'admet donc pas de limite en $+\infty$. $\square$

Dans l'exemple précédent 9.17, la fonction dérivée de f est bornée ; dans l'exemple suivant elle ne l'est pas.

9.18. Fonction f dérivable sur $]0, +\infty[$ qui tend vers 0 en $+\infty$ et telle que $\liminf_{x \rightarrow +\infty} f'(x) = -\infty$ et $\limsup_{x \rightarrow +\infty} f'(x) = +\infty$.

Nous considérons l'application :

$$\left| \begin{array}{l} f :]0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \frac{\sin(x^3)}{x}. \end{array} \right.$$

Pour tout réel $x > 0$, $|f(x)| \leq 1/x$, donc f admet pour limite 0 dans $\mathbb{R}$ en $+\infty$. Par ailleurs f est dérivable sur l'intervalle $]0, +\infty[$ et, pour tout réel $x > 0$:

$$f'(x) = 3x \cos(x^3) - \frac{\sin(x^3)}{x^2}.$$

Nous introduisons les suites $(u_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$, de termes généraux :

$$u_n = \sqrt[3]{2n\pi} \text{ et } v_n = \sqrt[3]{(2n+1)\pi}.$$

Pour tout $n \in \mathbb{N}^*$, $\cos((u_n)^3) = 1$, $\cos((v_n)^3) = -1$ et $\sin((u_n)^3) = \sin((v_n)^3) = 0$ donc $f'(u_n) = 3u_n$ et $f'(v_n) = -3v_n$. Par conséquent la suite $(f'(u_n))$ tend vers $+\infty$ et $(f'(v_n))$ vers $-\infty$, donc $\liminf_{x \rightarrow +\infty} f'(x) = -\infty$ et $\limsup_{x \rightarrow +\infty} f'(x) = +\infty$. $\square$

Nous donnons maintenant un exemple où la fonction étudiée est décroissante.

9.19. Fonction dérivable et décroissante sur $[0, +\infty[$ qui tend vers 0 en $+\infty$ mais dont la fonction dérivée n'admet pas de limite en $+\infty$.

Nous posons $I = \left[-\frac{1}{2}, \frac{1}{2}\right]$ et nous considérons l'application :

$$\left| \begin{array}{l} g : I \longrightarrow I \\ x \longmapsto g(x) = \frac{1}{2} \sin(\pi x). \end{array} \right.$$

La fonction g est dérivable sur I et, pour tout point x de I , $g'(x) = \frac{\pi}{2} \cos(\pi x)$.

Par suite $g'(x) > 0$ pour tout point x de l'intérieur $] -1/2, 1/2[$ de I , donc g est strictement croissante sur I . La fonction g est continue sur I , $g(-1/2) = -1/2$ et $g(1/2) = 1/2$, donc $g(I) = I$, ce qui montre que g est une bijection de I sur I .

On définit la suite $(g_n)_{n \geq 1}$ d'applications de I dans I , par récurrence sur $n \in \mathbb{N}^*$, par $g_1 = g$ et, pour tout entier $n \geq 1$, $g_{n+1} = g_n \circ g$.

Un récurrence immédiate montre que, pour tout entier $n \geq 1$, g_n est une bijection strictement croissante de I sur I et que g_n est dérivable sur I . Nous prouvons, par récurrence sur $n \in \mathbb{N}^*$, que, pour tout entier $n \geq 1$:

$$g_n(0) = 0, \quad g_n\left(-\frac{1}{2}\right) = -\frac{1}{2}, \quad g_n\left(\frac{1}{2}\right) = \frac{1}{2}, \quad g'_n(0) = \frac{\pi^n}{2^n} \text{ et } g'_n\left(-\frac{1}{2}\right) = g'_n\left(\frac{1}{2}\right) = 0.$$

Comme $g_1 = g$ et que $g(0) = 0$, c'est vrai au rang 1. Soit $n \in \mathbb{N}^*$ et supposons la propriété vraie au rang n . Pour tout point x de I , $g_{n+1}(x) = g_n(g(x))$ donc $g'_{n+1}(x) = g'_n(g(x))g'(x)$. Posons $a = \pm 1/2$. On a $g(0) = 0$, $g(a) = a$, $g'(0) = \pi/2$ et $g'(a) = 0$, donc $g_{n+1}(0) = g_n(0) = 0$, $g_{n+1}(a) = g_n(a) = a$, $g'_{n+1}(a) = 0$ et :

$$g'_{n+1}(0) = \frac{\pi}{2} g'_n(0) = \frac{\pi}{2} \frac{\pi^n}{2^n} = \frac{\pi^{n+1}}{2^{n+1}},$$

ce qui achève le raisonnement par récurrence.

Nous posons $J_0 = \left[0, \frac{1}{2}\right[$ et, pour tout entier $n \geq 1$, $J_n = \left[n - \frac{1}{2}, n + \frac{1}{2}\right[$.

La famille $(J_n)_{n \in \mathbb{N}}$ est une partition de l'intervalle $\mathbb{R}_+ = [0, +\infty[$, ce qui justifie

l'existence de l'application :

$$\left| \begin{array}{l} f : [0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \left(\frac{2}{3}\right)^n \left(\frac{5}{2} - g_{n+1}(x-n)\right) \\ \text{où } n \text{ est l'unique entier} \\ \text{naturel tel que } x \in J_n. \end{array} \right.$$

Si n est un entier ≥ 1 , $((n-(1/2))-(n-1)) = 1/2$ donc, la fonction g_n étant continue sur le segment $I = [-1/2, 1/2]$:

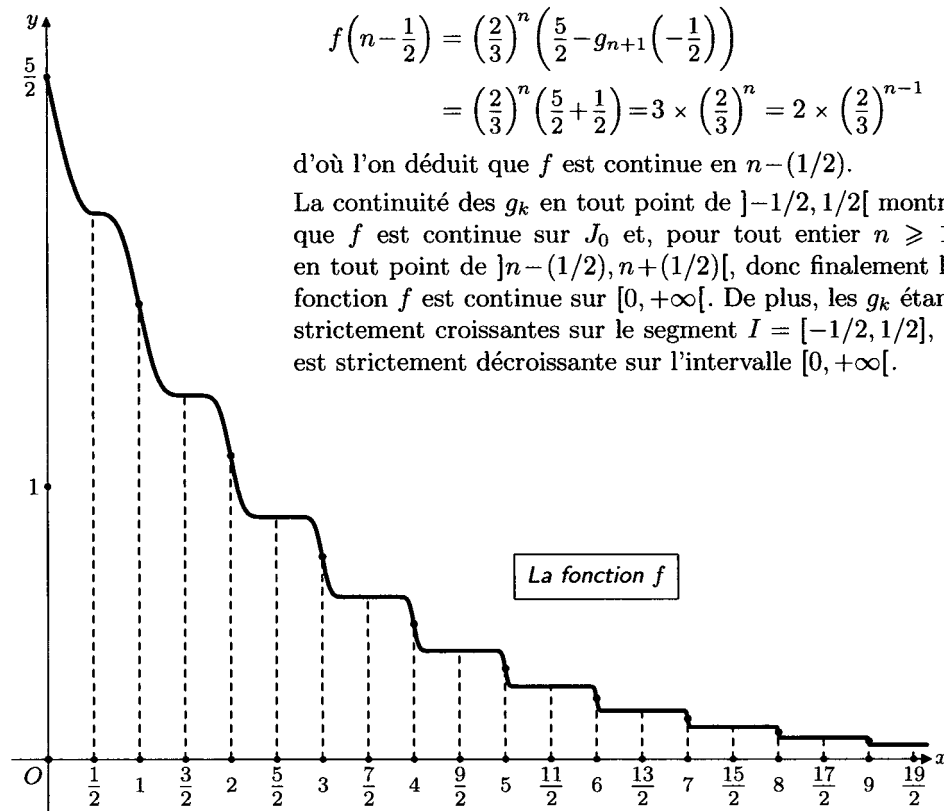
$$\lim_{x \rightarrow (n-(1/2))^-} f(x) = \left(\frac{2}{3}\right)^{n-1} \left(\frac{5}{2} - g_n\left(\frac{1}{2}\right)\right) = \left(\frac{2}{3}\right)^{n-1} \left(\frac{5}{2} - \frac{1}{2}\right) = 2 \times \left(\frac{2}{3}\right)^{n-1}$$

et $(n-(1/2))-n = -1/2$, ce qui donne :

$$\begin{aligned} f\left(n-\frac{1}{2}\right) &= \left(\frac{2}{3}\right)^n \left(\frac{5}{2} - g_{n+1}\left(-\frac{1}{2}\right)\right) \\ &= \left(\frac{2}{3}\right)^n \left(\frac{5}{2} + \frac{1}{2}\right) = 3 \times \left(\frac{2}{3}\right)^n = 2 \times \left(\frac{2}{3}\right)^{n-1} \end{aligned}$$

d'où l'on déduit que f est continue en $n-(1/2)$.

La continuité des g_k en tout point de $] -1/2, 1/2[$ montre que f est continue sur J_0 et, pour tout entier $n \geq 1$, en tout point de $]n-(1/2), n+(1/2)[$, donc finalement la fonction f est continue sur $[0, +\infty[$. De plus, les g_k étant strictement croissantes sur le segment $I = [-1/2, 1/2]$, f est strictement décroissante sur l'intervalle $[0, +\infty[$.



Pour tout $n \in \mathbb{N}^*$, on a, pour tout $x \in [n-1, (n-1)+(1/2)] = [n-1, (n-(1/2))]$, $f(x) = (2/3)^{n-1}((5/2) - g_n(x-(n-1)))$, et, pour tout $x \in [(n-(1/2), n+(1/2)]$, $f(x) = (2/3)^n((5/2) - g_{n+1}(x-n))$, donc, comme $g_n'(1/2) = g_{n+1}'(-1/2) = 0$, la fonction f est dérivable à gauche et à droite en $n-(1/2)$ et ses dérivées à gauche et à droite en $n-(1/2)$ valent 0, ce qui montre que f est dérivable en $n-(1/2)$ et que $f'(n-(1/2)) = 0$. De plus, f est dérivable sur $[0, 1/2[$ et, pour tout entier $n \geq 1$, sur $]n-(1/2), n+(1/2)[$, donc f est dérivable sur $[0, +\infty[$. On a, pour tout entier $n \geq 1$:

$$f(n) = \left(\frac{2}{3}\right)^n \left(\frac{5}{2} - g_{n+1}(0)\right) = \frac{5}{2} \left(\frac{2}{3}\right)^n$$

et :

$$f'(n) = -\left(\frac{2}{3}\right)^n g'_{n+1}(0) = -\left(\frac{2}{3}\right)^n \frac{\pi^{n+1}}{2^{n+1}} = -\frac{\pi}{2} \left(\frac{\pi}{3}\right)^n.$$

Pour tout entier $n \geq 1$, $f'(n - (1/2)) = 0$, donc la suite $(f'(n + (1/2)))$ converge vers 0. Or $0 < 2/3 < 1$ et $\pi/3 > 1$, donc la suite $(f(n))$ tend vers 0 et la suite $(f'(n))$ vers $-\infty$. Il en résulte que la fonction dérivée f' n'admet pas de limite dans $\mathbb{R}$ et ne tend ni vers $+\infty$ ni vers $-\infty$ en $+\infty$, et, f étant décroissante sur $\mathbb{R}$, que la fonction f tend vers 0 en $+\infty$. $\square$

Considérons de nouveau un réel a et une fonction f définie et dérivable sur $]a, +\infty[$. Dans le cas où f' tend vers 0 en $+\infty$, on pourrait s'attendre à ce que f admette une limite en $+\infty$. Il n'en est rien comme le montrent les deux exemples suivants.

9.20. Fonction f dérivable sur $]0, +\infty[$ telle que f tend vers $+\infty$ et f' vers 0 en $+\infty$.

La fonction $\sqrt{\cdot}$ est dérivable sur $]0, +\infty[$ et, pour tout réel $x > 0$, $\sqrt{\cdot}'(x) = \frac{1}{2\sqrt{x}}$, donc $\lim_{+\infty} \sqrt{\cdot} = +\infty$ et $\lim_{+\infty} \sqrt{\cdot}' = 0$. $\square$

9.21. Fonction f dérivable sur $\mathbb{R}$ telle que f' tend vers 0 en $+\infty$ alors que f n'admet pas de limite en $+\infty$.

Nous considérons l'application :

$$\left| \begin{array}{l} f :]0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \cos(\sqrt{x}). \end{array} \right.$$

La fonction f est dérivable sur $]0, +\infty[$ et, pour tout nombre réel $x > 0$:

$$f'(x) = -\frac{\sin(\sqrt{x})}{2\sqrt{x}}$$

donc $|f'(x)| \leq 1/(2\sqrt{x})$, ce qui montre que f' tend vers 0 en $+\infty$.

La suite $(u_n)_{n \geq 1}$ de terme général $u_n = n^2 \pi^2$ tend vers $+\infty$ et, pour tout entier $n \geq 1$, $f(u_n) = (-1)^n$, terme général d'une suite divergente, donc la fonction f n'admet pas de limite en $+\infty$. $\square$

■ Dérivées et extremums

DÉFINITION 9.3. — Si f est une fonction définie sur un intervalle I et a un point intérieur à I , f admet un maximum (resp. un minimum) relatif en a s'il existe un nombre réel $r > 0$ tel que $f(x) \leq f(a)$ (resp. $f(x) \geq f(a)$) pour tout point x de l'intervalle $]a - r, a + r[\cap I$.

DÉFINITION 9.4. — Si f est une fonction définie sur un intervalle I et a un point intérieur à I , f admet un extremum relatif en a si f admet un maximum ou un minimum relatif en a .

THÉORÈME 9.3. — Si f est une fonction définie sur un intervalle I et a un point intérieur à I , si f admet un extremum relatif en a et si la fonction f est dérivable en a , alors $f'(a) = 0$.

La réciproque est fausse comme le montre l'exemple suivant.

9.22. Fonction dérivable sur $\mathbb{R}$ dont la dérivée en 0 est nulle mais n'admettant pas d'extremum relatif en 0.

L'application $f : x \mapsto f(x) = x^3$ de $\mathbb{R}$ dans $\mathbb{R}$ est dérivable sur $\mathbb{R}$ et, pour tout nombre réel x , $f'(x) = 3x^2$; en particulier $f'(0) = 0$. On a, quels que soient les nombres réels x et y tels que $x < 0 < y$, $f(x) < f(0) = 0 < f(y)$, donc la fonction f n'admet pas d'extremum relatif en 0. $\square$

Une fonction f dérivable qui admet en un point a un minimum relatif n'est pas forcément croissante sur un voisinage de a à droite.

9.23. Fonction dérivable sur $\mathbb{R}$ qui admet un minimum relatif en 0 mais qui n'est croissante sur aucun voisinage de 0 à droite.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ x^2 \left(2 + \sin \frac{1}{x} \right) & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

En utilisant les résultats de l'exemple 9.9 (pages 166 et 167) concernant la fonction $g : x \mapsto g(x) = f(x) - 2x^2$, nous obtenons que la fonction f est dérivable sur $\mathbb{R}$, que $f'(0) = 0$ et que, pour tout nombre réel $x \neq 0$:

$$f'(x) = 4x + 2x \sin \frac{1}{x} - \cos \frac{1}{x}$$

Pour tout $x \in \mathbb{R} \setminus \{0\}$, $2 + \sin(1/x) \geq 2 - 1 = 1$ et $x^2 > 0$, donc $f(x) \geq x^2 > f(0)$, ce qui montre que f admet un minimum absolu strict en 0. À l'aide des suites $(u_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$, de termes généraux $u_n = 1/(n\pi + (\pi/3))$ et $v_n = 1/(n\pi)$, et d'un raisonnement analogue à celui de l'exemple 9.16 (pages 171 et 172), on prouve l'existence d'un entier $N \geq 1$ tel que, pour tout entier $n \geq N$, f est strictement croissante sur le segment $[u_n, v_n]$ si n est impair et strictement décroissante sur $[u_n, v_n]$ dans le cas où n est pair. Par conséquent f n'est croissante sur aucun voisinage de 0 à droite. $\square$

■ Fonctions indéfiniment dérivables

Soit Ω un ouvert de $\mathbb{R}$ et f une fonction définie sur Ω . On définit, par récurrence sur $n \in \mathbb{N}$, l'assertion « f est n fois dérivable sur Ω » et la fonction dérivée $f^{(n)}$ d'ordre n de f sur Ω de la manière suivante : f est 0 fois dérivable sur Ω si f est définie sur Ω , et alors $f^{(0)}$ est la restriction de f à Ω ; pour tout entier naturel n , f est $(n+1)$ fois dérivable sur Ω si f est n fois dérivable sur Ω et $f^{(n)}$ dérivable sur Ω , et dans ce cas on pose $f^{(n+1)} = (f^{(n)})'$. Si $n \in \mathbb{N}$, f est n fois continûment dérivable sur Ω , ou de classe $\mathcal{C}^n$ sur Ω , si f est n fois dérivable sur Ω et $f^{(n)}$ continue sur Ω . La fonction f est indéfiniment dérivable sur Ω , ou de classe $\mathcal{C}^\infty$ sur Ω si f est n fois dérivable sur Ω pour tout $n \in \mathbb{N}$ ou encore, ce qui est équivalent, si, pour tout entier naturel n , f est de classe $\mathcal{C}^n$ sur Ω . On note enfin, pour tout $n \in \mathbb{N} \cup \{\infty\}$, $\mathcal{C}^n(\Omega, \mathbb{R})$ l'espace vectoriel réel des applications de Ω dans $\mathbb{R}$ de classe $\mathcal{C}^n$ sur Ω .

Dans toutes ces définitions, on peut remplacer l'ouvert Ω par un intervalle.

Si a est un nombre réel et f une fonction définie sur un voisinage de a , et si n est un entier ≥ 1 , f est n fois dérivable en a si f est $(n-1)$ fois dérivable sur un voisinage ouvert de a et si $f^{(n-1)}$ est dérivable en a et, dans ce cas, $f^{(n)}(a) = (f^{(n-1)})'(a)$.

9.24. Fonction indéfiniment dérivable en 0 mais qui n'est indéfiniment dérivable sur aucun voisinage de 0.

Nous utilisons dans cet exemple les propriétés de la convergence uniforme des suites et séries de fonctions, rappelées aux chapitres 12 et 13.

Nous posons, pour tout entier $n \geq 1$ et tout réel x , $f_n(x) = x^{n-1}|x|$. On a donc, pour tout $n \in \mathbb{N}^*$, $f_n(x) = x^n$ pour tout $x \geq 0$ et $f_n(x) = -x^n$ pour tout $x \leq 0$. Pour tout entier naturel $k < n$, f_n est k fois dérivable sur $\mathbb{R}$ et, pour tout réel x :

$$f_n^{(k)}(x) = \begin{cases} \frac{n!}{(n-k)!} x^{n-k} & \text{si } x > 0, \\ 0 & \text{si } x = 0, \\ -\frac{n!}{(n-k)!} x^{n-k} & \text{si } x < 0. \end{cases}$$

Par contre f_n n'est n fois dérivable que sur l'ouvert $\mathbb{R} \setminus \{0\}$ et la fonction f_n est n fois dérivable à droite et à gauche en 0, $f_n^{(n)}(0) = n!$ et $f_n^{(n)}(0) = -n!$, donc f_n n'est pas n fois dérivable en 0.

Nous associons à tout entier $n \geq 1$ l'application :

$$\left| \begin{array}{l} g_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto g_n(x) = f_n\left(x - \frac{1}{n}\right) + f_n\left(x + \frac{1}{n}\right). \end{array} \right.$$

Si n est entier naturel et si $n \geq 1$, la fonction g_n est n fois dérivable sur l'ouvert $\Omega_n =]-\infty, -1/n[\cup]-1/n, 1/n[\cup]1/n, +\infty[$ et $(n-1)$ fois dérivable en $-1/n$ et en $1/n$. Nous introduisons, pour tout entier $n \geq 1$, l'application :

$$\left| \begin{array}{l} u_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_n(x) = \frac{g_n(x)}{n!}. \end{array} \right.$$

Soit $n \in \mathbb{N}^*$. Pour tout nombre réel y , $|f_n(y)| = |y|^n$ et, pour tout nombre réel x , $|x \pm (1/n)| \leq |x| + |1/n| \leq 1 + |x|$ donc :

$$|g_n(x)| \leq \left|x - \frac{1}{n}\right|^n + \left|x + \frac{1}{n}\right|^n \leq 2(1 + |x|)^n.$$

Par conséquent, pour tout réel $a > 0$, on a, pour tout $n \in \mathbb{N}^*$ et tout $x \in [-a, a]$:

$$|u_n(x)| \leq 2 \frac{(1+a)^n}{n!}.$$

Pour tout nombre réel $a > 0$, la série $\sum ((1+a)^n/n!)$ converge, donc la série de fonctions $\sum_{n \geq 1} u_n$ converge uniformément sur $[-a, a]$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} h : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto h(x) = \sum_{n=1}^{+\infty} u_n(x). \end{array} \right.$$

Soit $p \in \mathbb{N}^*$. Soit n un entier $> p$. On a $|f_n^{(p)}(y)| \leq (n!/(n-p)!)|y|^{n-p}$ pour tout $y \in \mathbb{R}$ et $g_n^{(p)}(x) = f_n^{(p)}(x - (1/n)) + f_n^{(p)}(x + (1/n))$ pour tout réel x . Pour tout $x \in]-1/p, 1/p[$, $|x \pm (1/n)| \leq 2$ donc $|u_n^{(p)}(x)| \leq 2 \times 2^{n-p}/(n-p)!$. Comme

la série $\sum_{n \geq p+1} 2^{n-p}/(n-p)!$ converge, la série de fonctions $\sum_{n \geq 1} u_n^{(p)}$ converge uniformément sur l'intervalle $] -1/p, 1/p[$. Par suite la fonction h est p fois dérivable sur $] -1/p, 1/p[$. Finalement, h est indéfiniment dérivable en 0.

Pour tout entier $p \geq 1$, f_p est, pour tout entier $n \neq p$, n fois dérivable sur un voisinage ouvert de $1/p$ et, comme $f_p^{(p)}(0) \neq f_p^{(p)}(0)$, les dérivées à droite et à gauche de h en $1/p$ diffèrent. Ainsi la fonction h n'est indéfiniment dérivable sur aucun voisinage de 0. $\square$

9.25. Fonction indéfiniment dérivable, strictement positive sauf en 0, et dont toutes les dérivées successives en 0 sont nulles.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ e^{-\frac{1}{x^2}} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

La fonction f est continue en tout réel $x \neq 0$ et, comme $\lim_{(x \rightarrow 0)} (-1/x^2) = -\infty$, $f(x)$ admet pour limite $0 = f(0)$ quand x tend vers 0, d'où l'on déduit que f est continue sur $\mathbb{R}$. La fonction f est de classe $\mathcal{C}^\infty$ sur l'ouvert $\Omega = \mathbb{R} \setminus \{0\}$ et, pour tout point x de Ω , $f'(x) = (2/x^3)e^{-1/x^2}$.

Nous prouvons, par récurrence sur $n \in \mathbb{N}^*$, que, pour tout entier $n \geq 1$, il existe une fonction polynomiale P_n telle que, pour tout point x de Ω :

$$f^{(n)}(x) = \frac{P_n(x)}{x^{2n+1}} e^{-\frac{1}{x^2}}.$$

C'est acquis au rang 1 pour la fonction polynomiale $P_1 : x \mapsto P_1(x) = 2$.

Soit $n \in \mathbb{N}^*$ et supposons la propriété vraie au rang n . On a, pour tout $x \in \Omega$:

$$f^{(n+1)}(x) = \left(\frac{1}{x^{2n+1}} P_n'(x) - \frac{2n+1}{x^{2n+2}} P_n(x) + \frac{2}{x^{2n+3}} P_n(x) \right) e^{-\frac{1}{x^2}} = \frac{P_{n+1}(x)}{x^{2n+3}} e^{-\frac{1}{x^2}}$$

pour la fonction $P_{n+1} : x \mapsto P_{n+1}(x) = x^2 P_n'(x) - (2n+1)x P_n(x) + 2P_n(x)$, qui est clairement une fonction polynomiale. Ceci achève le raisonnement par récurrence.

Soit $n \in \mathbb{N}^*$. On a, pour tout point x de Ω :

$$|f^{(n)}(x)| = |P_n(x)| \frac{1}{|x|^{2n+1}} e^{-\frac{1}{x^2}} = |P_n(x)| \frac{\left(\frac{1}{x^2}\right)^{n+\frac{1}{2}}}{e^{\frac{1}{x^2}}}.$$

Or P_n est continue en 0, $\lim_{(x \rightarrow 0)} (1/x^2) = +\infty$ et $\lim_{(y \rightarrow +\infty)} (y^{n+\frac{1}{2}}/e^y) = 0$, donc $f^{(n)}(x)$ tend vers 0 quand x tend vers 0.

Une récurrence facile montre alors que, pour tout $n \in \mathbb{N}^*$, f est n fois dérivable sur $\mathbb{R}$ et $f^{(n)}(0) = 0$. En conclusion, f est indéfiniment dérivable sur $\mathbb{R}$, $f(x) > 0$ pour tout réel $x \neq 0$ et les dérivées successives de f en 0 sont toutes nulles. $\square$

■ Développements limités

Pour les définitions concernant les développements limités, voir [ARN2], § VI.4.

Soit a un réel et f une fonction définie sur un voisinage de a . La fonction f admet un développement limité à l'ordre 1 au voisinage de a si, et seulement si, f est

dérivable en a . Si $n \geq 2$ et si f est n fois dérivable en a , f admet un développement limité à l'ordre n au voisinage de a donné par la formule de Taylor-Young :

$$f(x) = f(a) + \sum_{k=1}^n \frac{f^{(k)}(a)}{k!} (x-a)^k + o_{x \rightarrow a}(x-a)^n,$$

mais la réciproque est fautive.

9.26. Fonction admettant un développement limité à l'ordre 2 au voisinage de 0 mais qui n'est pas deux fois dérivable en 0.

Nous introduisons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ x^3 \sin \frac{1}{x} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

La fonction f est dérivable en tout point de $\mathbb{R} \setminus \{0\}$ et, pour tout réel x non nul :

$$(1) \quad f'(x) = 3x^2 \sin \frac{1}{x} - x \cos \frac{1}{x}.$$

Pour tout réel $x \neq 0$, $|f(x)| \leq x^3$ et $Qf(0, x) = f(x)/x$, donc $|Qf(0, x)| \leq x^2$, ce qui montre que $Qf(0, \cdot)$ admet pour limite 0 en 0 : f est dérivable en 0 et $f'(0) = 0$. Ainsi f est dérivable sur $\mathbb{R}$. On a, pour tout nombre réel $x \neq 0$:

$$(2) \quad Qf'(0, x) = \frac{f'(x)}{x} = 3x \sin \frac{1}{x} - \cos \frac{1}{x}.$$

Pour tout $x \in \mathbb{R} \setminus \{0\}$, $|x \sin(1/x)| \leq |x|$, donc $x \sin(1/x)$ admet pour limite 0 quand x tend vers 0. On déduit donc de (2) que si f était deux fois dérivable en 0, $\cos(1/x)$ admettrait pour limite $-f''(0)$ quand x tend vers 0, en contradiction avec ce que l'on sait de la fonction $x \mapsto \cos(1/x)$: elle n'admet pas de limite en 0 — pour le prouver, on utilise par exemple la suite $(1/(n\pi))$. Comme f , la fonction ω , définie par $\omega(0) = 0$ et $\omega(x) = x \sin(1/x)$ pour $x \neq 0$, admet pour limite 0 en 0. On a $f(x) = x^2 \omega(x) = 0 + 0x + 0x^2 + o_{(x \rightarrow 0)}(x^2)$, donc f admet un développement limité à l'ordre 2 au voisinage de 0 dont la partie régulière est le polynôme nul. $\square$

9.27. Fonction admettant, pour tout $n \in \mathbb{N}$, un développement limité à l'ordre n au voisinage de 0, mais qui n'est dérivable sur aucun voisinage de 0.

Nous reprenons la fonction f de l'exemple 9.25 (page 178) :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ e^{-\frac{1}{x^2}} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

Nous avons vu que f est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$ et que $f^{(n)}(0) = 0$ pour tout $n \in \mathbb{N}^*$. Si n un entier naturel on a, pour tout point x de l'ouvert $\Omega = \mathbb{R} \setminus \{0\}$:

$$\left| \frac{f(x)}{x^n} \right| = \frac{1}{|x|^n} e^{-\frac{1}{x^2}} = \frac{\left(\frac{1}{x^2}\right)^{\frac{n}{2}}}{e^{\frac{1}{x^2}}}$$

ce qui, puisque $\lim_{(x \rightarrow 0)} (1/x^2) = +\infty$ et $\lim_{(y \rightarrow +\infty)} (y^{\frac{n}{2}}/e^y) = 0$, montre que $f(x)/x^n$ admet pour limite 0 quand x tend vers 0.

Nous déduisons de f l'application :

$$\left| \begin{array}{l} g : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} f(x) & \text{si } x \in \mathbb{Q}, \\ 0 & \text{si } x \in \mathbb{R} \setminus \mathbb{Q}. \end{cases} \end{array} \right.$$

Pour tout réel x , $|g(x)| \leq |f(x)|$, donc, pour tout $n \in \mathbb{N}$, $\omega_n(x) = g(x)/x^n$ admet pour limite 0 quand x tend vers 0, ce qui s'écrit $g(x) = o_{(x \rightarrow 0)}(x^n)$, d'où l'on déduit que g admet un développement limité à l'ordre n au voisinage de 0 dont la partie régulière est le polynôme nul.

Soit a un rationnel différent de zéro. Nous posons $\varepsilon_0 = g(a) = f(a)$; alors $\varepsilon_0 > 0$. Pour tout nombre réel $\eta > 0$, l'intervalle $]a - \eta, a + \eta[$ contient au moins un irrationnel x , pour lequel $|a - x| < \eta$ et $|g(x) - g(a)| = |g(a)| = g(a) \geq \varepsilon_0$. Il en résulte que g est discontinue en a .

Tout voisinage de 0 contenant des rationnels non nuls, g n'est continue sur aucun voisinage ouvert de 0, donc g n'est dérivable sur aucun voisinage ouvert de 0. $\square$

Si une fonction paire admet un développement limité à l'ordre n au voisinage de 0, les coefficients d'indice impair de la partie régulière de ce développement limité sont tous nuls. Cependant la réciproque est fausse.

9.28. Fonction admettant, pour tout $n \in \mathbb{N}$, un développement limité à l'ordre n au voisinage de 0 tel que les coefficients d'indice impair de sa partie régulière sont tous nuls, alors qu'elle n'est pas paire.

Nous utilisons une nouvelle fois la fonction f des exemples 9.25 (page 178) et 9.27 (pages 179 et 180) et nous introduisons l'application :

$$\left| \begin{array}{l} g : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} f(x) + \cos x & \text{si } x > 0, \\ \cos x & \text{si } x \leq 0 \end{cases} \end{array} \right.$$

Comme $f(0) = 0$, la fonction g est continue sur $\mathbb{R}$, $g(0) = 1$ et g est indéfiniment dérivable sur l'ouvert $\Omega = \mathbb{R} \setminus \{0\}$.

Soit $n \in \mathbb{N}^*$. Pour tout réel x , $g^{(n)}(x) = f^{(n)}(x) + \cos^{(n)}(x) = f^{(n)}(x) + \cos(x + \frac{n\pi}{2})$ si $x > 0$ et $g^{(n)}(x) = \cos^{(n)}(x) = \cos(x + \frac{n\pi}{2})$ si $x < 0$. Nous avons vu dans l'exemple 9.25 que $f^{(n)}$ admet pour limite 0 en 0, donc $g^{(n)}(x)$ admet pour limite $\cos \frac{n\pi}{2}$ quand x tend vers 0. On en déduit, comme pour la fonction f , que g est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$ et que, pour tout entier $n \geq 1$, $g^{(n)}(0) = \cos \frac{n\pi}{2}$, ce qui montre que $g^{(n)}(0) = 0$ si n est impair et $g^{(n)}(0) = (-1)^{n/2}$ si n est pair. Par la formule de Taylor-Young, g admet, pour tout entier naturel n , un développement limité à l'ordre n au voisinage de 0, donné par :

$$g(x) = 1 + \sum_{k=1}^n \frac{g^{(k)}(0)}{k!} x^k + o_{x \rightarrow 0}(x^n),$$

donc les coefficients d'indice impair de la partie régulière de ce développement limité sont tous nuls. Or, pour tout réel $x \neq 0$, $g(x) - g(-x) = e^{-1/x^2} > 0$, donc g n'est pas une fonction paire. $\square$

Remarquons que la fonction g de l'exemple précédent 9.28 et la fonction cosinus ont des développements limités identiques à tout ordre au voisinage de 0 bien qu'elles diffèrent sur tout voisinage de 0.

■ Dérivées supérieures et inférieures

Nous considérons une fonction f définie sur un voisinage d'un réel a . Il existe alors un nombre réel $r > 0$ tel que f est définie sur l'intervalle ouvert $]a - r, a + r[$. Pour tout $\varepsilon \in]0, r[$, les bornes supérieure et inférieure :

$$\sup_{0 < |x-a| < \varepsilon} \frac{f(x) - f(a)}{x - a} \quad \text{et} \quad \inf_{0 < |x-a| < \varepsilon} \frac{f(x) - f(a)}{x - a}$$

existent dans la droite numérique achevée $\overline{\mathbb{R}}$, et les fonctions :

$$\varepsilon \mapsto \sup_{0 < |x-a| < \varepsilon} \frac{f(x) - f(a)}{x - a} \quad \text{et} \quad \varepsilon \mapsto \inf_{0 < |x-a| < \varepsilon} \frac{f(x) - f(a)}{x - a}$$

sont respectivement croissante et décroissante sur $]0, r[$, d'où l'existence dans $\overline{\mathbb{R}}$ des limites, appelées respectivement la dérivée supérieure et la dérivée inférieure⁵ de f en a :

$$L(f, a) = \lim_{\varepsilon \rightarrow 0^+} \left(\sup_{0 < |x-a| < \varepsilon} \frac{f(x) - f(a)}{x - a} \right) \quad \text{et} \quad \ell(f, a) = \lim_{\varepsilon \rightarrow 0^+} \left(\inf_{0 < |x-a| < \varepsilon} \frac{f(x) - f(a)}{x - a} \right).$$

La fonction f est dérivable en a si, et seulement si, la dérivée supérieure $L(f, a)$ et inférieure $\ell(f, a)$ sont finies et égales, et, si la fonction f est dérivable en a , alors $L(f, a) = \ell(f, a) = f'(a)$. La différence entre $L(f, a)$ et $\ell(f, a)$ exprime l'irrégularité du quotient $Qf(a, x)$ quand x tend vers a .

9.29. Fonction continue en 0 telle que $L(f, 0) = +\infty$ et $\ell(f, 0) = -\infty$.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ \sqrt{|x|} \sin \frac{1}{x} & \text{si } x \neq 0 \end{cases} \end{array} \right.$$

Pour tout nombre réel x , $|f(x)| \leq \sqrt{|x|}$, donc f admet pour limite $0 = f(0)$ en 0. Par suite f est continue en 0.

Nous introduisons les suites $(u_n)_{n \geq 0}$ et $(v_n)_{n \geq 0}$, de termes généraux :

$$u_n = \frac{1}{\frac{\pi}{2} + 2n\pi} \quad \text{et} \quad v_n = \frac{1}{\frac{3\pi}{2} + 2n\pi}.$$

Pour tout entier naturel n , $\sin(1/u_n) = 1$, donc :

$$Qf(0, u_n) = \frac{f(u_n) - f(0)}{u_n - 0} = \frac{f(u_n)}{u_n} = \sqrt{\frac{\pi}{2} + 2n\pi},$$

ce qui montre que la suite de terme général $Qf(0, u_n)$ tend vers $+\infty$. On prouve de même que la suite de terme général $Qf(0, v_n)$ tend vers $-\infty$. Or les suites (u_n) et (v_n) sont à valeurs dans $]0, +\infty[$ et convergent vers 0, donc :

$$L(f, 0) = +\infty \quad \text{et} \quad \ell(f, 0) = -\infty. \quad \square$$

5. Les dérivées supérieure et inférieure d'une fonction en un point sont évoqués dès 1870 par Paul du Bois-Reymond ; il distinguait, contrairement à ce que nous faisons ici, les limites à droite et à gauche. Cependant la première étude précise en est faite en 1878 par le mathématicien italien Ulisse Dini dans son ouvrage *Fondamenti per la teoria delle funzioni di variabili reali*.

L'exemple 9.6 (pages 161 et 162) nous a fourni une fonction continue sur $\mathbb{R}$ qui n'est dérivable en aucun point. De manière analogue nous construisons une fonction continue sur $\mathbb{R}$ telle que $L(f, a) = +\infty$ et $\ell(f, a) = -\infty$ pour tout réel a .

9.30. Fonction continue sur $\mathbb{R}$ telle que, pour tout nombre réel a , $L(f, a) = +\infty$ et $\ell(f, a) = -\infty$.

Nous utilisons dans cet exemple les propriétés de la convergence uniforme des suites et séries de fonctions, rappelées aux chapitres 12 et 13.

Nous considérons la fonction $g : x \mapsto g(x) = d(x, \mathbb{Z})$ (distance de x à $\mathbb{Z}$) introduite dans l'exemple 9.6 (pages 161 et 162) et représentée page 161. La fonction g est lipschitzienne de rapport 1 sur $\mathbb{R}$ et on a, pour tout réel x , $0 \leq g(x) \leq 1/2$.

Nous associons à tout entier naturel n l'application :

$$\left| \begin{array}{l} u_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_n(x) = \frac{g(12^n x)}{2^n}, \end{array} \right.$$

lipschitzienne de rapport 6^n donc continue sur $\mathbb{R}$. On a, pour tout entier naturel n et tout nombre réel x , $|u_n(x)| \leq 1/2^{n+1}$, donc la série de fonctions $\sum_n u_n$ converge normalement sur $\mathbb{R}$, donc uniformément sur $\mathbb{R}$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \sum_{n=0}^{+\infty} u_n(x) \end{array} \right.$$

et montre que f est continue sur $\mathbb{R}$.

Soit a un nombre réel. Soit δ et μ des réels strictement positifs. Nous choisissons un entier naturel m tel que $1/12^m < \delta$ et $6^{m-2} > \mu$. La propriété d'Archimède nous fournit un entier relatif k tel que :

$$\frac{k}{2 \times 12^m} \leq a < \frac{k+1}{2 \times 12^m}.$$

Nous posons $x_i = \frac{k+i-2}{2 \times 12^m}$ pour $i = 1, 2, 3$ et 4. On a alors $x_1 < x_2 \leq a < x_3 < x_4$.

Soit $i \in \{1, 2, 3, 4\}$. Comme $g(p) = 0$ si $p \in \mathbb{Z}$, on obtient, pour tout entier $j > m$:

$$u_j(x_i) = \frac{1}{2^j} g\left(\frac{12^j(k+i-2)}{2 \times 12^m}\right) = 0$$

donc :

$$f(x_i) = \sum_{n=0}^m u_n(x_i).$$

Pour tout $p \in \mathbb{Z}$, $g(p + (1/2)) = 1/2$, donc :

$$u_m(x_i) = \frac{1}{2^m} g\left(\frac{12^m(k+i-2)}{2 \times 12^m}\right) = \frac{1}{2^m} g\left(\frac{k+i-2}{2}\right) = \begin{cases} 0 & \text{si } k+i \text{ est pair,} \\ \frac{1}{2^{m+1}} & \text{si } k+i \text{ est impair.} \end{cases}$$

On a :

$$|u_m(x_3) - u_m(x_2)| = \frac{1}{2^{m+1}} = |u_m(x_4) - u_m(x_1)| \text{ et } x_3 - x_2 = \frac{1}{2 \times 12^m}$$

donc $|u_m(x_3) - u_m(x_2)| = 6^m(x_3 - x_2)$ et $|u_m(x_4) - u_m(x_1)| = (6^m/3)(x_4 - x_1)$.

De plus, pour $i, j \in \{1, 2, 3, 4\}$:

$$\sum_{n=0}^{m-1} |u_n(x_i) - u_n(x_j)| \leq \sum_{n=0}^{m-1} 6^n |x_i - x_j| = \frac{6^m - 1}{6 - 1} |x_i - x_j| < \frac{6^m}{5} |x_i - x_j|.$$

On a $f(x_3) - f(x_2) = u_m(x_3) - u_m(x_2) + \sum_{n=0}^{m-1} (u_n(x_3) - u_n(x_2))$ donc :

$$\begin{aligned} |f(x_3) - f(x_2)| &\geq |u_m(x_3) - u_m(x_2)| - \sum_{n=0}^{m-1} |u_n(x_3) - u_n(x_2)| \\ &\geq 6^m(x_3 - x_2) - \frac{6^m}{5}(x_3 - x_2) = \frac{144}{5} 6^{m-2}(x_3 - x_2) > \mu(x_3 - x_2). \end{aligned}$$

De même :

$$|f(x_4) - f(x_1)| \geq \frac{6^m}{3}(x_4 - x_1) - \frac{6^m}{5}(x_4 - x_1) = \frac{24}{5} 6^{m-2}(x_4 - x_1) > \mu(x_4 - x_1).$$

D'autre part $f(x_4) - f(x_1)$ est du signe de $u_m(x_4) - u_m(x_1)$, donc positif si k est pair et négatif si k est impair ; de même $f(x_3) - f(x_2)$ et $u_m(x_3) - u_m(x_2)$ sont de même signe et de signe opposé à celui de $f(x_4) - f(x_1)$. Supposons par exemple que $f(x_4) - f(x_1) > 0$. On a :

$$\frac{a - x_1}{x_4 - x_1} \frac{f(x_4) - f(x_1)}{a - x_1} + \frac{x_4 - a}{x_4 - x_1} \frac{f(x_4) - f(a)}{x_4 - a} = \frac{f(x_4) - f(x_1)}{x_4 - x_1} > \mu.$$

Nous posons $\lambda = \frac{a - x_1}{x_4 - x_1}$. Alors $\lambda \in]0, 1[$, $1 - \lambda = \frac{x_4 - a}{x_4 - x_1}$ et $\frac{f(x_4) - f(x_1)}{x_4 - x_1}$ est compris entre

$$Q_1 = \frac{f(a) - f(x_1)}{a - x_1} \text{ et } Q_4 = \frac{f(x_4) - f(a)}{x_4 - a}$$

donc $Q_1 > \mu$ ou bien $Q_4 > \mu$. De même, puisque $f(x_3) - f(x_2) < 0$:

$$\frac{f(a) - f(x_2)}{a - x_2} < -\mu \text{ ou bien } \frac{f(x_3) - f(a)}{x_3 - a} < -\mu.$$

Ainsi on a trouvé, pour tout réel $\delta > 0$ et tout réel $\mu > 0$, des nombres réels y et z tels que $|y - a| < \delta$, $|z - a| < \delta$, $Qf(a, y) > \mu$ et $Qf(a, z) < -\mu$.

En conclusion, $L(f, a) = +\infty$ et $\ell(f, a) = -\infty$. $\square$

■ Equations différentielles

DÉFINITION 9.5. — Si f est une fonction définie et continue sur un ouvert Ω de $\mathbb{R}^2$ et si (x_0, y_0) est un point de Ω , on dit que l'équation différentielle :

$$(E) \mid y' = f(x, y)$$

vérifie le problème de Cauchy en (x_0, y_0) si elle admet une unique solution maximale⁶ φ définie sur un intervalle contenant x_0 telle que $\varphi(x_0) = y_0$.

THÉORÈME 9.4. — **Théorème de Cauchy-Lipschitz.**

Si f est une fonction définie et de classe $\mathcal{C}^1$ sur un ouvert Ω de $\mathbb{R}^2$, l'équation différentielle :

$$(E) \mid y' = f(x, y)$$

vérifie le problème de Cauchy⁷ en tout point (x_0, y_0) de Ω .

6. Ceci signifie que cette solution ne peut être prolongée sur un intervalle strictement plus grand.

7. On peut affaiblir la condition : f est de classe $\mathcal{C}^1$. Cauchy démontre ce théorème en 1824

9.31. Equation différentielle pour laquelle il n'y a pas unicité du problème de Cauchy.

Nous considérons l'équation différentielle (E) $| y' = 3\sqrt[3]{y^2}$ et nous posons $x_0 = 0$ et $y_0 = 0$. Les quatre applications de $\mathbb{R}$ dans $\mathbb{R}$ suivantes :

- $f_1 : x \mapsto f_1(x) = x^3$,
- $f_2 : x \mapsto f_2(x) = 0$,
- $f_3 : x \mapsto f_3(x) = \begin{cases} 0 & \text{si } x < 0, \\ x^3 & \text{si } x \geq 0, \end{cases}$
- $f_4 : x \mapsto f_4(x) = \begin{cases} x^3 & \text{si } x \leq 0, \\ 0 & \text{si } x > 0, \end{cases}$

sont des solutions deux à deux distinctes de l'équation différentielle (E) sur $\mathbb{R}$ et on a $f_i(x_0) = y_0$ pour $i = 1, 2, 3$ et 4 . $\square$

L'exemple précédent 9.31 concerne une équation différentielle qui n'est pas linéaire. Pour une équation différentielle linéaire, on dispose du théorème suivant.

THÉORÈME 9.5. — Si a et b sont des fonctions définies et continues sur un intervalle I , l'équation différentielle linéaire :

$$(E) \mid y' = a(x)y + b(x)$$

vérifie le problème de Cauchy en tout point (x_0, y_0) de $I \times \mathbb{R}$.

Ceci devient faux si y' est multiplié par une fonction de x qui s'annule en au moins un point de I .

9.32. Equation différentielle linéaire qui ne vérifie pas le problème de Cauchy.

Nous considérons l'équation différentielle linéaire (E) $| xy' = y$.

Sur $I =]0, +\infty[$ et $]-\infty, 0[$ elle s'écrit (E) $| y' = y/x$ et on la résout par les méthodes traditionnelles ; on obtient que les solutions de (E) sur $]0, +\infty[$ sont les applications $f_1 : x \mapsto f_1(x) = \lambda_1 x$ de $]0, +\infty[$ dans $\mathbb{R}$ où λ_1 est une constante réelle et les solutions de (E) sur $]-\infty, 0[$ les applications $f_2 : x \mapsto f_2(x) = \lambda_2 x$ de $]-\infty, 0[$ dans $\mathbb{R}$ où λ_2 est une constante réelle.

Soit f une éventuelle solution de (E) sur $\mathbb{R}$. Alors f est une application de $\mathbb{R}$ dans $\mathbb{R}$ dérivable sur $\mathbb{R}$, sa restriction à $]0, +\infty[$ est une solution de (E) sur $]0, +\infty[$ et sa restriction à $]-\infty, 0[$ une solution de (E) sur $]-\infty, 0[$, donc il existe des réels λ_1 et λ_2 tels que $f(x) = \lambda_1 x$ pour tout $x > 0$ et $f(x) = \lambda_2 x$ pour tout $x < 0$. En substituant 0 à x dans l'égalité $xf'(x) = f(x)$, valable pour tout $x \in \mathbb{R}$, on obtient : $f(0) = 0$. Par suite $f(x) = \lambda_1 x$ pour tout $x \in [0, +\infty[$ et $f(x) = \lambda_2 x$ pour tout $x \in]-\infty, 0]$, donc $\lambda_1 = f'_d(0) = f'(0) = f'_g(0) = \lambda_2$. En posant $\lambda = \lambda_1 (= \lambda_2)$,

lorsque f et sa fonction dérivée partielle par rapport à y sont continues. Lipschitz en 1868 puis Picard en 1890 affaiblissent ces hypothèses.

λ est une constante réelle et $f(x) = \lambda x$ pour tout réel x . La réciproque étant évidente, les solutions de (E) sur $\mathbb{R}$ sont les applications $f : x \mapsto f(x) = \lambda x$ de $\mathbb{R}$ dans $\mathbb{R}$ où λ est une constante réelle. Ainsi toute solution f de (E) sur $\mathbb{R}$ vérifie $f(0) = 0$. Il n'y a donc pas unicité d'une solution f telle que $f(0) = 0$, et l'existence d'une solution f telle que $f(0) = 1$ est mise en défaut. L'équation (E) ne vérifie donc pas le problème de Cauchy quand on en recherche des solutions sur $\mathbb{R}$. $\square$

THÉORÈME 9.6. — Une équation différentielle linéaire du premier ordre :

$$(E) \mid y' = a(x)y + b(x)$$

où les fonctions a et b sont continues, admet toujours des solutions maximales et l'ensemble de ces solutions est une droite affine réelle.

Ceci devient faux si y' est multiplié par une fonction de x qui s'annule en au moins un point.

9.33. Equation différentielle linéaire du premier ordre dont l'ensemble des solutions sur $\mathbb{R}$ est un plan affine réel.

Nous considérons l'équation différentielle (E) $\mid xy' - 3y = 2x^5$.

Sur $I =]-\infty, 0[$ ou $]0, +\infty[$ elle s'écrit (E) $\mid y' = (3/x)y + 2x^4$ et on la résout par les méthodes traditionnelles ; on voit que les solutions de (E) sur I sont les applications $f : x \mapsto f(x) = x^5 + \lambda x^3$ de I dans $\mathbb{R}$ où λ est une constante réelle.

Soit f une éventuelle solution de (E) sur $\mathbb{R}$. C'est une application de $\mathbb{R}$ dans $\mathbb{R}$ dérivable sur $\mathbb{R}$ telle que $xf'(x) - 3f(x) = 2x^5$ pour tout réel x , ce qui, en substituant 0 à x , montre que $f(0) = 0$; de plus la restriction de f à $I =]-\infty, 0[$ et $I =]0, +\infty[$ est une solution de (E) sur I . Il existe donc des réels λ et μ tels que f est l'application de $\mathbb{R}$ dans $\mathbb{R}$:

$$(1) \quad f : x \mapsto f(x) = \begin{cases} x^5 + \lambda x^3 & \text{si } x < 0, \\ x^5 + \mu x^3 & \text{si } x \geq 0. \end{cases}$$

Soit f une telle fonction. Alors $f(x) = x^5 + \lambda x^3$ pour tout point x de $]-\infty, 0[$ et $f(x) = x^5 + \mu x^3$ pour tout $x \in [0, +\infty[$. Ainsi f est dérivable sur $]-\infty, 0[$ et sur $[0, +\infty[$, $f'(x) = 5x^4 + 3\lambda x^2$ pour tout $x \in]-\infty, 0[$ et $f'(x) = 5x^4 + 3\mu x^2$ pour tout $x \in [0, +\infty[$, donc f est dérivable à gauche et à droite en 0 et $f'_g(0) = f'_d(0) = 0$. Par suite, f est dérivable sur $\mathbb{R}$ et $f'(0) = 0$. Or $0 \times f'(0) - 3f(0) = 0 = 2 \times 0^5$ et ce qui précède montre que la restriction de f à $I =]-\infty, 0[$ et $I =]0, +\infty[$ est une solution de (E) sur I , donc f est une solution de (E) sur $\mathbb{R}$.

En conclusion, les solutions de (E) sur $\mathbb{R}$ sont les applications f de $\mathbb{R}$ dans $\mathbb{R}$ définies par (1) où λ et μ sont des constantes réelles. En définissant les applications :

$$\varphi : x \mapsto \varphi(x) = \begin{cases} x^3 & \text{si } x \leq 0, \\ 0 & \text{si } x > 0 \end{cases} \quad \text{et} \quad \psi : x \mapsto \psi(x) = \begin{cases} 0 & \text{si } x < 0, \\ x^3 & \text{si } x \geq 0 \end{cases}$$

de $\mathbb{R}$ dans $\mathbb{R}$ et l'application $g : x \mapsto g(x) = x^5$ de $\mathbb{R}$ dans $\mathbb{R}$, on voit facilement que (φ, ψ) est une famille libre de vecteurs de $\mathcal{C}^1(\mathbb{R}, \mathbb{R})$ et que l'ensemble des solutions de (E) sur $\mathbb{R}$ est $\mathcal{P} = \{g + \lambda\varphi + \mu\psi \mid \lambda, \mu \in \mathbb{R}\}$, c'est-à-dire le plan affine $\mathcal{P} = g + P$ où P est le plan vectoriel engendré par la famille (φ, ψ) . $\square$

9.34. Equation différentielle linéaire du premier ordre dont l'ensemble des solutions sur $\mathbb{R}$ est un singleton.

Nous travaillons sur l'équation différentielle (E) $| xy' + (x+1)y = x^2$.

Sur $I =]-\infty, 0[$ ou $]0, +\infty[$ elle s'écrit (E) $| y' + \frac{x+1}{x}y = x$.

On résout l'équation homogène associée sur $I =]-\infty, 0[$ ou $]0, +\infty[$ et on trouve que l'ensemble de ses solutions sur I est la droite vectorielle de $\mathcal{C}^1(I, \mathbb{R})$ engendrée par l'application :

$$\left| \begin{array}{l} \varphi : I \longrightarrow \mathbb{R} \\ x \longmapsto \varphi(x) = \frac{e^{-x}}{x} = \frac{1}{xe^x}. \end{array} \right.$$

On applique alors la méthode de variation de la constante, qui conduit au calcul d'une primitive sur I , par deux intégrations par parties ou la méthode des coefficients indéterminés :

$$G(x) = \int x^2 e^x dx = (x^2 - 2x + 2)e^x,$$

d'où la solution particulière $f_0 : x \mapsto f_0(x) = \varphi(x)G(x) = \frac{x^2 - 2x + 2}{x}$ de (E) sur I .

Les solutions de (E) sur I sont donc les applications de I dans $\mathbb{R}$:

$$f : x \mapsto f(x) = \frac{x^2 - 2x + 2 + \lambda e^{-x}}{x}$$

où λ est une constante réelle.

Soit f une éventuelle solution de (E) sur $\mathbb{R}$. Alors f est une application de $\mathbb{R}$ dans $\mathbb{R}$ dérivable sur $\mathbb{R}$ et, pour tout réel x , $xf'(x) + (x+1)f(x) = x^2$. En substituant 0 à x , on obtient : $f(0) = 0$. De plus la restriction de f à $I =]0, +\infty[$ et à $I =]-\infty, 0[$ est une solution de (E) sur I . Par suite il existe des réels λ et μ tels que :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} \frac{x^2 - 2x + 2 + \lambda e^{-x}}{x} & \text{si } x < 0, \\ 0 & \text{si } x = 0, \\ \frac{x^2 - 2x + 2 + \mu e^{-x}}{x} & \text{si } x > 0. \end{cases} \end{array} \right.$$

Si $\alpha \in \mathbb{R}$, $x^2 - 2x + 2 + \alpha e^{-x}$ admet pour limite $2 + \alpha$ quand x tend vers 0, donc, si $\lambda \neq -2$, $f(x)$ diverge vers $\pm\infty$ quand x tend vers 0 à gauche, et si $\mu \neq -2$, $f(x)$ diverge vers $\pm\infty$ quand x tend vers 0 à droite. La continuité de f en 0 impose que $\lambda = \mu = -2$, donc la seule solution possible de (E) sur $\mathbb{R}$ est l'application f de $\mathbb{R}$ dans $\mathbb{R}$:

$$(1) \quad f : x \mapsto f(x) = \begin{cases} \frac{x^2 - 2x + 2 - 2e^{-x}}{x} & \text{si } x \neq 0, \\ 0 & \text{si } x = 0 \end{cases}.$$

On obtient facilement que $x^2 - 2x + 2 - 2e^{-x} = o_{(x \rightarrow 0)}(x^2)$, d'où l'on déduit que $f(x) = o_{(x \rightarrow 0)}(x)$. Or $f(0) = 0$, donc f est dérivable en 0 et $f'(0) = 0$. Par suite f est dérivable sur $\mathbb{R}$. De plus $0 \times f'(0) + (0+1)f(0) = 0 = 0^2$ et ce qui précède montre que la restriction de f à $I =]-\infty, 0[$ et $I =]0, +\infty[$ est une solution de (E) sur I , donc f est une solution de (E) sur $\mathbb{R}$.

En conclusion, la fonction f définie par (1) est une solution de (E) sur $\mathbb{R}$ et c'est la seule : l'ensemble des solutions de (E) sur $\mathbb{R}$ est le singleton $\{f\}$. $\square$

9.35. Equation différentielle linéaire du premier ordre n'ayant pas de solution sur $\mathbb{R}$ et dont l'ensemble des solutions sur $] -\infty, 1[$ est un singleton.

Nous considérons l'équation différentielle (E) $| 2x(1-x)y' + (1-x)y = 1$.

Sur $I =]-\infty, 0[,]0, 1[$ ou $]1, +\infty[$, elle s'écrit : (E) $| y' + \frac{1}{2x}y = \frac{1}{2x(1-x)}$.

On résout l'équation homogène associée sur $I =]-\infty, 0[,]0, 1[$ ou $]1, +\infty[$, et on trouve que l'ensemble de ses solutions sur I est la droite vectorielle de $\mathcal{C}^1(I, \mathbb{R})$ engendrée par l'application :

$$\left| \begin{array}{l} \varphi : I \longrightarrow \mathbb{R} \\ x \longmapsto \varphi(x) = \frac{1}{\sqrt{|x|}}. \end{array} \right.$$

On applique alors la méthode de variation de la constante, qui conduit au calcul d'une primitive sur $I =]0, 1[$ ou $]1, +\infty[$:

$$G_1(x) = \int \frac{\sqrt{x}}{2x(1-x)} dx = \int \frac{1}{1-(\sqrt{x})^2} \frac{1}{2\sqrt{x}} dx = \frac{1}{2} \ln \left| \frac{1+\sqrt{x}}{1-\sqrt{x}} \right|$$

et d'une primitive sur $I =]-\infty, 0[$:

$$G_2(x) = \int \frac{\sqrt{-x}}{2x(1-x)} dx = \int \frac{1}{1+(\sqrt{-x})^2} \frac{-1}{2\sqrt{-x}} dx = \arctan \sqrt{-x}$$

et on en déduit :

- une solution particulière $f_1 : x \mapsto f_1(x) = \varphi(x)G_1(x) = \frac{1}{2\sqrt{x}} \ln \left| \frac{1+\sqrt{x}}{1-\sqrt{x}} \right|$ de (E) sur $]0, 1[$ ou $]1, +\infty[$;
- une solution particulière $f_2 : x \mapsto f_2(x) = \varphi(x)G_2(x) = \frac{1}{\sqrt{-x}} \arctan \sqrt{-x}$ de (E) sur $] -\infty, 0[$.

L'ensemble des solutions de (E) sur $I =]-\infty, 0[$ ou $]0, 1[$ est donc la droite affine $\mathcal{D}_1$ de $\mathcal{C}^1(I, \mathbb{R})$ passant par f_1 et dirigée par φ et l'ensemble des solutions de (E) sur $I =]1, +\infty[$ la droite affine $\mathcal{D}_2$ de $\mathcal{C}^1(I, \mathbb{R})$ passant par f_2 et dirigée par φ , ce qui confirme l'affirmation du théorème 9.6 (page 185).

Si l'équation (E) admettait une solution f sur $\mathbb{R}$, on aurait, pour tout réel x , $2x(1-x)f'(x) + (1-x)f(x) = 1$, et la substitution de 1 à x donnerait l'égalité $0 = 1$, évidemment fausse. L'équation (E) n'admet donc pas de solution sur $\mathbb{R}$.

Il nous reste à chercher d'éventuelles solutions de (E) sur $] -\infty, 1[$. Pour ceci nous introduisons les applications :

$$\left| \begin{array}{l} \ell :]-1, 1[\longrightarrow \mathbb{R} \\ x \longmapsto \ell(x) = \frac{1}{2} \ln \frac{1+x}{1-x} \end{array} \right.$$

et :

$$\left| \begin{array}{l} g :]-\infty, 1[\longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} \frac{1}{\sqrt{-x}} \arctan \sqrt{-x} & \text{si } x < 0, \\ 1 & \text{si } x = 0, \\ \frac{1}{\sqrt{x}} \ell(\sqrt{x}) & \text{si } x > 0. \end{cases} \end{array} \right.$$

La fonction ℓ est de classe $\mathcal{C}^\infty$ sur $] -1, 1[$ et, pour tout point x de $] -1, 1[$, $\ell'(x) = 1/(1-x^2)$. Comme $\ell'(0) = 1$ et que $\ell(0) = 0$, $\ell(t) \sim_{(t \rightarrow 0)} t$, et nous savons que $\arctan t \sim_{(t \rightarrow 0)} t$, donc $\ell(t)/t$ et $(\arctan t)/t$ admettent pour limite 1 quand t tend vers 0. Par conséquent $g(x)$ admet pour limite 1 quand x tend vers 0, donc g est continue en 0 — c'est pour cela que l'on a choisi $g(0) = 1$.

Soit f une éventuelle solution de (E) sur l'intervalle $] -\infty, 1[$. On a alors, pour tout $x \in] -\infty, 1[$, $2x(1-x)f'(x) + (1-x)f(x) = 1$; en substituant 0 à x , on obtient l'égalité $f(0) = 1$. La restriction de f à $I =] -\infty, 0[$ et $I =] 0, 1[$ étant une solution de (E) sur I , il existe des réels λ et μ tels que, pour tout $x \in] -\infty, 1[$:

$$f(x) = g(x) + \frac{\lambda}{\sqrt{-x}} \text{ si } x < 0 \text{ et } f(x) = g(x) + \frac{\mu}{\sqrt{x}} \text{ si } x > 0.$$

La fonction g admet pour limite 1 en 0, donc, si $\lambda \neq 0$ ou $\mu \neq 0$, f tend vers $\pm\infty$ à gauche ou à droite en 0, en contradiction avec la continuité de f en 0; on a donc nécessairement $\lambda = \mu = 0$, ce qui montre que $f = g$. Il nous reste donc simplement à prouver que g est bien une solution de (E) sur $] -\infty, 1[$. En partant des développements limités à l'ordre 2 au voisinage de 0 des fonctions :

$$\arctan' : t \mapsto \arctan'(t) = \frac{1}{1+t^2} \text{ et } \ell' : t \mapsto \ell'(t) = \frac{1}{1-t^2}$$

on obtient les développements limités à l'ordre 3 au voisinage de 0 :

$$\arctan t = t - \frac{1}{3}t^3 + o_{t \rightarrow 0}(t^3) \text{ et } \ell(t) = t + \frac{1}{3}t^3 + o_{t \rightarrow 0}(t^3)$$

qui montrent que $\lim_{t \rightarrow 0} \frac{\arctan t - t}{t^3} = -\frac{1}{3}$ et $\lim_{t \rightarrow 0} \frac{\ell(t) - t}{t^3} = \frac{1}{3}$. Pour tout réel $x \neq 0$:

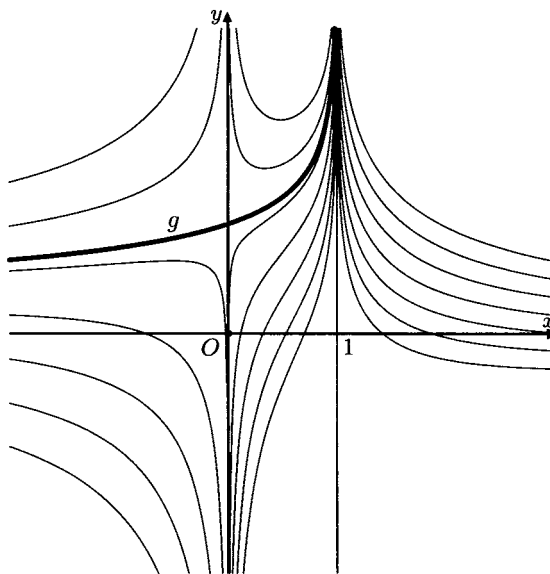
$$Qg(0, x) = \frac{g(x) - g(0)}{x - 0} = \frac{g(x) - 1}{x} = \begin{cases} -\frac{\arctan \sqrt{-x} - \sqrt{-x}}{(\sqrt{-x})^3} & \text{si } x < 0, \\ \frac{\ell(\sqrt{x}) - \sqrt{x}}{(\sqrt{x})^3} & \text{si } x > 0 \end{cases}$$

donc $Qg(0, x)$ admet pour limite $1/3$ quand x tend vers 0. Par suite g est dérivable en 0 et $g'(0) = 1/3$. Or g est dérivable en tout point de $] -\infty, 0[$ et en tout point de $] 0, 1[$, donc la fonction g est dérivable sur $] -\infty, 1[$. Enfin la restriction de g à $I =] -\infty, 0[$ et $I =] 0, 1[$ est une solution de (E) sur I et $2 \times 0 \times (1-0)g'(0) + (1-0)g(0)$ est égal à 1, donc g est une solution de (E) sur $] -\infty, 1[$.

Par conséquent la fonction g est l'unique solution de (E) sur l'intervalle $] -\infty, 1[$.

Sur le dessin ci-contre sont représentées plusieurs solutions

de (E) sur chacun des intervalles $] -\infty, 0[$, $] 0, 1[$ et $] 1, +\infty[$ et, en gras, l'unique solution g de l'équation différentielle (E) sur l'intervalle $] -\infty, 1[$. $\square$



THÉORÈME 9.7. — Une équation différentielle linéaire homogène d'ordre $n \in \mathbb{N}^*$:

$$(E) \mid y^{(n)} + a_{n-1}(x)y^{(n-1)} + \cdots + a_1(x)y' + a_0(x)y = 0$$

où $a_0, a_1, \dots, a_{n-1}$ sont continues, admet toujours des solutions maximales et l'ensemble de ces solutions est un espace vectoriel réel de dimension finie n .

Ce résultat, qui généralise le théorème 9.6 dans le cas où $b = 0$, devient faux si $y^{(n)}$ est multiplié par une fonction de x qui s'annule en au moins un point.

9.36. Equation différentielle linéaire d'ordre 2 dont l'ensemble des solutions sur $\mathbb{R}$ est un espace vectoriel de dimension 4.

Nous considérons l'équation différentielle $(E) \mid x^2y'' - 6xy' + 12y = 0$.

A l'aide de calculs faciles, on voit que les applications $g : x \mapsto g(x) = x^2$ et $h : x \mapsto h(x) = x^3$ de $\mathbb{R}$ dans $\mathbb{R}$ sont des solutions de (E) sur $\mathbb{R}$.

Sur les intervalles $I_1 =]0, +\infty[$ et $I_2 =]-\infty, 0[$, on peut diviser par x^2 et appliquer le théorème 9.7. Les solutions de (E) sur I_1 sont donc les combinaisons linéaires des restrictions $g|_{I_1}$ et $h|_{I_1}$ de g et de h à I_1 et celles de (E) sur I_2 les combinaisons linéaires des restrictions $g|_{I_2}$ et $h|_{I_2}$ de g et de h à I_2 ; il est clair en effet que la famille $(g|_{I_1}, h|_{I_1})$ est libre dans l'espace vectoriel réel $\mathcal{C}^2(I_1, \mathbb{R})$ et la famille $(g|_{I_2}, h|_{I_2})$ libre dans l'espace vectoriel réel $\mathcal{C}^2(I_2, \mathbb{R})$.

Soit f une solution de (E) sur $\mathbb{R}$ tout entier. La restriction de f à I_1 est une solution de (E) sur I_1 et sa restriction à I_2 une solution sur I_2 , donc il existe des constantes réelles α, β, λ et μ telles que $f(x) = \alpha x^2 + \beta x^3$ pour tout réel $x < 0$ et $f(x) = \lambda x^2 + \mu x^3$ pour tout réel $x > 0$, et en substituant 0 à x dans l'égalité $x^2 f''(x) - 6x f'(x) + 12f(x) = 0$, valable pour tout $x \in \mathbb{R}$, on obtient : $f(0) = 0$.

Si α, β, λ et μ sont des constantes réelles et f l'application de $\mathbb{R}$ dans $\mathbb{R}$:

$$(1) \quad f : x \mapsto f(x) = \begin{cases} \alpha x^2 + \beta x^3 & \text{si } x < 0, \\ \lambda x^2 + \mu x^3 & \text{si } x \geq 0, \end{cases}$$

alors $f(x) = \alpha x^2 + \beta x^3$ pour tout $x \in [0, +\infty[$ et $f(x) = \lambda x^2 + \mu x^3$ pour tout $x \in]-\infty, 0]$, ce qui montre que f est dérivable à gauche et à droite en 0 et que $f'_g(0) = 0$ et $f'_d(0) = 0$, donc f est dérivable en 0 et $f'(0) = 0$, d'où l'on déduit que f est dérivable sur $\mathbb{R}$ et, comme la restriction de f à I_1 (resp. à I_2) est une solution de (E) sur I_1 (resp. sur I_2), f est une solution de (E) sur $\mathbb{R}$.

Les solutions de (E) sur $\mathbb{R}$ sont donc les applications f de $\mathbb{R}$ dans $\mathbb{R}$ définies par (1) où α, β, λ et μ sont des constantes réelles. Nous introduisons les applications :

$$\left| \begin{array}{l} u_1 : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_1(x) = \begin{cases} x^2 & \text{si } x \leq 0, \\ 0 & \text{si } x > 0, \end{cases} \end{array} \right| \quad \left| \begin{array}{l} u_2 : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_2(x) = \begin{cases} x^3 & \text{si } x \leq 0, \\ 0 & \text{si } x > 0, \end{cases} \end{array} \right|$$

$$\left| \begin{array}{l} u_3 : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_3(x) = \begin{cases} 0 & \text{si } x < 0, \\ x^2 & \text{si } x \geq 0, \end{cases} \end{array} \right| \quad \left| \begin{array}{l} u_4 : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_4(x) = \begin{cases} 0 & \text{si } x < 0, \\ x^3 & \text{si } x \geq 0. \end{cases} \end{array} \right|$$

Clairement, (u_1, u_2, u_3, u_4) est une famille libre de solutions de (E) sur $\mathbb{R}$ et, quels que soient $\alpha, \beta, \lambda, \mu \in \mathbb{R}$, la fonction f définie par (1) est la combinaison linéaire $\alpha u_1 + \beta u_2 + \lambda u_3 + \mu u_4$. En conclusion, l'espace vectoriel réel des solutions de (E) sur $\mathbb{R}$ est de dimension finie 4. $\square$

Chapitre 10

Fonctions d'une variable réelle monotones, périodiques, convexes, bornées

Nous étudions dans ce chapitre des fonctions possédant des propriétés couramment utilisées : fonctions monotones, périodiques, convexes, bornées et à variation bornée. Les fonctions périodiques ont été introduites par l'étude de phénomènes physiques. Dans la première moitié du XIX^e siècle, Fourier puis Dirichlet ont travaillé à les exprimer sous la forme de séries trigonométriques. Les fonctions réglées et à variation bornée ont été étudiées principalement dans la deuxième moitié du XIX^e siècle, après les définitions rigoureuses des notions de continuité et de dérivabilité, pour trouver des conditions suffisantes d'intégrabilité.

Dans tout le chapitre, les fonctions sont des fonctions de $\mathbb{R}$ dans $\mathbb{R}$.

■ Monotonie et continuité

Une fonction monotone — c'est-à-dire croissante ou décroissante — n'a bien sûr aucune raison d'être continue ; cependant elle admet, en tout point en lequel ceci a un sens, une limite à droite et une limite à gauche. Ceci permet de démontrer que l'ensemble des points de discontinuité d'une fonction définie et monotone sur un intervalle est au plus dénombrable. Il est intéressant de construire une fonction strictement monotone f dont l'ensemble des points de discontinuité est dense ; de ce fait f n'est continue sur aucun intervalle d'intérieur non vide contenu dans son domaine de définition.

10.1. Fonction strictement croissante sur $[0, 1]$ discontinue en tous les points d'un ensemble dense dans $[0, 1]$.

Nous utilisons dans cet exemple les propriétés de la convergence uniforme des suites et séries de fonctions, rappelées aux chapitre 12 et 13.

L'ensemble $\mathcal{Q} = \mathbb{Q} \cap]0, 1[$ est dénombrable. Nous choisissons une bijection φ de $\mathbb{N}$ sur $\mathcal{Q}$, nous posons, pour tout point x de $[0, 1]$, $\mathcal{A}_x = \{n \in \mathbb{N} \mid \varphi(n) < x\}$, et nous

associons à tout entier naturel n l'application :

$$\left| \begin{array}{l} u_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto u_n(x) = \begin{cases} \frac{1}{2^n} & \text{si } n \in \mathcal{A}_x, \\ 0 & \text{sinon.} \end{cases} \end{array} \right.$$

Comme $0 \leq u_n(x) \leq 1/2^n$ pour tout entier naturel n et tout point x de $[0, 1]$, la série de fonctions $\sum_n u_n$ converge normalement, donc uniformément, sur le segment $[0, 1]$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \sum_{n=0}^{+\infty} u_n(x). \end{array} \right.$$

Montrons tout d'abord que la fonction f ainsi construite est strictement croissante. Soit x et y des points de $[0, 1]$ tels que $x < y$. Clairement, $\mathcal{A}_x \subset \mathcal{A}_y$, donc, si $n \in \mathcal{A}_x$, $u_n(x) = 1/2^n = u_n(y)$, alors que si $n \in \mathbb{N} \setminus \mathcal{A}_x$, $u_n(x) = 0$ et $u_n(y) \geq 0$; par suite, $u_n(x) \leq u_n(y)$ pour tout $n \in \mathbb{N}$. Nous choisissons un rationnel r tel que $x < r < y$; on a $0 < r < 1$, donc $r \in \mathbb{Q}$. Nous posons $q = \varphi^{-1}(r)$; q est l'unique entier naturel tel que $\varphi(q) = r$. Alors $x < \varphi(q)$ donc $q \notin \mathcal{A}_x$, et $\varphi(q) < y$ donc $q \in \mathcal{A}_y$, d'où l'on déduit que $0 = u_q(x) < 1/2^q = u_q(y)$; par suite $f(y) - f(x) \geq 1/2^q > 0$, donc $f(x) < f(y)$. La fonction f est donc strictement croissante sur $[0, 1]$.

Soit r un point de $\mathbb{Q}$. Alors $q = \varphi^{-1}(r)$ est l'unique entier naturel tel que $\varphi(q) = r$. On prouve comme ci-dessus que, pour tout point x du segment $[0, 1]$ tel que $x > r$, $f(x) - f(r) \geq 1/2^q$. Il en résulte que $f(x) - f(r)$ n'admet pas pour limite 0 quand x tend vers r à droite, donc que f n'est pas continue à droite en r . En conclusion, f est discontinue en tous les points de $\mathbb{Q}$, ensemble dense dans $[0, 1]$. $\square$

En complément nous démontrons que la fonction f de l'exemple précédent 10.1 est continue à gauche¹ en tous les points de $]0, 1]$.

Soit a un nombre réel tel que $0 < a \leq 1$. Soit n un entier naturel. Si $\varphi(n) \geq a$, alors, pour tout x du segment $[0, 1]$ tel que $x < a$, $\varphi(n) \geq a > x$ donc $u_n(x) = 0 = u_n(a)$. Si $\varphi(n) < a$, alors $\eta = a - \varphi(n)$ est un réel strictement positif et, pour tout point x de $[0, 1]$ tel que $a - \eta < x < a$, on a $\varphi(n) < x < a$ donc $u_n(x) = 1/2^n = u_n(a)$. Par conséquent, pour tout entier naturel n , la fonction u_n est continue à gauche en a . La convergence uniforme de la série de fonctions $\sum_n u_n$ sur $[0, 1]$ montre alors que la fonction f est continue à gauche en a . $\square$

Nous avons vu que la monotonie n'impose la continuité sur aucun intervalle ouvert. Démontrons qu'inversement, la continuité n'entraîne pas non plus la monotonie.

10.2. Fonction continue sur $\mathbb{R}$ qui n'est monotone sur aucun intervalle d'intérieur non vide.

Nous reprenons la fonction g , la suite $(u_n)_{n \geq 0}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$ et la fonction f définies dans l'exemple 9.6 (pages 161 et 162).

La fonction f est continue sur $\mathbb{R}$.

1. On obtient ainsi, comme dans l'exemple 8.13 (page 143), une fonction continue à gauche en tout point mais n'admettant pas de limite à droite en tous les points d'un ensemble dense.

$$a = \frac{k}{4^n}, \quad b = a + \frac{1}{4^{2n+1}} = \frac{4^{n+1}k + 1}{4^{2n+1}} \quad \text{et} \quad c = a - \frac{1}{4^{2n+1}} = \frac{4^{n+1}k - 1}{4^{2n+1}}.$$
$$f(b) - f(a) = \sum_{p=0}^{n-1} (u_p(b) - u_p(a)) + \sum_{p=n}^{2n} u_p(b)$$
$$\sum_{p=0}^{n-1} (u_p(b) - u_p(a)) \geq -n(b-a).$$
$$b - a = \frac{1}{4^{2n+1}} < \frac{1}{2 \times 4^{2n}} \leq \frac{1}{2 \times 4^p},$$
$$S_{n,k} = \left[\frac{k-1}{4^n}, \frac{k+1}{4^n} \right].$$
$$d - r < \frac{k-1}{4^n} < \frac{k}{4^n} \leq d < \frac{k+1}{4^n} < d + r.$$

Nous savons que la continuité d'une fonction sur un segment entraîne sa continuité uniforme mais pas sa continuité absolue, comme nous l'a montré l'exemple 8.16 (pages 145 et 146). Nous démontrons que si l'on ajoute la monotonie de la fonction, ceci n'entraîne toujours pas la continuité absolue.

10.3. Fonction définie, continue et croissante sur $[0, 1]$ qui n'est pas absolument continue sur $[0, 1]$.

Une telle fonction peut être obtenue directement, mais nous préférons la construire par étapes pour que le lecteur se familiarise mieux avec elle ; s'il connaît bien l'ensemble de Cantor, il fera le lien avec celui-ci (voir les exemples 5.30 pages 96 et 97 et 5.34 page 99).

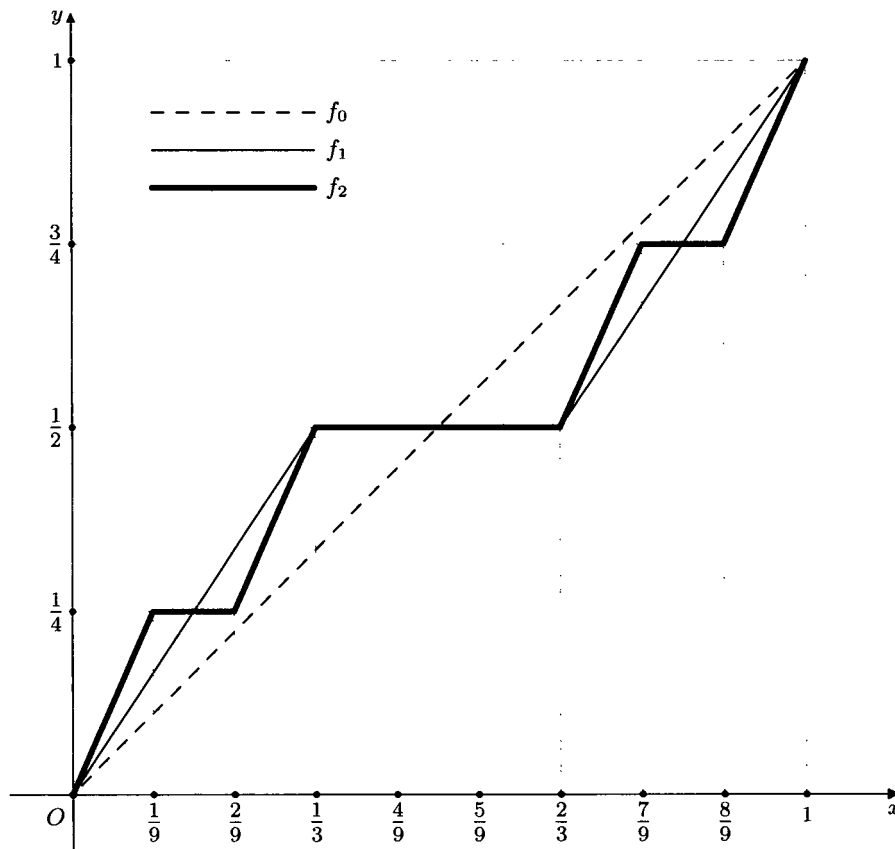
Nous construisons par récurrence la suite $(f_n)_{n \geq 0}$ d'applications de $[0, 1]$ dans $\mathbb{R}$ de la manière suivante :

- $f_0(x) = x$ pour tout point x de $[0, 1]$.
- Pour tout entier naturel n :

$$f_{n+1} : [0, 1] \longrightarrow \mathbb{R} \quad \left| \quad x \longmapsto f_{n+1}(x) = \begin{cases} \frac{1}{2} f_n(3x) & \text{si } 0 \leq x < \frac{1}{3}, \\ \frac{1}{2} & \text{si } \frac{1}{3} \leq x \leq \frac{2}{3}, \\ \frac{1}{2} f_n(3x-2) + \frac{1}{2} & \text{si } \frac{2}{3} < x \leq 1. \end{cases}$$

Nous avons ainsi, pour tout point x du segment $[0, 1]$:

$$f_1(x) = \frac{3x}{2} \text{ si } 0 \leq x < \frac{1}{3}, \quad f_1(x) = \frac{1}{2} \text{ si } \frac{1}{3} \leq x \leq \frac{2}{3} \text{ et } f_1(x) = \frac{3x}{2} - \frac{1}{2} \text{ si } \frac{2}{3} < x \leq 1.$$



Nous prouvons, par récurrence sur $n \in \mathbb{N}$, que, pour tout entier naturel n , f_n est continue sur $[0, 1]$, $f_n(0) = 0$ et $f_n(1) = 1$. On a $f_0(0) = 0$, $f_0(1) = 1$ et la fonction f_0 est évidemment continue sur $[0, 1]$. Soit $n \in \mathbb{N}$ et supposons f_n continue sur $[0, 1]$ et que $f_n(0) = 0$ et $f_n(1) = 1$. La définition de f_{n+1} montre que $f_{n+1}(0) = (1/2)f_n(0) = 0$ et $f_{n+1}(1) = (1/2)f_n(1) + (1/2) = (1/2) + (1/2) = 1$, que f_{n+1} est continue sur les intervalles $[0, 1/3[$, $[1/3, 2/3]$ et $]2/3, 1]$, et que :

$$\frac{1}{2}f_n\left(3\frac{1}{3}\right) = \frac{1}{2}f_n(1) = \frac{1}{2} = f_{n+1}\left(\frac{1}{3}\right)$$

et :

$$\frac{1}{2}f_n\left(3\frac{2}{3} - 2\right) + \frac{1}{2} = \frac{1}{2}f_n(0) + \frac{1}{2} = \frac{1}{2} = f_{n+1}\left(\frac{2}{3}\right),$$

d'où l'on déduit que f_{n+1} est continue sur $[0, 1]$. Ceci achève le raisonnement par récurrence. On a, pour tout point x du segment $[0, 1]$:

$$f_1(x) - f_0(x) = \begin{cases} \frac{x}{2} & \text{si } 0 \leq x < \frac{1}{3}, \\ \frac{1}{2} - x & \text{si } \frac{1}{3} \leq x \leq \frac{2}{3}, \\ \frac{x}{2} - \frac{1}{2} & \text{si } \frac{2}{3} < x \leq 1, \end{cases} \quad \text{donc } |f_1(x) - f_0(x)| \leq \frac{1}{6} = \frac{1}{3} \times \frac{1}{2}.$$

Nous démontrons, par récurrence sur $n \in \mathbb{N}^*$, que, pour tout entier $n \geq 1$:

$$|f_n(x) - f_{n-1}(x)| \leq \frac{1}{3} \times \frac{1}{2^n} \quad \text{pour tout } x \in [0, 1].$$

Le résultat est acquis au rang 1. Soit n un entier ≥ 1 et supposons la propriété vraie au rang n . Soit x un point de $[0, 1]$. Si x appartient à $[0, 1/3[$, on a :

$$|f_{n+1}(x) - f_n(x)| = \frac{1}{2}|f_n(3x) - f_{n-1}(3x)| \leq \frac{1}{2} \times \frac{1}{3} \times \frac{1}{2^n} = \frac{1}{3} \times \frac{1}{2^{n+1}}$$

et on opère de même pour les points x de $[1/3, 2/3]$ et ceux de $[2/3, 1]$. Ceci achève le raisonnement par récurrence.

Quels que soient $n \in \mathbb{N}$ et $p \in \mathbb{N}^*$, on a, pour tout point x du segment $[0, 1]$:

$$|f_{n+p}(x) - f_n(x)| \leq \sum_{k=1}^p \underbrace{|f_{n+k}(x) - f_{n+k-1}(x)|}_{\leq \frac{1}{3 \times 2^{n+k}}} \leq \frac{1}{3 \times 2^{n+1}} \times \underbrace{\sum_{k=1}^p \frac{1}{2^{k-1}}}_{\leq 2} \leq \frac{1}{3 \times 2^n}.$$

Ceci montre que (f_n) est une suite de Cauchy pour la norme de la convergence uniforme ; elle converge donc uniformément vers une application f de $[0, 1]$ dans $\mathbb{R}$, continue sur $[0, 1]$. Or f_n est, pour tout entier naturel n , une application croissante de $[0, 1]$ dans $\mathbb{R}$ et les inégalités larges « passent à la limite », donc f est une application croissante² de $[0, 1]$ dans $\mathbb{R}$. De plus, $f(0) = 0$ et $f(1) = 1$.

2. Cette fonction, connue sous le nom *d'escalier du diable*, a été introduite par Cantor. On peut aussi la définir de la manière suivante. Soit x un élément de $[0, 1[$ dont le développement illimité propre en base 3 est $x = 0, \alpha_1 \alpha_2 \alpha_3 \dots \alpha_n \dots$. Si aucun des α_i ne vaut 1, on note, pour tout entier $i \geq 1$, β_i le quotient exact de α_i par 2 et on pose, en base 2, $f(x) = 0, \beta_1 \beta_2 \beta_3 \dots \beta_n \dots$. Si l'un au moins des α_i est égal à 1, on note n le plus petit des entiers i tel que $\alpha_i = 1$ et, pour tout $i \in [1, n-1]$, β_i le quotient exact de α_i par 2 et on pose, en base 2, $f(x) = 0, \beta_1 \beta_2 \beta_3 \dots \beta_{n-1} 1$. Enfin, on pose $f(1) = 1$. Voir [HAIR], chapitre III, §6, exercice 6.5 et [BOUA], chapitre 23, §23.4 ; voir aussi l'exemple 11.20, page 223.

Nous démontrons enfin que f n'est pas absolument continue sur $[0, 1]$.

Nous notons, pour tout entier naturel n , A_n l'ensemble des entiers $k \in \llbracket 0, 3^n - 1 \rrbracket$ ne contenant aucun 1 dans leur écriture en base 3. On a ainsi $A_0 = \{0\}$, $A_1 = \{0, 2\}$ et $A_2 = \{0, 2, 6, 8\}$. Pour tout $n \in \mathbb{N}$, A_n est fini et de cardinal 2^n ; en effet, $\text{card } A_0 = 1$ et, si $n \geq 1$, tout entier $k \in \llbracket 0, 3^n - 1 \rrbracket$ s'écrit avec au plus n chiffres en base 3 et, pour chacun d'entre eux, nous n'avons que deux possibilités : 0 ou 2. Nous posons alors, pour tout entier naturel n et tout $k \in A_n$:

$$x_{n,k} = \frac{k}{3^n} \text{ et } y_{n,k} = x_{n,k} + \frac{1}{3^n}.$$

Nous prouvons que toute la croissance des f_n , ainsi que celle de f , se concentre sur les intervalles $I_k = [x_{n,k}, y_{n,k}]$ pour $n \in \mathbb{N}$ et $k \in A_n$, en montrant, par récurrence sur $n \in \mathbb{N}$, que, pour tout $n \in \mathbb{N}$, on a $f(y_{n,k}) - f(x_{n,k}) = 1/2^n$ pour tout $k \in A_n$. C'est évident au rang 0. Au rang 1, on a $x_{1,0} = 0$, $y_{1,0} = 1/3$, $x_{1,2} = 2/3$ et $y_{1,2} = 1$, donc $f(y_{1,0}) - f(x_{1,0}) = f(y_{1,2}) - f(x_{1,2}) = 1/2$.

Soit n un entier ≥ 2 et supposons la propriété vraie au rang $n-1$. Soit $k \in A_n$. Alors $0 \leq k < 3^{n-1} - 1$ ou $2 \times 3^{n-1} \leq k \leq 3^n - 1$, car sinon le premier chiffre de l'écriture de k en base 3 serait 1. Dans le premier cas, $0 \leq x_{n,k} < y_{n,k} \leq 1/3$ donc $f(y_{n,k}) - f(x_{n,k}) = (1/2)(f_{n-1}(3x_{n,k}) - f_{n-1}(3y_{n,k}))$. Ainsi $k \in A_{n-1}$, son écriture en base 3 ne contient aucun 1, $3x_{n,k} = x_{n-1,k}$ et $3y_{n,k} = y_{n-1,k}$, donc :

$$\begin{aligned} f(y_{n,k}) - f(x_{n,k}) &= \frac{1}{2} (f_{n-1}(3x_{n,k}) - f_{n-1}(3y_{n,k})) \\ &= \frac{1}{2} (f_{n-1}(x_{n-1,k}) - f_{n-1}(y_{n-1,k})) = \frac{1}{2} \frac{1}{2^{n-1}} = \frac{1}{2^n}. \end{aligned}$$

Dans le second cas, on a, en base 3, $k = 2\alpha_2\alpha_3 \cdots \alpha_n$ et $x_{n,k} = 0, 2\alpha_2\alpha_3 \cdots \alpha_n$ donc $3x_{n,k} - 2 = 0, \alpha_2\alpha_3 \cdots \alpha_n$. On remarque que $k' = \alpha_2\alpha_3 \cdots \alpha_n$ appartient à A_{n-1} et on conclut alors comme dans le premier cas. Ceci achève le raisonnement par récurrence.

Nous allons démontrer que, pour tout entier $n \geq 1$, les points anguleux de la courbe représentative de f_n sont tous des points de la courbe représentative de la fonction limite f . Plus précisément, nous posons, pour tout entier naturel n :

$$Z_n = \left\{ \frac{k}{3^n} \mid k \in \llbracket 0, 3^n \rrbracket \right\}$$

et nous prouvons que, pour tout $x \in Z_n$, $f_p(x) = f_n(x)$ pour tout entier $p \geq n$.

Remarquons d'abord que, pour tout $n \in \mathbb{N}$, les $x_{n,k}$ et les $y_{n,k}$ définis plus haut sont des éléments de Z_n . Par ailleurs $Z_n \subset Z_{n+1}$ pour tout $n \in \mathbb{N}$; en effet, si $x = k/3^n$ où $k \in \llbracket 0, 3^n \rrbracket$, alors $x = (3k)/3^{n+1}$ et $3k \in \llbracket 0, 3^{n+1} \rrbracket$.

Il suffit donc de montrer que, pour tout $n \in \mathbb{N}$ et tout $x \in Z_n$, $f_{n+1}(x) = f_n(x)$; en effet, si $n \in \mathbb{N}$ et si $x \in Z_n$, alors $x \in Z_{n+1}$ donc $f_{n+2}(x) = f_{n+1}(x)$, et on obtient alors le résultat de proche en proche.

Nous prouvons donc, par récurrence sur $n \in \mathbb{N}$, que, pour tout entier naturel n , $f_{n+1}(x) = f_n(x)$ pour tout $x \in Z_n$. Au rang 0, on a $Z_0 = \{0, 1\}$, $f_1(0) = f_0(0) = 0$ et $f_1(1) = f_0(1) = 1$. Soit n un entier ≥ 1 et supposons la propriété vraie au rang $n-1$. Soit x un point de Z_n . On a $x = k/3^n$ où $k \in \llbracket 0, 3^n \rrbracket$. Supposons d'abord que $k < 3^{n-1}$, c'est-à-dire que $0 \leq x < 1/3$; alors $f_{n+1}(x) = (1/2)f_n(3x)$, on a $3x = (3k)/3^n = k/3^{n-1}$, donc $3x$ appartient à Z_{n-1} ce qui, par l'hypothèse de récurrence, donne : $f_n(3x) = f_{n-1}(3x)$ donc $f_{n+1}(x) = (1/2)f_{n-1}(3x) = f_n(x)$.

Supposons maintenant que $3^{n-1} \leq k \leq 2 \times 3^{n-1}$, c'est-à-dire que $1/3 \leq x \leq 2/3$; alors $f_{n+1}(x) = f_n(x) = 1/2$. Supposons enfin que $2 \times 3^{n-1} < k \leq 3^n$, c'est-à-dire que $2/3 < x \leq 1$; alors $3x - 2 = (3k - 2 \times 3^n)/3^n = (k - 2 \times 3^{n-1})/3^{n-1}$, donc $3x - 2$ appartient à Z_{n-1} , ce qui, en appliquant l'hypothèse de récurrence, montre que $f_{n+1}(x) = (1/2)f_n(3x - 2) = (1/2)f_{n-1}(3x - 2) = f_n(x)$. Ceci achève le raisonnement par récurrence.

On en déduit que, pour tout entier naturel n et tout point x de Z_n , on a, pour tout entier $p \geq n$, $f_p(x) = f_n(x)$, d'où, en passant à la limite quand p tend vers l'infini, l'égalité $f(x) = f_n(x)$.

Soit n un entier naturel. Comme A_n possède 2^n éléments et que, pour tout $k \in A_n$, $y_{n,k} - x_{n,k} = 1/3^n$, il vient :

$$\sum_{k \in A_n} (y_{n,k} - x_{n,k}) = \frac{2^n}{3^n} = \left(\frac{2}{3}\right)^n.$$

Par ailleurs, pour tout $k \in A_n$, $f(y_{n,k}) - f(x_{n,k}) = 1/2^n$ donc :

$$\sum_{k \in A_n} |f(y_{n,k}) - f(x_{n,k})| = \sum_{k \in A_n} (f(y_{n,k}) - f(x_{n,k})) = \frac{2^n}{2^n} = 1.$$

Nous posons $\varepsilon_0 = 1$. Soit η un réel > 0 . La suite $((2/3)^n)$ converge vers 0, ce qui justifie le choix d'un entier $p \geq 1$ tel que $(2/3)^p < \eta$. Alors, quels que soient les éléments k et ℓ de A_p tels que $k < \ell$, on a $x_{p,k} < y_{p,k} \leq x_{p,\ell} < y_{p,\ell}$ et les inégalités :

$$\sum_{k \in A_p} (y_{p,k} - x_{p,k}) < \eta \quad \text{et} \quad \sum_{k \in A_p} |f(y_{p,k}) - f(x_{p,k})| = 1 \geq \varepsilon_0.$$

En conclusion, la fonction f n'est pas absolument continue sur $[0, 1]$. $\square$

■ Fonctions périodiques

DÉFINITION 10.1. — Si f est une application de $\mathbb{R}$ dans $\mathbb{R}$, une période de f est un nombre réel P tel que $f(x + P) = f(x)$ pour tout réel x .

THÉORÈME 10.1. — L'ensemble des périodes d'une application de $\mathbb{R}$ dans $\mathbb{R}$ est un sous-groupe du groupe additif $(\mathbb{R}, +)$ de $\mathbb{R}$.

DÉFINITION 10.2. — Une application de $\mathbb{R}$ dans $\mathbb{R}$ est périodique si elle possède au moins une période différente de zéro.

L'opposé d'une période étant une période (théorème 10.1), toute application périodique de $\mathbb{R}$ dans $\mathbb{R}$ admet au moins une période strictement positive.

Une application f de $\mathbb{R}$ dans $\mathbb{R}$ continue, périodique et non constante admet une plus petite période T strictement positive, qui est appelée la période de f , et le groupe des périodes de f est alors $T\mathbb{Z} = \{nT \mid n \in \mathbb{Z}\}$. C'est le cas de sinus et de cosinus pour $T = 2\pi$.

10.4. Fonction périodique non constante n'admettant pas de plus petite période strictement positive.

Nous considérons la fonction de Dirichlet — voir l'exemple 8.1, page 134 — notée f et définie sur $\mathbb{R}$ par : $f(x) = 1$ si x est rationnel ; $f(x) = 0$ si x est irrationnel.

Si P est une période de f , alors $f(P) = f(0) = 1$ donc P est un nombre rationnel. Inversement, si $P \in \mathbb{Q}$, alors, pour tout réel x , x et $x+P$ sont tous deux rationnels ou bien tous deux irrationnels, donc $f(x) = f(x+P)$, ce qui montre que P est une période de f . Le groupe des périodes de f est donc $\mathbb{Q}$. Or la suite $(1/n)_{n \geq 1}$ est à valeurs dans $\mathbb{Q} \cap]0, +\infty[$ et converge vers 0, donc f n'admet pas de plus petite période strictement positive. $\square$

10.5. Fonction périodique qui n'est pas bornée.

Nous introduisons l'application :

$$f : \mathbb{R} \longrightarrow \mathbb{R} \quad \left| \quad \begin{aligned} x \longmapsto f(x) = \begin{cases} 1 & \text{si } x = 0, \\ 0 & \text{si } x \text{ est irrationnel,} \\ q & \text{si } x \in \mathbb{Q} \setminus \{0\} \text{ et si } (p, q) \text{ est le représentant} \\ & \text{irréductible de } x \text{ tel que } q \in \mathbb{N}^*. \end{cases} \end{aligned} \right.$$

On a, pour tout entier $n \geq 1$, $f\left(\frac{1}{n}\right) = n$, donc la fonction f n'est pas bornée sur $\mathbb{R}$.

Si P est une période de f , alors $f(P) = f(0) = 1$, donc P est un entier relatif. Si x est un nombre réel irrationnel, il en est de même de $x+1$, donc $f(x+1) = 0 = f(x)$. Si $x \in \mathbb{Q} \setminus \{0\}$ et si l'on note (p, q) le représentant irréductible de x de dénominateur strictement positif, $f(x) = q$ et, comme $x+1 = (p+q)/q$ et que q et $p+q$ sont premiers entre eux, $f(x+1) = q$ donc $f(x+1) = f(x)$. En conclusion, 1 est une période de f , donc f est périodique. Nous avons même établi que 1 est la plus petite période strictement positive de f et que le groupe de ses périodes est $\mathbb{Z}$. $\square$

■ Fonctions convexes

Rappelons que si E est un espace vectoriel réel et si a et b sont des vecteurs de E , le segment a, b est $[a, b] = \{(1-t)a + tb \mid t \in [0, 1]\}$ et qu'une partie $\mathcal{C}$ de E est un convexe de E si, quels que soient les « points » a et b de $\mathcal{C}$, le segment $[a, b]$ est inclus dans $\mathcal{C}$. Ceci s'applique à $E = \mathbb{R}$; on retrouve ainsi la notion habituelle de segment et les convexes de $\mathbb{R}$ sont les intervalles.

Une fonction f définie sur un intervalle I est convexe sur I si l'ensemble :

$$\mathcal{C}_f^+ = \{(x, y) \mid x \in I \text{ et } y \geq f(x)\}$$

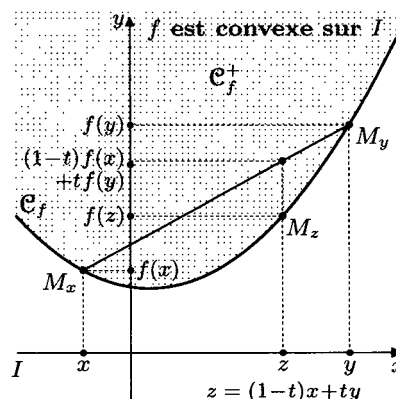
des points de $\mathbb{R}^2$ situés « au-dessus » de la courbe représentative :

$$\mathcal{C}_f = \{(x, y) \mid x \in I \text{ et } y = f(x)\}$$

de f sur I est un convexe du plan $\mathbb{R}^2$, ou encore si, quels que soient les points distincts x et y de I , la corde $[M_x, M_y]$, où $M_x = (x, f(x))$ et $M_y = (y, f(y))$, est au-dessus de l'arc de courbe :

$$(M_x M_y) = \{M_z = (z, f(z)) \mid z \in [x, y]\} \quad (\text{voir le dessin ci-dessus}).$$

Ces définitions géométriques se traduisent par la définition analytique suivante.



DÉFINITION 10.3. — a) Une fonction f définie sur un intervalle I est convexe sur I si, pour tout couple (x, y) de points de I et tout point t du segment $[0, 1]$:

$$f((1-t)x + ty) \leq (1-t)f(x) + tf(y).$$

b) Une fonction f définie sur un intervalle I est concave sur I si sa fonction opposée $-f$ est convexe sur I .

Pour *démontrer* qu'une fonction f définie sur un intervalle I est convexe sur I , il *suffit* de prouver que, pour tout couple (x, y) de points de I tels que $x < y$ et tout $t \in]0, 1[$, $f((1-t)x + ty) \leq (1-t)f(x) + tf(y)$.

Une fonction est convexe et concave sur un intervalle si, et seulement si, c'est la restriction à cet intervalle d'une fonction affine. Une fonction définie et convexe sur un intervalle ouvert est continue sur cet intervalle. Si une fonction f est définie et continue sur un intervalle I et deux fois dérivable sur l'intérieur I_0 de I , f est convexe sur I si, et seulement si, $f''(x) \geq 0$ pour tout point x de I_0 .

10.6. Fonction convexe discontinue.

Nous considérons l'application :

$$\left| \begin{array}{l} f : [0, +\infty[\rightarrow \mathbb{R} \\ x \mapsto f(x) = \begin{cases} 1 & \text{si } x = 0, \\ x^2 & \text{si } x > 0. \end{cases} \end{array} \right.$$

La fonction f n'est pas continue à droite en 0. Soit $x, y \in [0, +\infty[$ tels que $x < y$ et t un point de $]0, 1[$. Si $x > 0$, alors $y > 0$ et $(1-t)x + ty > 0$, donc :

$$\begin{aligned} f((1-t)x + ty) &= ((1-t)x + ty)^2 = (1-t)^2x^2 + t^2y^2 + 2t(1-t)xy \\ &= (1-t)x^2 + ty^2 - t(1-t)(x^2 - 2xy + y^2) \\ &= (1-t)f(x) + tf(y) - \underbrace{t(1-t)(x-y)^2}_{\geq 0} \leq (1-t)f(x) + tf(y). \end{aligned}$$

Si $x = 0$, on déduit de $ty > 0$ que $f(ty) = t^2y^2$ et de $0 < t < 1$ que $t^2 < t$, donc :

$$f((1-t)x + ty) = f(ty) = \underbrace{t^2}_{< t} y^2 \leq ty^2 < \underbrace{(1-t)}_{> 0} + ty^2 = (1-t)f(x) + tf(y).$$

En conclusion, f est discontinue et convexe sur l'intervalle $[0, +\infty[$. $\square$

Une fonction f définie et continue sur un intervalle I est convexe sur I dès que, pour tout couple (x, y) de points de I :

$$(1) \quad f\left(\frac{x+y}{2}\right) \leq \frac{f(x) + f(y)}{2}.$$

Il suffit donc pour une fonction continue de réaliser la définition pour $t = \frac{1}{2}$.

10.7. Fonction qui vérifie (1) quels que soient les réels x et y mais qui n'est pas convexe sur $\mathbb{R}$.

La définition de cette fonction nécessite l'utilisation de l'axiome du choix³, par l'intermédiaire de l'une de ses conséquences : *Dans un espace vectoriel de dimension infinie, tout sous-espace vectoriel admet au moins un supplémentaire.*

3. Voir le chapitre 1, page 14.

Le corps $\mathbb{Q}$ des nombres rationnels étant un sous-corps de $\mathbb{R}$, $\mathbb{R}$ est canoniquement un espace vectoriel sur $\mathbb{Q}$; les scalaires sont alors les nombres rationnels et les vecteurs les nombres réels. L'espace vectoriel $\mathbb{R}$ est de dimension infinie⁴ sur $\mathbb{Q}$. Comme $\mathbb{Q}$ est un sous-espace vectoriel de l'espace vectoriel $\mathbb{R}$ sur $\mathbb{Q}$, il admet au moins un supplémentaire dans $\mathbb{R}$. Nous choisissons un supplémentaire V de $\mathbb{Q}$ dans $\mathbb{R}$. Ainsi V est un sous-espace vectoriel de $\mathbb{R}$ et $\mathbb{R} = \mathbb{Q} \oplus V$. Nous notons p la projection vectorielle sur $\mathbb{Q}$ parallèlement à V et nous introduisons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = (p(x))^2. \end{array} \right.$$

Soit x et y des nombres réels. Nous décomposons x et y dans la somme directe $\mathbb{R} = \mathbb{Q} \oplus V$: $x = r + v$ et $y = s + w$ où $r, s \in \mathbb{Q}$ et $v, w \in V$. Alors :

$$\frac{x+y}{2} = \frac{r+s}{2} + \frac{v+w}{2}, \quad \frac{r+s}{2} \in \mathbb{Q} \text{ et } \frac{v+w}{2} = \frac{1}{2}v + \frac{1}{2}w \in V$$

donc :

$$f(x) = r^2, \quad f(y) = s^2 \text{ et } f\left(\frac{x+y}{2}\right) = \left(\frac{r+s}{2}\right)^2 = \frac{r^2+s^2}{4} + \frac{rs}{2}.$$

De $(r-s)^2 = r^2 + s^2 - 2rs \geq 0$ on déduit que $\frac{rs}{2} \leq \frac{r^2+s^2}{4}$. Par conséquent :

$$f\left(\frac{x+y}{2}\right) = \frac{r^2+s^2}{4} + \frac{rs}{2} \leq \frac{r^2+s^2}{2} = \frac{f(x)+f(y)}{2}$$

ce qui montre que l'assertion (1) est vraie.

On a $\mathbb{Q} \subsetneq \mathbb{R}$, donc $V \neq \{0\}$. Nous choisissons un élément u de V tel que $u > 0$ — il en existe, car si $x \in V$, $-x \in V$ — et un nombre rationnel r tel que $0 < r < u$, et nous posons :

$$t = \frac{r}{u}.$$

Alors $0 < t < 1$ et $r = (1-t) \times 0 + tu$. On a $f(0) = (p(0))^2 = 0$ et $f(u) = (p(u))^2 = 0$ donc $f((1-t) \times 0 + tu) = f(r) = r^2 > 0 = (1-t)f(0) + tf(u)$. Il en résulte que la fonction f n'est pas convexe sur $\mathbb{R}$ — en fait la restriction de f à aucun intervalle d'intérieur non vide n'est convexe sur cet intervalle. $\square$

Nous avons rappelé qu'une fonction définie et convexe sur un intervalle ouvert est continue sur cet intervalle; il en est donc de même pour une fonction concave. Par contre, la continuité d'une fonction sur un intervalle I ne permet pas d'affirmer qu'elle est convexe ou concave sur I , ni même convexe ou concave sur au moins un intervalle d'intérieur non vide inclus dans I .

10.8. Fonction continue sur $\mathbb{R}$ mais qui, sur tout intervalle d'intérieur non vide, n'est ni convexe ni concave.

Nous reprenons la fonction g , la suite $(u_n)_{n \geq 0}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$ et la fonction f définies dans l'exemple 9.6 (pages 161 et 162). La fonction f ainsi obtenue est continue sur $\mathbb{R}$.

Soit n un entier naturel et $k \in \mathbb{Z}$. Nous posons :

$$a = \frac{k}{4^n}, \quad b = a + \frac{1}{4^{2n+1}} = \frac{4^{n+1}k + 1}{4^{2n+1}} \text{ et } c = a - \frac{1}{4^{2n+1}} = \frac{4^{n+1}k - 1}{4^{2n+1}}.$$

4. Voir, dans le chapitre 5, au bas de la page 94.

Nous avons établi dans l'exemple 10.2 (pages 191 et 192), que $f(b) - f(a) \geq b - a$ et $f(c) - f(a) \geq a - c = b - a$. Or $b - a > 0$, donc $f(a) < f(b)$ et $f(a) < f(c)$. De plus, a est le milieu de (c, b) , donc :

$$f\left(\frac{c+b}{2}\right) = f(a) = \frac{f(a) + f(a)}{2} < \frac{f(c) + f(b)}{2}.$$

Par suite f n'est pas concave sur $[c, b]$. Or $[c, b] \subset S_{n,k} = \left[\frac{k-1}{4^n}, \frac{k+1}{4^n}\right]$, donc la fonction f n'est pas concave sur le segment $S_{n,k}$.

Nous posons :

$$x = \frac{k+1}{4^n}, \quad z = \frac{4k+1}{4^{n+1}} \quad \text{et} \quad t = \frac{1}{4}.$$

Alors $t \in]0, 1[$ et $z = (1-t)a + tx$.

Soit p un entier naturel. Si $p \geq n+1$, alors, pour tout $m \in \mathbb{Z}$:

$$\frac{m}{4^{n+1}} = \frac{4^{p-n-1}m}{4^p} \quad \text{et} \quad 4^{p-n-1}m \in \mathbb{Z}$$

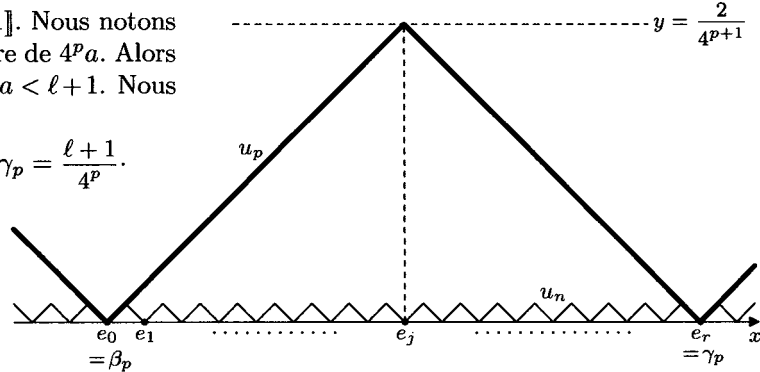
donc $u_p\left(\frac{m}{4^{n+1}}\right) = 0$. Par suite : (1) $u_p(a) = u_p(z) = u_p(x) = 0$ pour tout $p \geq n+1$.

On a $u_n(a) = u_n(x) = 0$ et $u_n(z) = 1/4^{n+1}$ (voir le dessin ci-dessus) donc :

$$(2) \quad u_n(z) > 0 = (1-t)u_n(a) + tu_n(x).$$

Soit $p \in [0, n-1]$. Nous notons ℓ la partie entière de $4^p a$. Alors $\ell \in \mathbb{Z}$ et $\ell \leq 4^p a < \ell + 1$. Nous posons :

$$\beta_p = \frac{\ell}{4^p} \quad \text{et} \quad \gamma_p = \frac{\ell+1}{4^p}.$$



Comme $4^p a = 4^p \frac{k}{4^n} = \frac{k}{4^{n-p}}$, $4^{n-p}\ell \leq k < 4^{n-p}(\ell+1)$ donc $k \leq 4^{n-p}(\ell+1) - 1$.

On pose $m_0 = 4^{n-p}\ell$, $r = 4^{n-p}$ et, pour tout $i \in [0, r]$, $m_i = m_0 + i$ et $e_i = \frac{m_i}{4^n}$.

On a alors $[\beta_p, \gamma_p] = [e_0, e_r]$. Comme $n - p \geq 1$, le quotient $j = r/2$ est un entier naturel et e_j est le milieu du segment $[\beta_p, \gamma_p]$ (voir le dessin ci-dessus). De plus $k \in \{m_0, m_1, \dots, m_j, \dots, m_{r-1}\}$, donc il existe un indice $i \in [0, r-1]$ tel que $[a, x] = [k/4^n, (k+1)/4^n] = [e_i, e_{i+1}]$. Par conséquent la restriction de u_p à $[a, x]$ est affine, donc $u_p(z) = u_p((1-t)a + tx) = (1-t)u_p(a) + tu_p(x)$. Finalement :

$$(3) \quad u_p(z) = (1-t)u_p(a) + tu_p(x) \quad \text{pour tout } p \in [0, n-1].$$

On a :

$$f(z) = \sum_{p=0}^{+\infty} u_p(z), \quad f(a) = \sum_{p=0}^{+\infty} u_p(a) \quad \text{et} \quad f(x) = \sum_{p=0}^{+\infty} u_p(x),$$

on déduit de (1) que :

$$f(z) = u_n(z) + \sum_{p=0}^{n-1} u_p(z), \quad f(a) = u_n(a) + \sum_{p=0}^{n-1} u_p(a) \quad \text{et} \quad f(x) = u_n(x) + \sum_{p=0}^{n-1} u_p(x)$$

et de (2) et (3) que $f((1-t)a+tx) = f(z) > (1-t)f(a) + tf(x)$. La fonction f n'est donc pas convexe sur $[a, x]$. Or $[a, x]$ est inclus dans $S_{n,k}$ (segment défini à la page précédente), donc la fonction f n'est pas convexe sur $S_{n,k}$.

Si I est un intervalle d'intérieur non vide, il existe un entier naturel n et un entier relatif k tels que $S_{n,k}$ est inclus dans I (voir page 192, à la fin de l'exemple 10.2) donc f n'est ni convexe ni concave sur I .

En conclusion, la fonction f est continue sur $\mathbb{R}$ mais, sur tout intervalle d'intérieur non vide, f n'est ni convexe ni concave. $\square$

Si f est une fonction définie sur un intervalle I , f est convexe sur I si, et seulement si, f est continue sur I , dérivable à droite et à gauche en tout point de l'intérieur de I et, quels que soient les points a et b de l'intérieur de I tels que $a < b$, $f'_g(a) \leq f'_d(a) \leq f'_g(b)$. On en déduit que si une fonction f définie sur un intervalle I est continue et convexe sur I , l'ensemble des points de l'intérieur de I en lesquels f n'est pas dérivable est au plus dénombrable.

10.9. Fonction convexe sur un intervalle I qui n'est pas dérivable en les points d'une partie dénombrable de l'intérieur de I dense dans I .

Nous définissons l'application f de $[0, 1]$ dans $\mathbb{R}$ comme dans l'exemple 10.1 (pages 190 et 191). Comme f est (strictement) croissante sur $[0, 1]$, f est intégrable sur le segment $[0, 1]$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} F : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto F(x) = \int_0^x f(t) dt. \end{array} \right.$$

La fonction F est continue sur le segment $[0, 1]$. En tout point a de $]0, 1]$, f admet pour limite $f(a^-) = \sup\{f(x) \mid 0 \leq x < a\}$ à gauche et on a $f(a^-) \leq f(a)$; en tout point a de $[0, 1[$, f admet pour limite $f(a^+) = \inf\{f(x) \mid a < x \leq 1\}$ à droite et on a $f(a) \leq f(a^+)$. De plus, quels que soient les points a et b de I tels que $a < b$, on a $f(a) \leq f(a^+) \leq f(b^-) \leq f(b)$.

Par suite F est dérivable à gauche et à droite en tout point de l'intérieur $]0, 1[$ de $[0, 1]$ et, pour tout point a de $]0, 1[$, $F'_g(a) = f(a^-)$ et $F'_d(a) = f(a^+)$, d'où l'on déduit que, quels que soient les points a et b de $]0, 1[$ tels que $a < b$, $F'_g(a) \leq F'_d(a) \leq F'_g(b)$. Il en résulte que F est convexe sur $[0, 1]$.

L'ensemble $\mathcal{Q} = \mathbb{Q} \cap]0, 1[$ est une partie dénombrable de l'intérieur $]0, 1[$ de $[0, 1]$ et $\mathcal{Q}$ est dense dans $[0, 1]$. De plus, pour tout point a de $\mathcal{Q}$, f est discontinue en a donc F n'est pas dérivable en a . $\square$

■ Fonctions bornées

10.10. Fonction qui n'est bornée sur aucun intervalle d'intérieur non vide.

Nous reprenons la fonction f définie à l'exemple 10.5 (page 197). Alors $f(0) = 1$, $f(x) = 0$ si x est irrationnel et, si x est un nombre rationnel non nul et (p, q) son représentant irréductible de dénominateur positif, $f(x) = q$.

Soit a et b des réels tels que $a < b$. Soit M un réel strictement positif. L'intervalle

ouvert $]a, b[$ contient une infinité de rationnels, et pour seulement un nombre fini d'entre eux, le dénominateur q de leur représentant irréductible (p, q) tel que $q \in \mathbb{N}^*$ est inférieur ou égal à M ; en effet, si (p, q) est le représentant irréductible de dénominateur positif d'un rationnel r appartenant à $[a, b]$ et si $q \leq M$, alors q appartient à $\llbracket 1, m \rrbracket$ où m est la partie entière de M et l'entier naturel $|p|$ est inférieur ou égal à la partie entière de $M \times \max(|a|, |b|)$. Nous pouvons donc choisir un rationnel r appartenant à $]a, b[$ de représentant irréductible (p, q) tel que $q > M$, et alors $f(r) > M$. Par conséquent la fonction f n'est pas majorée sur l'intervalle $]a, b[$, donc f n'est pas bornée sur $]a, b[$. Tout intervalle d'intérieur non vide contenant au moins un intervalle ouvert $]a, b[$ où a et b sont des réels tels que $a < b$, la fonction f n'est bornée sur aucun intervalle d'intérieur non vide. $\square$

Une fonction continue sur un compact est bornée et atteint ses bornes sur ce compact. Par contre une fonction bornée sur un intervalle compact n'atteint en général pas sa borne supérieure ni sa borne inférieure sur cet intervalle.

10.11. Fonction bornée sur $\mathbb{R}$ qui n'atteint ses bornes sur aucun intervalle compact d'intérieur non vide.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x \text{ est irrationnel ou égal à } 0, \\ \frac{(-1)^q(q-1)}{q} & \text{si } x \in \mathbb{Q} \setminus \{0\} \text{ et si } (p, q) \text{ est le} \\ & \text{représentant irréductible de} \\ & x \text{ tel que } q \in \mathbb{N}^*. \end{cases} \end{array} \right.$$

On a, pour tout rationnel $r \neq 0$ dont le représentant irréductible de dénominateur positif est (p, q) :

$$|f(r)| = \frac{q-1}{q} < 1$$

donc f est minorée par -1 et majorée par 1 sur $\mathbb{R}$. Par suite f est bornée sur $\mathbb{R}$. Soit I un intervalle d'intérieur non vide. Nous notons $\mathcal{D}$ l'ensemble des dénominateurs des représentants irréductibles de dénominateur positif des rationnels appartenant à I . En raisonnant comme dans l'exemple précédent 10.10, nous voyons que $\mathcal{D}$ n'est pas majoré dans $\mathbb{N}$. Il en résulte que $\mathcal{D}$ contient des entiers naturels pairs et impairs aussi grands que l'on veut. De plus :

$$\frac{n-1}{n} \underset{n \rightarrow \infty}{\sim} 1$$

donc -1 est la borne inférieure et 1 la borne supérieure de f sur I . Or, pour tout entier $q \geq 1$, $(q-1)/q < 1$, donc f n'atteint ni sa borne supérieure, ni sa borne inférieure sur l'intervalle I . $\square$

■ Fonctions à variation bornée

Soit a et b des nombres réels tels que $a < b$. Une subdivision du segment $[a, b]$ est un $(n+1)$ -uplet $\sigma = (x_0, x_1, \dots, x_n)$ de réels où n est un entier ≥ 1 et :

$$a = x_0 < x_1 < \dots < x_i < x_{i+1} < \dots < x_{n-1} < x_n = b,$$

les bornes de σ sont alors les points $a_0, a_1, \dots, a_n$ et, à toute fonction f définie sur

le segment $[a, b]$, nous associons le nombre réel :

$$v(f, \sigma) = \sum_{i=0}^{n-1} |f(x_{i+1}) - f(x_i)|.$$

DÉFINITION 10.4. — Une fonction f définie sur $[a, b]$ est à variation bornée⁵ sur le segment $[a, b]$ si l'ensemble des $v(f, \sigma)$ quand σ décrit l'ensemble des subdivisions de $[a, b]$ est majoré dans $\mathbb{R}$.

THÉORÈME 10.2. — Toute fonction définie et lipschitzienne sur le segment $[a, b]$ est à variation bornée sur $[a, b]$ et toute fonction définie et de classe $\mathcal{C}^1$ sur $[a, b]$ est à variation bornée sur $[a, b]$.

En effet, si f est lipschitzienne de rapport k sur $[a, b]$, on a $v(f, \sigma) \leq k(b-a)$ pour toute subdivision σ de $[a, b]$ et, si f est de classe $\mathcal{C}^1$ sur $[a, b]$, f' est continue et par suite bornée sur $[a, b]$, donc, k étant un majorant de $|f'|$ sur $[a, b]$, l'inégalité des accroissements finis montre que f est lipschitzienne de rapport k sur $[a, b]$.

Le théorème 10.2 peut tomber en défaut si l'on suppose seulement que la fonction est continue sur le segment $[a, b]$.

10.12. Fonction continue sur un segment qui n'est pas à variation bornée sur ce segment.

Nous considérons l'application :

$$\left| \begin{array}{l} f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} x \sin \frac{\pi}{x} & \text{si } x \neq 0, \\ 0 & \text{si } x = 0. \end{cases} \end{array} \right.$$

Pour tout point x de $[0, 1]$, $|f(x)| \leq |x|$, donc $f(x)$ admet pour limite $0 = f(0)$ quand x tend vers 0. Par suite f est continue en 0 et, comme f est clairement continue en tout point de $]0, 1]$, la fonction f est continue sur le segment $[0, 1]$.

Nous introduisons la suite $(\sigma_n)_{n \geq 1}$ de subdivisions de $[0, 1]$, de terme général :

$$\sigma_n = \left(0, \frac{1}{n + \frac{1}{2}}, \frac{1}{(n-1) + \frac{1}{2}}, \dots, \frac{1}{k + \frac{1}{2}}, \frac{1}{(k-1) + \frac{1}{2}}, \dots, \frac{1}{1 + \frac{1}{2}}, 1 \right).$$

On a, pour tout entier $k \geq 1$:

$$f\left(\frac{1}{(k-1) + \frac{1}{2}}\right) = f\left(\frac{1}{k - \frac{1}{2}}\right) = \frac{2(-1)^{k-1}}{2k-1} \text{ et } f\left(\frac{1}{k + \frac{1}{2}}\right) = \frac{2(-1)^k}{2k+1}$$

donc, pour tout entier $n \geq 2$:

$$\begin{aligned} v(f, \sigma_n) &\geq \sum_{k=2}^n \left| f\left(\frac{1}{(k-1) + \frac{1}{2}}\right) - f\left(\frac{1}{k + \frac{1}{2}}\right) \right| = 2 \sum_{k=2}^n \left(\frac{1}{2k-1} + \frac{1}{2k+1} \right) \\ &\geq 2 \sum_{k=2}^n \left(2 \times \frac{1}{2k+2} \right) = 2 \sum_{k=2}^n \frac{1}{k+1} = 2 \sum_{p=3}^{n+1} \frac{1}{p}. \end{aligned}$$

Comme la série harmonique $\sum_{n \geq 1} (1/n)$ diverge, la suite $(v(f, \sigma_n))$ tend vers $+\infty$, donc f n'est pas à variation bornée sur le segment $[0, 1]$. $\square$

5. Cette notion est introduite par Camille Jordan, en 1881, à propos de la rectification des courbes ; voir dans le chapitre 18 ce qui suit l'exemple 18.11, page 348.

Si a et b sont des réels tels que $a < b$, si $p \in \mathbb{N}^*$ et si $f_1, \dots, f_p$ sont des fonctions définies sur $[a, b]$ et $\lambda_1, \dots, \lambda_p$ des réels, alors, pour toute subdivision σ de $[a, b]$:

$$v\left(\sum_{i=1}^p \lambda_i f_i\right) \leq \sum_{i=1}^p |\lambda_i| v(f_i, \sigma).$$

On en déduit le théorème suivant.

THÉORÈME 10.3. — Si a et b sont des réels tels que $a < b$, si $p \in \mathbb{N}^*$ et si $f_1, \dots, f_p$ sont des fonctions définies et à variation bornée sur le segment $[a, b]$ et $\lambda_1, \dots, \lambda_p$ des nombres réels, alors la combinaison linéaire $\lambda_1 f_1 + \dots + \lambda_p f_p$ est une fonction à variation bornée sur le segment $[a, b]$.

10.13. Fonction continue sur $\mathbb{R}$ dont la restriction n'est à variation bornée sur aucun segment d'intérieur non vide.

Nous utilisons dans cet exemple les propriétés de la convergence uniforme des suites et séries de fonctions, rappelées aux chapitre 12 et 13.

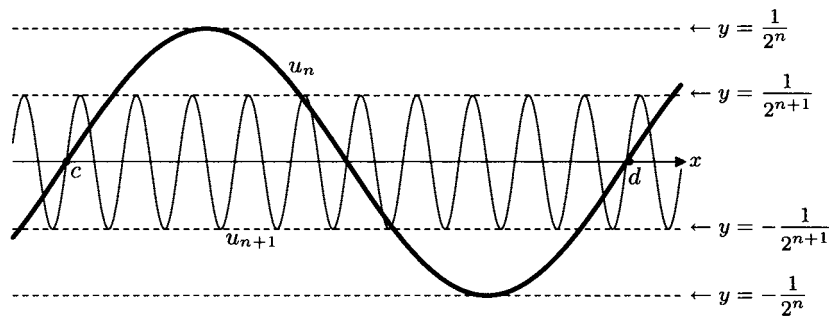
Nous introduisons la suite $(u_n)_{n \geq 0}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} u_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_n(x) = \frac{\sin(10^n x)}{2^n}. \end{array} \right.$$

Pour tout entier naturel n , la fonction u_n est continue sur $\mathbb{R}$, $\frac{2\pi}{10^n}$ est la période de u_n et, pour tout $\ell \in \mathbb{Z}$:

$$u_n\left(\frac{(4\ell-1)\pi}{2 \times 10^n}\right) = -\frac{1}{2^n} \text{ et } u_n\left(\frac{(4\ell+1)\pi}{2 \times 10^n}\right) = \frac{1}{2^n},$$

u_n est croissante sur $\left[\frac{(4\ell-1)\pi}{2 \times 10^n}, \frac{(4\ell+1)\pi}{2 \times 10^n}\right]$ et décroissante sur $\left[\frac{(4\ell+1)\pi}{2 \times 10^n}, \frac{(4\ell+3)\pi}{2 \times 10^n}\right]$.



Pour tout réel x et tout entier naturel n , $|u_n(x)| \leq 1/2^n$, donc la série de fonctions $\sum_n u_n$ converge uniformément sur $\mathbb{R}$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \sum_{n=0}^{+\infty} u_n(x) \end{array} \right.$$

et prouve que la fonction f est continue sur $\mathbb{R}$.

Soit a et b des nombres réels tels que $a < b$. La suite de terme général $(3\pi)/10^n$ converge vers 0, ce qui justifie le choix d'un entier $n \geq 1$ tel que $(3\pi)/10^n < (b-a)$. La somme k de 1 et de la partie entière de $(10^n a)/(2\pi)$ est un entier relatif et :

$$a \leq c = \frac{2k\pi}{10^n} < d = \frac{2(k+1)\pi}{10^n} < b.$$

Pour tout entier naturel i , la fonction u_i est de classe $\mathcal{C}^1$ sur le segment $[c, d]$, donc on déduit du théorème 10.2 (page 203) que u_i est à variation bornée sur $[c, d]$.

Soit $q \in \mathbb{N}^*$. Nous introduisons la subdivision σ_q de $[c, d]$ dont les bornes sont c, d et, pour tout $p \in \llbracket 0, q \rrbracket$, tous les points de $[c, d]$ en lesquels u_{n+p} présente un maximum ou un minimum, ces points étant mis en ordre croissant.

Soit $p \in \llbracket 0, q \rrbracket$. En posant $m = 10^p$ et :

$$a_0 = c, \quad a_i = c + \frac{(2i-1)\pi}{2 \times 10^{n+p}} \quad \text{pour tout } i \in \llbracket 1, m \rrbracket \quad \text{et } a_{m+1} = d,$$

alors $c = a_0 < a_1 < a_2 < \dots < a_q < a_{m+1} = d$, $u_{n+p}(c) = u_{n+p}(d) = 0$, les points $a_1, \dots, a_m$ sont les points de $[c, d]$ en lesquels u_{n+p} admet alternativement un maximum, égal à $1/2^{n+p}$, et un minimum, égal à $-1/2^{n+p}$, u_{n+p} est monotone sur $[c, a_1]$, sur les $[a_i, a_{i+1}]$ pour tout $i \in \llbracket 1, m-1 \rrbracket$ et sur $[a_m, d]$. Les points $a_0, a_1, \dots, a_m$ sont des bornes de la subdivision σ_q et u_{n+p} est, pour tout $i \in \llbracket 0, m \rrbracket$, monotone sur $[a_i, a_{i+1}]$. Par suite, pour tout $i \in \llbracket 0, m \rrbracket$, $|u_{n+p}(a_{i+1}) - u_{n+p}(a_i)|$ est la somme des $|u_{n+p}(b_{j+1}) - u_{n+p}(b_j)|$ où $b_1, b_2, \dots$ sont les bornes de σ_q appartenant au segment $[a_i, a_{i+1}]$. Il en résulte que $v(u_{n+p}, \sigma_q)$ est la somme des $|u_{n+p}(a_{i+1}) - u_{n+p}(a_i)|$ pour i décrivant $\llbracket 0, m \rrbracket$; or $|u_{n+p}(a_{i+1}) - u_{n+p}(a_i)|$ est égal à $1/2^{n+p}$ pour $i = 0$ et $i = m$, et à $2/2^{n+p}$ pour tout $i \in \llbracket 1, m-1 \rrbracket$, donc :

$$v(u_{n+p}, \sigma_q) = \frac{1}{2^{n+p}} + \frac{2(m-1)}{2^{n+p}} + \frac{1}{2^{n+p}} = \frac{2m}{2^{n+p}} = \frac{2 \times 10^p}{2^{n+p}} = \frac{5^p}{2^{n-1}}.$$

Nous posons, pour tout nombre réel x :

$$g(x) = \sum_{\ell=0}^{n-1} u_\ell(x) \quad \text{et} \quad h(x) = \sum_{\ell=n}^{+\infty} u_\ell(x).$$

Nous construisons ainsi des applications g et h de $\mathbb{R}$ dans $\mathbb{R}$ telles que :

$$f = g + h.$$

Pour tout entier $\ell > n + q$, u_ℓ s'annule sur toutes les bornes de σ_q , donc $v(h, \sigma_q)$ est égal à $v(h_0, \sigma_q)$ pour la fonction :

$$h_0 = \sum_{\ell=n}^{n+q} u_\ell = \sum_{p=0}^q u_{n+p} = \left(\sum_{p=0}^{q-1} u_{n+p} \right) + u_{n+q}.$$

Remarquons que si h_1 et h_2 sont des fonctions définies sur $[c, d]$ et σ une subdivision de $[c, d]$, alors $v(h_1, \sigma) \leq v(h_1 + h_2, \sigma) + v(h_2, \sigma)$ puisque $h_1 = (h_1 + h_2) + (-h_2)$ et que $v(-h_2, \sigma) = v(h_2, \sigma)$, donc $v(h_1 + h_2, \sigma) \geq v(h_1, \sigma) - v(h_2, \sigma)$.

On a $v(u_{n+q}, \sigma_q) = \frac{5^q}{2^{n-1}}$. Pour tout $p \in \llbracket 0, q-1 \rrbracket$, $v(u_{n+p}, \sigma_q) = \frac{5^p}{2^{n-1}}$, donc :

$$v(h_0 - u_{n+q}, \sigma_q) \leq \sum_{p=0}^{q-1} v(u_{n+p}, \sigma_q) \leq \sum_{p=0}^{q-1} \frac{5^p}{2^{n-1}} = \frac{1}{2^{n-1}} \times \frac{5^q - 1}{5 - 1} < \frac{5^q}{2^{n+1}}.$$

En utilisant la remarque précédente et les fonctions u_{n+q} et $h_0 - u_{n+q}$, on obtient :

$$v(h, \sigma_q) = v(h_0, \sigma_q) \geq v(u_{n+q}, \sigma_q) - v(h_0 - u_{n+q}) \geq \frac{5^q}{2^{n-1}} - \frac{5^q}{2^{n+1}} = 3 \frac{5^q}{2^{n+1}}.$$

Soit A un nombre réel. La suite $(\beta_q)_{q \geq 1}$, de terme général $\beta_q = 5^q/2^{n+1}$, diverge vers $+\infty$, ce qui justifie le choix d'un entier $q \geq 1$ tel que $5^q/2^{n+1} > A/3$. En définissant la subdivision σ_q de $[c, d]$ comme dans ce qui précède, on obtient l'inégalité $v(h, \sigma_q) > A$. Par conséquent h n'est pas à variation bornée sur $[c, d]$.

Le théorème 10.3 (page 204) montre que la fonction $g = \sum_{\ell=0}^{n-1} u_\ell$ est à variation bornée sur $[c, d]$ comme combinaison linéaire d'une famille de fonctions à variation bornée sur le segment $[c, d]$. Or $f = g + h$ et h n'est pas à variation bornée sur $[c, d]$, donc f n'est pas à variation bornée sur $[c, d]$. *A fortiori* la fonction f n'est pas à variation bornée sur le segment $[a, b]$. $\square$

Chapitre 11

Intégration

Le calcul d'intégral est apparu dans l'Antiquité avec Archimède, à l'occasion des calculs des aires et de volumes. Cependant la notion d'intégrale d'une fonction est née avec Isaac Newton et Gottfried Leibniz à la fin du XVII^e siècle. Augustin-Louis Cauchy est le premier à construire l'intégrale, en se limitant aux fonctions continues sur un intervalle. Bernhard Riemann propose une construction plus générale basée sur des encadrements par des suites de fonctions en escalier dont les intégrales convergent vers le même réel. En 1902, Henri Lebesgue, partant de la notion de mesure, définit une classe plus large de fonctions intégrables. Nous nous intéresserons à ces deux dernières familles de fonctions intégrables.

Dans tout le chapitre, les fonctions sont des fonctions de $\mathbb{R}$ dans $\mathbb{R}$.

■ Intégrale de Riemann

Soit a et b des nombres réels tels que $a < b$. Une subdivision du segment $[a, b]$ est un $(n + 1)$ -uplet $\sigma = (x_0, x_1, \dots, x_n)$ de réels où n est un entier ≥ 1 et :

$$a = x_0 < x_1 < \dots < x_i < x_{i+1} < \dots < x_{n-1} < x_n = b,$$

les bornes d'une telle subdivision σ sont alors les points $x_0, x_1, \dots, x_n$ et les intervalles ouverts créés par σ sont les $]x_i, x_{i+1}[$ pour tout $i \in \llbracket 0, n-1 \rrbracket$.

Une application f de $[a, b]$ dans $\mathbb{R}$ est en escalier sur $[a, b]$ s'il existe au moins une subdivision de $[a, b]$ telle que f est constante sur chacun des intervalles ouverts créés par cette subdivision, que l'on appelle une subdivision de $[a, b]$ adaptée à f .

Si f est une application de $[a, b]$ dans $\mathbb{R}$ en escalier sur le segment $[a, b]$, le réel $\sum_{i=0}^{n-1} (x_{i+1} - x_i) \alpha_i$, où $\sigma = (x_0, x_1, \dots, x_n)$ est une subdivision de $[a, b]$ adaptée à f et, pour tout $i \in \llbracket 0, n-1 \rrbracket$, α_i la valeur constante de f sur $]x_i, x_{i+1}[$, est indépendant de la subdivision σ adaptée à f ; nous notons $I_{[a,b]}(f)$ cette valeur commune à toutes les subdivisions adaptées à f — c'est l'intégrale de f sur $[a, b]$.

DÉFINITION 11.1. — Une fonction f définie sur le segment $[a, b]$ est intégrable au sens de Riemann sur $[a, b]$ si, pour tout nombre réel $\varepsilon > 0$, il existe des applications φ et ψ de $[a, b]$ dans $\mathbb{R}$ en escalier sur $[a, b]$ telles que :

$$\varphi \leq f \leq \psi \text{ sur } [a, b] \text{ et } I_{[a,b]}(\psi - \varphi) \leq \varepsilon.$$

Si f est une fonction définie et intégrable au sens de Riemann sur le segment $[a, b]$, f est bornée sur $[a, b]$ et la borne supérieure des $I_{[a,b]}(\varphi)$ pour $\varphi \leq f$ en escalier sur $[a, b]$ est égale à la borne inférieure des $I_{[a,b]}(\psi)$ pour $\psi \geq f$ en escalier sur $[a, b]$, et leur valeur commune est l'intégrale de f sur $[a, b]$, notée :

$$\int_a^b f(t) dt.$$

A partir de là, on prouve tous les résultats classiques de l'intégrale de Riemann. Rappelons que si f est une application de $[a, b]$ dans $\mathbb{R}$ en escalier sur $[a, b]$, alors f est intégrable au sens de Riemann sur $[a, b]$ et $\int_a^b f(t) dt = I_{[a,b]}(f)$.

Pour l'étude de l'intégrabilité et de l'intégrale au sens de Lebesgue, on pourra voir par exemple [BR1A], chapitres 7 et 8.

En pratique la presque totalité des fonctions courantes bornées sur un segment y sont intégrables au sens de Riemann. Il est cependant aisé de trouver des exemples de fonctions bornées qui ne sont pas intégrables au sens de Riemann.

11.1. Fonction bornée qui n'est pas intégrable au sens de Riemann, mais qui l'est sens de Lebesgue.

Nous considérons la fonction de Dirichlet¹, c'est-à-dire l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 1 & \text{si } x \text{ est rationnel,} \\ 0 & \text{si } x \text{ est irrationnel,} \end{cases} \end{array} \right.$$

qui est évidemment bornée sur $\mathbb{R}$.

Soit a et b des nombres réels tels que $a < b$.

Soit φ une application de $[a, b]$ dans $\mathbb{R}$ en escalier sur le segment $[a, b]$ telle que $\varphi \leq f$ (resp. $\varphi \geq f$) sur $[a, b]$. Nous choisissons une subdivision $\sigma = (x_0, x_1, \dots, x_n)$ de $[a, b]$ adaptée à φ et nous notons, pour tout $i \in \llbracket 0, n-1 \rrbracket$, α_i la valeur constante de φ sur $]x_i, x_{i+1}[$. Si i appartient à $\llbracket 0, n-1 \rrbracket$, l'intervalle ouvert $]x_i, x_{i+1}[$ contient au moins un irrationnel u (resp. un rationnel r), donc on a $\alpha_i = \varphi(u) \leq f(u) = 0$ (resp. $\alpha_i = \varphi(r) \geq f(r) = 1$). Il en résulte que $I_{[a,b]}(\varphi) \leq 0$ (resp. $I_{[a,b]}(\varphi) \geq b-a$). Ainsi, quelles que soient les applications φ et ψ de $[a, b]$ dans $\mathbb{R}$ en escalier sur $[a, b]$ telles que $\varphi \leq f \leq \psi$ sur $[a, b]$, on a $I_{[a,b]}(\psi - \varphi) = I_{[a,b]}(\psi) - I_{[a,b]}(\varphi) \geq b-a$. Or $b-a > 0$, donc f n'est pas intégrable au sens de Riemann sur le segment $[a, b]$. Comme f est nulle sauf sur un ensemble de mesure nulle, f est intégrable au sens de Lebesgue sur $[a, b]$ et son intégrale est nulle. $\square$

Si une fonction est intégrable au sens de Riemann sur un segment, son carré l'est aussi. Cependant la réciproque est fausse.

11.2. Fonction de carré intégrable au sens de Riemann qui n'est pas intégrable.

Nous introduisons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 1 & \text{si } x \text{ est rationnel,} \\ -1 & \text{si } x \text{ est irrationnel.} \end{cases} \end{array} \right.$$

1. Voir les exemples 8.1 page 134 et 10.4 page 196.

De même que dans l'exemple précédent 11.1, f n'est pas intégrable au sens de Riemann sur le segment $[0, 1]$. Cependant f^2 est constante égale à 1 sur $[0, 1]$, donc f^2 est intégrable au sens de Riemann sur $[0, 1]$. $\square$

La fonction f de l'exemple précédent 11.2 est, comme la fonction de Dirichlet de l'exemple 11.1, intégrable au sens de Lebesgue sur $[0, 1]$.

Soit a et b des réels tels que $a < b$. Une classe importante de fonctions intégrables au sens de Riemann sur $[a, b]$ est celle des fonctions réglées sur $[a, b]$, c'est-à-dire les fonctions qui admettent dans $\mathbb{R}$ une limite à droite en tout point de $[a, b[$ et une limite gauche en tout point de $]a, b]$. Une fonction continue ou monotone est réglée sur $[a, b]$. Les fonctions réglées sur $[a, b]$ sont celles dont la restriction à $[a, b]$ est limite uniforme d'une suite de fonctions en escalier sur $[a, b]$. Cependant une fonction qui n'est pas réglée sur $[a, b]$ peut y être intégrable au sens de Riemann.

11.3. Fonction intégrable sur $[0, 1]$ qui n'est pas réglée sur $[0, 1]$.

Nous considérons l'application :

$$\left| \begin{array}{l} f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ \sin \frac{1}{x} & \text{si } x \in]0, 1]. \end{cases} \end{array} \right.$$

La suite $(a_n)_{n \geq 0}$ de terme général $a_n = 1/((\pi/2) + n\pi)$ est à valeurs dans $]0, 1[$ et converge vers 0. Comme $f(a_n) = (-1)^n$ pour tout n , la suite $(f(a_n))$ diverge, donc f n'admet pas de limite à droite en 0. Par suite f n'est pas réglée sur $[0, 1]$.

Soit ε un réel > 0 . Quitte à le remplacer par $\min(1, \varepsilon)$, nous supposons que $\varepsilon \leq 1$. La fonction f étant continue sur $[\varepsilon/4, 1]$, f est intégrable au sens de Riemann sur $[\varepsilon/4, 1]$. Il existe donc des applications φ_0 et ψ_0 de $[\varepsilon/4, 1]$ dans $\mathbb{R}$, en escalier sur le segment $[\varepsilon/4, 1]$, telles que $\varphi_0 \leq f \leq \psi_0$ sur $[\varepsilon/4, 1]$ et $I_{[\varepsilon/4, 1]}(\psi_0 - \varphi_0) \leq \varepsilon/2$.

Nous introduisons les applications :

$$\left| \begin{array}{l} \varphi : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto \varphi(x) = \begin{cases} -1 & \text{si } x \in \left[0, \frac{\varepsilon}{4}\right[, \\ \varphi_0(x) & \text{si } x \in \left[\frac{\varepsilon}{4}, 1\right] \end{cases} \end{array} \right.$$

et :

$$\left| \begin{array}{l} \psi : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto \psi(x) = \begin{cases} 1 & \text{si } x \in \left[0, \frac{\varepsilon}{4}\right[, \\ \psi_0(x) & \text{si } x \in \left[\frac{\varepsilon}{4}, 1\right]. \end{cases} \end{array} \right.$$

Clairement, φ et ψ sont en escalier sur $[0, 1]$ et on a :

$$I_{[0, 1]}(\varphi) = -\frac{\varepsilon}{4} + I_{[\varepsilon/4, 1]}(\varphi_0) \quad \text{et} \quad I_{[0, 1]}(\psi) = \frac{\varepsilon}{4} + I_{[\varepsilon/4, 1]}(\psi_0),$$

donc :

$$\begin{aligned} I_{[0, 1]}(\psi - \varphi) &= I_{[0, 1]}(\psi) - I_{[0, 1]}(\varphi) = \frac{\varepsilon}{2} + I_{[\varepsilon/4, 1]}(\psi_0) - I_{[\varepsilon/4, 1]}(\varphi_0) \\ &= \frac{\varepsilon}{2} + I_{[\varepsilon/4, 1]}(\psi_0 - \varphi_0) \leq \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon. \end{aligned}$$

Enfin, sinus étant minorée par -1 et majorée par 1 , on a, sur $[0, 1]$, $\varphi \leq f \leq \psi$. En conclusion, f est intégrable au sens de Riemann² sur $[0, 1]$. $\square$

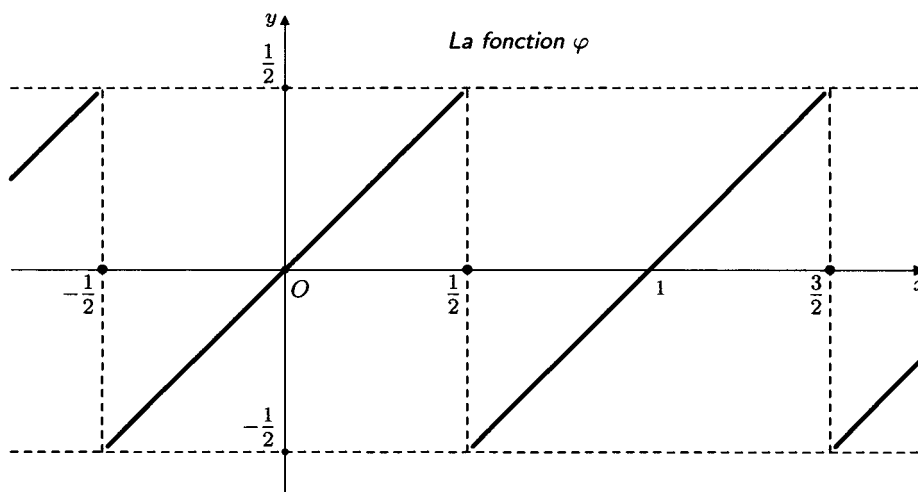
2. Comme $|f(x)| \leq 1$ pour tout $x \in [0, 1]$, il est immédiat de montrer par le critère de majoration que f est intégrable au sens de Lebesgue sur $[0, 1]$.

11.4. Fonction intégrable sur $[0, 1]$ telle que l'ensemble des points où elle est discontinue est dénombrable et dense³ dans $[0, 1]$.

Nous utilisons dans cet exemple les propriétés de la convergence uniforme des suites et séries de fonctions, rappelées aux chapitres 12 et 13.

Nous notons, pour tout réel x n'appartenant pas à $\frac{1}{2} + \mathbb{Z}$, $\delta(x)$ l'entier le plus proche de x , nous introduisons l'application :

$$\left| \begin{array}{l} \varphi : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto \varphi(x) = \begin{cases} 0 & \text{si } x \in \frac{1}{2} + \mathbb{Z}, \\ x - \delta(x) & \text{si } x \notin \frac{1}{2} + \mathbb{Z} \end{cases} \end{array} \right.$$



et nous associons à tout entier $n \geq 1$ l'application :

$$\left| \begin{array}{l} u_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_n(x) = \frac{\varphi(nx)}{n^2}. \end{array} \right.$$

Pour tout $p \in \mathbb{Z}$, φ est affine donc continue sur $]p - (1/2), p + (1/2)[$. Par suite φ est continue en tout point de $\mathbb{R} \setminus ((1/2) + \mathbb{Z})$. De plus, pour tout $p \in \mathbb{Z}$, φ admet pour limite $-1/2$ en $p + (1/2)$ à gauche et $1/2$ en $p + (1/2)$ à droite, donc φ est réglée sur $\mathbb{R}$. Pour tout $n \in \mathbb{N}^*$, l'application $x \mapsto nx$ est strictement croissante sur $\mathbb{R}$, donc la fonction u_n est réglée sur $\mathbb{R}$. Pour tout point x du segment $[0, 1]$ et tout entier $n \geq 1$, $|u_n(x)| \leq 1/(2n^2)$, donc, puisque la série $\sum_{n \geq 1} 1/(2n^2)$ converge, la série de fonctions $\sum u_n$ converge normalement, donc uniformément, sur $[0, 1]$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \sum_{n=1}^{+\infty} u_n(x). \end{array} \right.$$

De plus, pour tout $n \in \mathbb{N}^*$, u_n est intégrable au sens de Riemann sur $[0, 1]$, donc

3. Riemann propose cet exemple en 1854 pour montrer que la théorie de l'intégration qu'il vient d'introduire permet d'intégrer des fonctions qui sont fort loin d'être continues.

la convergence uniforme de la série $\sum u_n$ montre que la fonction f est intégrable au sens de Riemann sur $[0, 1]$.

Pour tout point x de $[0, 1]$, tout nombre entier relatif p et tout entier $n \geq 1$, $nx = p + (1/2)$ si, et seulement si, $x = (p/n) + (1/(2n))$. L'ensemble :

$$\mathcal{A} = [0, 1] \cap \left\{ \frac{p}{n} + \frac{1}{2n} \mid p \in \mathbb{Z} \text{ et } n \in \mathbb{N}^* \right\}$$

est clairement dénombrable et dense dans $[0, 1]$.

Si un point x de $[0, 1]$ n'appartient pas à $\mathcal{A}$, alors, pour tout entier $n \geq 1$, la fonction $t \mapsto \varphi(nt)$ est continue en x , donc aussi u_n , d'où l'on déduit que f est continue en x .

Pour tout entier $n \geq 1$, u_n est réglée sur $[0, 1]$, donc la convergence uniforme permet d'affirmer que f est réglée sur $[0, 1]$ et que, pour tout point a de $]0, 1[$, (resp. de $[0, 1[$), la limite à gauche (resp. à droite) de f en a la somme de la série $\sum \lim_{(x \rightarrow a^-)} u_n(x)$ (resp. $\sum \lim_{(x \rightarrow a^+)} u_n(x)$). En tout point de discontinuité de φ , la limite à gauche est plus grande que la limite à droite, ce qui est donc vérifiée pour toutes les u_n . On en déduit que si $a \in]0, 1[$ et si l'une au moins des fonctions u_n admet une limite à gauche différente de la limite à droite, il en est de même pour la fonction f . Par conséquent f n'est continue en aucun point de $\mathcal{A}$.

En conclusion, l'ensemble des points de discontinuité de la fonction f est $\mathcal{A}$, ensemble dénombrable et dense dans $[0, 1]$. $\square$

Soit a et b des nombres réels tels que $a < b$. L'intégrale d'une fonction f définie, positive et intégrable sur $[a, b]$ est un réel positif ou nul. Si une fonction f est définie, positive et continue sur $[a, b]$ et son intégrale y est nulle, alors f est nulle sur $[a, b]$. Enfin, si une fonction f est définie, positive et continue sur $[a, b]$, et s'il existe au moins un point t_0 de $[a, b]$ tel que $f(t_0) \neq 0$, alors $\int_a^b f(t) dt > 0$.

11.5. Fonction positive et intégrale sur $[0, 1]$ et d'intégrale nulle qui n'est pas identiquement nulle sur $[0, 1]$.

L'application :

$$\left| \begin{array}{l} f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x \in [0, 1[, \\ 1 & \text{si } x = 1 \end{cases} \end{array} \right.$$

est en escalier, donc intégrable sur $[0, 1]$, son intégrale sur $[0, 1]$ est $(1 - 0) \times 0 = 0$, f est positive sur $[0, 1]$ et $f(1) > 0$. $\square$

Cet exemple est trivial, aussi nous imposons à l'ensemble des points en lesquels la valeur de f est strictement positive d'être dense dans $[0, 1]$.

11.6. Fonction positive et intégrale sur $[0, 1]$, d'intégrale nulle sur le segment $[0, 1]$, telle que l'ensemble des points où elle n'est pas nulle est dénombrable et dense dans $[0, 1]$.

Nous considérons l'application, à valeurs réelles positives ou nulles :

$$\left| \begin{array}{l} f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x \text{ est irrationnel ou égal à } 0, \\ \frac{1}{q} & \text{si } x \in \mathbb{Q} \setminus \{0\} \text{ et si } (p, q) \text{ est le} \\ & \text{représentant irréductible de} \\ & x \text{ de dénominateur positif.} \end{cases} \end{array} \right.$$

Nous posons, pour tout réel $\varepsilon > 0$, $A_\varepsilon = \{x \in [0, 1] \mid f(x) \geq \varepsilon\}$. Si ε est un nombre réel > 0 et si nous notons n la partie entière de $1/\varepsilon$, les éléments de l'ensemble A_ε sont des rationnels p/q où p est un entier naturel et q un entier supérieur ou égal à 1 tels que $p \leq q$ et $(1/q) \geq \varepsilon$, ce qui équivaut à $q \leq (1/\varepsilon)$, donc à $q \leq n$; par suite A_ε est fini et de cardinal inférieur ou égal à n^2 .

Nous associons à tout nombre réel $\varepsilon > 0$ l'application :

$$\left| \begin{array}{l} f_\varepsilon : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_\varepsilon(x) = \begin{cases} f(x) & \text{si } x \in A_\varepsilon, \\ \varepsilon & \text{si } x \notin A_\varepsilon \end{cases} \end{array} \right.$$

et la subdivision σ_ε de $[0, 1]$ obtenue en ordonnant les points de $\{0, 1\} \cup A_\varepsilon$. Pour tout réel $\varepsilon > 0$, f_ε est une fonction en escalier sur $[0, 1]$ et σ_ε une subdivision adaptée à f_ε , on a $0 \leq f \leq f_\varepsilon$ sur $[0, 1]$ et $I_{[0,1]}(f_\varepsilon) \leq \varepsilon$. Il en résulte que f est intégrable sur $[0, 1]$ et que, pour tout nombre réel $\varepsilon > 0$, l'intégrale de f sur $[0, 1]$ est positive et inférieure ou égale à celle de f_ε , qui vaut ε , donc :

$$\int_0^1 f(t) dt = 0.$$

De plus l'ensemble des points du segment $[0, 1]$ en lesquels f est strictement positive est $\Omega = \mathbb{Q} \cap]0, 1]$, ensemble dénombrable dense dans $[0, 1]$. $\square$

Dans l'exemple précédent 11.6, l'ensemble des points x de $[0, 1]$ tels que $f(x) > 0$ est infini, mais il n'est « que » dénombrable.

11.7. Fonction positive et intégrable sur $[0, 1]$, d'intégrale nulle sur le segment $[0, 1]$, telle que l'ensemble des points où elle n'est pas nulle est équipotent $\mathbb{R}$.

Nous utilisons la suite $(C_n)_{n \in \mathbb{N}}$ de parties de $\mathbb{R}$ définie dans l'exemple 5.30 (pages 96 et 97) et nous associons à tout entier naturel n l'application :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \begin{cases} 1 & \text{si } x \in C_n, \\ 0 & \text{si } x \notin C_n. \end{cases} \end{array} \right.$$

Si n est un entier naturel, la fonction f_n est en escalier sur $[0, 1]$ et, comme C_n est la réunion de 2^n intervalles deux à deux disjoints d'amplitude $1/3^n$, on obtient :

$$I_{[0,1]}(f_n) = \frac{2^n}{3^n} = \left(\frac{2}{3}\right)^n.$$

Nous notons C l'intersection de la famille $(C_n)_{n \in \mathbb{N}}$ — c'est l'ensemble triadique de Cantor — et nous considérons l'application, à valeurs réelles positives ou nulles :

$$\left| \begin{array}{l} f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 1 & \text{si } x \in C, \\ 0 & \text{si } x \notin C. \end{cases} \end{array} \right.$$

Comme C est inclus dans tous les C_n , on a $0 \leq f \leq f_n$ sur $[0, 1]$ pour tout entier naturel n . Or la suite de terme général $I_{[0,1]}(f_n)$ converge vers 0, donc f est intégrable au sens de Riemann sur le segment $[0, 1]$ et :

$$\int_0^1 f(t) dt = 0.$$

Pourtant $f(x) = 1$ pour tout point x de l'ensemble C et nous avons établi dans l'exemple 5.30 que C est équipotent à $\mathbb{R}$. $\square$

La composée $g \circ f$ d'une fonction f intégrable au sens de Riemann et d'une fonction g continue est intégrable au sens de Riemann. Si l'application g est seulement intégrable, la conclusion n'est plus nécessairement vérifiée.

11.8. Composée de deux fonctions intégrables qui n'est pas une fonction intégrable.

Nous reprenons l'application f définie dans l'exemple 11.6 (pages 211 et 212) et nous introduisons l'application :

$$\left| \begin{array}{l} g : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} 0 & \text{si } x = 0, \\ 1 & \text{si } x \in]0, 1]. \end{cases} \end{array} \right.$$

La fonction f est intégrable sur le segment $[0, 1]$ — voir l'exemple 11.6 — et à valeurs dans $[0, 1]$ et g est en escalier donc intégrable sur $[0, 1]$. Pour tout t appartenant à $\mathcal{Q} = \mathbb{Q} \cap]0, 1]$, $f(t) \in]0, 1]$ donc $g \circ f(t) = 1$, et, pour tout point t de $\mathcal{A} = \mathbb{C}_{[0,1]} \setminus \mathcal{Q} = \{0\} \cup ([0, 1] \setminus \mathbb{Q})$, $f(t) = 0$ donc $g \circ f(t) = 0$. On prouve alors comme dans l'exemple 11.1 (page 208) que $g \circ f$ n'est pas intégrable sur $[0, 1]$. $\square$

■ Convergence des intégrales

DÉFINITION 11.2. — Une fonction f définie sur un intervalle I est localement intégrable sur I si elle est intégrable sur tout segment inclus dans I .

Si a et b sont des nombres réels tels que $a < b$ et f une fonction définie et intégrable sur le segment $[a, b]$, la fonction $x \mapsto \int_a^x f(t) dt$ est continue sur $[a, b]$, donc :

$$\lim_{x \rightarrow b^-} \int_a^x f(t) dt = \int_a^b f(t) dt.$$

Cette remarque justifie les définitions suivantes.

DÉFINITION 11.3. — Soit a et b des points de $\overline{\mathbb{R}}$ tels que $-\infty < a < b \leq +\infty$ et f une fonction définie et localement intégrable sur l'intervalle $[a, b[$.

a) On dit que l'intégrale $\int_a^b f(t) dt$ converge si $\int_a^x f(t) dt$ admet une limite dans $\mathbb{R}$ quand x tend vers b à gauche (vers $+\infty$ si $b = +\infty$) et, si l'intégrale $\int_a^b f(t) dt$ converge, l'intégrale de f sur l'intervalle $[a, b[$ est le nombre réel :

$$\int_a^b f(t) dt = \lim_{x \rightarrow b^-} \int_a^x f(t) dt.$$

b) L'intégrale $\int_a^b f(t) dt$ diverge si elle ne converge pas.

De même, si a et b sont des points de $\overline{\mathbb{R}}$ tels que $-\infty \leq a < b < +\infty$ et f une fonction définie et localement intégrable sur l'intervalle $]a, b]$, on dit que l'intégrale $\int_a^b f(t) dt$ converge si $\int_x^b f(t) dt$ admet une limite dans $\mathbb{R}$ quand x tend vers a à droite (vers $-\infty$ si $a = -\infty$) et, si l'intégrale $\int_a^b f(t) dt$ converge, l'intégrale de f sur l'intervalle $]a, b]$ est le réel $\int_a^b f(t) dt = \lim_{(x \rightarrow a^+)} \int_x^b f(t) dt$; enfin, l'intégrale $\int_a^b f(t) dt$ diverge si elle ne converge pas.

Le théorème et la définition qui suivent concernent des points a et b de $\overline{\mathbb{R}}$ tels que $-\infty < a < b \leq +\infty$ et une fonction f définie et localement intégrable sur $[a, b[$.

Ils s'appliquent sans changement au cas où a et b sont des points de $\overline{\mathbb{R}}$ tels que $-\infty \leq a < b < +\infty$ et f une fonction définie et localement intégrable sur $]a, b[$.

THÉORÈME 11.1. — Si l'intégrale $\int_a^b |f(t)| dt$ converge, alors l'intégrale $\int_a^b f(t) dt$ converge.

DÉFINITION 11.4. — a) La fonction f est intégrable sur $[a, b[$, ou encore l'intégrale $\int_a^b f(t) dt$ est absolument convergente, si l'intégrale $\int_a^b |f(t)| dt$ converge.

b) L'intégrale $\int_a^b f(t) dt$ est semi-convergente si l'intégrale $\int_a^b f(t) dt$ converge alors que l'intégrale $\int_a^b |f(t)| dt$ diverge.

L'implication du théorème 11.1 n'est pas une équivalence.

11.9. Fonction localement intégrable sur $[0, +\infty[$ dont l'intégrale est semi-convergente.

En posant $\frac{\sin 0}{0} = 1$, l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \frac{\sin x}{x} \end{array} \right.$$

est continue sur $\mathbb{R}$.

Nous considérons la suite $(u_n)_{n \geq 0}$ de terme général $u_n = \int_{n\pi}^{(n+1)\pi} f(x) dx$.

La fonction f garde un signe constant sur chacun des segments $[n\pi, (n+1)\pi]$, donc l'intégrabilité de f sur $[0, +\infty[$ équivaut à la convergence de la série $\sum |u_n|$, et l'intégrale $\int_0^{+\infty} f(x) dx$ converge si, et seulement si, la série $\sum u_n$ est convergente.

Soit $n \in \mathbb{N}$. On a, pour tout $x \in [n\pi + \frac{\pi}{6}, (n+1)\pi - \frac{\pi}{6}]$, $\frac{\pi}{6} \leq x - n\pi \leq \frac{5\pi}{6}$, donc :

$$|f(x)| = \frac{1}{x} |\sin x| = \frac{1}{x} |\sin(x - n\pi)| \geq \frac{1}{(n+1)\pi} \times \frac{1}{2} = \frac{1}{2\pi(n+1)}$$

ce qui donne :

$$\begin{aligned} |u_n| &= \int_{n\pi}^{(n+1)\pi} |f(x)| dx \geq \int_{n\pi + \frac{\pi}{6}}^{(n+1)\pi - \frac{\pi}{6}} |f(x)| dx \\ &\geq \int_{n\pi + \frac{\pi}{6}}^{(n+1)\pi - \frac{\pi}{6}} \frac{1}{2\pi(n+1)} dx = \frac{2\pi}{3} \times \frac{1}{2\pi(n+1)} = \frac{1}{3(n+1)} = v_n. \end{aligned}$$

Comme la série $\sum v_n$ diverge, il en est de même de la série $\sum |u_n|$, donc l'intégrale $\int_0^{+\infty} |f(x)| dx$ diverge.

Pour tout entier naturel n , le changement de variable $x = t + n\pi$ nous donne :

$$u_n = \int_{n\pi}^{(n+1)\pi} \frac{\sin x}{x} dx = (-1)^n w_n \text{ où } w_n = \int_0^\pi \frac{\sin t}{t + n\pi} dt.$$

Pour tout entier $n \geq 1$, on a $0 \leq \frac{\sin t}{t + n\pi} \leq \frac{1}{n\pi}$ pour tout point t de $[0, \pi]$, donc :

$$0 \leq w_n \leq \int_0^\pi \frac{1}{n\pi} dt = \frac{1}{n}.$$

De plus, pour tout $n \in \mathbb{N}$, on a $0 \leq \frac{\sin t}{t + (n+1)\pi} \leq \frac{\sin t}{t + n\pi}$ pour tout $t \in [0, \pi]$,

ce qui montre que $0 \leq w_{n+1} \leq w_n$. On déduit du critère des séries alternées⁴ que la série $\sum u_n$ converge.

En conclusion, l'intégrale $\int_0^{+\infty} \frac{\sin x}{x} dx$ est semi-convergente. $\square$

Nous donnons un autre exemple du même type, dans le cas d'un intervalle borné.

11.10. Fonction localement intégrable sur $]0, 1]$ dont l'intégrale est semi-convergente.

Nous considérons l'application :

$$\left| \begin{array}{ll} g :]0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = x^2 \sin \frac{\pi}{x^2} \end{array} \right.$$

et sa fonction dérivée $f = g'$ sur $]0, 1]$. Alors, pour tout point x de $]0, 1]$:

$$f(x) = \underbrace{2x \sin \frac{\pi}{x^2}}_{= f_1(x)} - \underbrace{\frac{2\pi}{x} \cos \frac{\pi}{x^2}}_{= f_2(x)}.$$

Comme $|g(x)| \leq x^2$ pour tout $x \in]0, 1]$, $g(x)$ admet pour limite 0 quand x tend vers 0 à droite. Or, pour tout point x de $]0, 1]$, $\int_x^1 f(t) dt = g(1) - g(x) = -g(x)$, donc l'intégrale $\int_0^1 f(t) dt$ converge et $\int_0^1 f(t) dt = 0$.

Pour tout point x de $]0, 1]$, $|f_1(x)| \leq 2x$, donc $f_1(x)$ admet pour limite 0 quand x tend vers 0 à droite. Nous prolongeons f_1 au segment $[0, 1]$ en posant $f_1(0) = 0$. Alors f_1 est continue sur $[0, 1]$, donc aussi $|f_1|$, d'où l'on déduit que l'intégrale $\int_0^1 |f_1(t)| dt$ converge. Nous travaillons maintenant sur la fonction $|f_2|$.

Nous introduisons les suites $(a_n)_{n \geq 1}$ et $(b_n)_{n \geq 1}$, de termes généraux :

$$a_n = \frac{1}{\sqrt{n + \frac{1}{3}}} \text{ et } b_n = \frac{1}{\sqrt{n - \frac{1}{3}}}.$$

Soit n un entier ≥ 2 . On a alors $0 < a_n < b_n < 1$ et, pour tout point x de $[a_n, b_n]$:

$$n\pi - \frac{\pi}{3} \leq \frac{\pi}{x^2} \leq n\pi + \frac{\pi}{3}, \text{ donc } -\frac{\pi}{3} \leq \frac{\pi}{x^2} - n\pi \leq \frac{\pi}{3},$$

ce qui montre que :

$$\left| \cos \frac{\pi}{x^2} \right| = \left| \cos \left(\frac{\pi}{x^2} - n\pi \right) \right| \geq \frac{1}{2}, \text{ donc } |f_2(x)| = \frac{2\pi}{x} \left| \cos \frac{\pi}{x^2} \right| \geq \frac{\pi}{x}.$$

Il en résulte que $\int_{a_n}^{b_n} |f_2(x)| dx \geq \pi \int_{a_n}^{b_n} \frac{1}{x} dx = \pi \ln \frac{a_n}{b_n} = \frac{\pi}{2} \times \underbrace{\ln \frac{n + \frac{1}{3}}{n - \frac{1}{3}}}_{= u_n}.$

La suite de terme général $\frac{n + (1/3)}{n - (1/3)}$ tend vers 1, donc :

$$u_n \underset{n \rightarrow \infty}{\sim} \frac{n + \frac{1}{3}}{n - \frac{1}{3}} - 1 = \frac{2}{3} \frac{1}{n - \frac{1}{3}} \underset{n \rightarrow \infty}{\sim} \frac{2}{3n}.$$

4. Théorème 7.3, page 112.

Par conséquent la série $\sum u_n$ diverge. On a, pour tout entier $n \geq 2$:

$$0 < a_n < b_n < a_{n-1} < b_{n-1} < \dots < a_p < b_p < \dots < a_2 < b_2 < 1$$

donc :

$$\int_{a_n}^1 |f_2(t)| dt \geq \sum_{p=2}^n \int_{a_p}^{b_p} |f_2(t)| dt \geq \pi \sum_{p=2}^n u_p,$$

ce qui montre que la suite de terme général $\int_{a_n}^1 |f_2(t)| dt$ tend vers $+\infty$.

Comme $\lim_{(n \rightarrow \infty)} a_n = 0^+$, $\int_x^1 |f_2(t)| dt$ n'admet pas de limite dans $\mathbb{R}$ quand x tend vers 0 à droite.

On a ainsi prouvé que l'intégrale $\int_0^1 |f_2(t)| dt$ diverge. Sur $]0, 1]$, $f_2 = f_1 - f$ donc :

$$|f_2| \leq |f_1| + |f|.$$

Or l'intégrale $\int_0^1 |f_1(t)| dt$ converge donc, si l'intégrale $\int_0^1 |f(t)| dt$ convergeait, l'intégrale $\int_0^1 |f_2(t)| dt$ serait convergente, en contradiction avec ce qui précède.

En conclusion, l'intégrale $\int_0^1 f(t) dt$ est semi-convergente. $\square$

Pour une fonction f localement intégrable, positive et décroissante sur un intervalle $[a, +\infty[$ où a est un nombre réel, l'intégrabilité de f sur $[a, +\infty[$ entraîne $\lim_{+\infty} f = 0$. Montrons d'une part que la réciproque est fautive et d'autre part la nécessité de la décroissance de f pour conclure.

11.11. Fonction localement intégrable, positive et décroissante sur $[1, +\infty[$, de limite nulle en $+\infty$, qui n'est pas intégrable sur l'intervalle $[1, +\infty[$.

L'application :

$$\left| \begin{array}{l} f : [1, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \frac{1}{x} \end{array} \right.$$

est continue, positive et décroissante sur $[1, +\infty[$ et $\lim_{+\infty} f = 0$. On a, pour tout nombre réel $x \geq 1$:

$$\int_1^x \frac{1}{t} dt = \ln x.$$

Or $\lim_{+\infty} \ln = +\infty$, donc la fonction f n'est pas intégrable sur $[1, +\infty[$. $\square$

Avant de détailler l'exemple suivant, remarquons que si a est un réel et f une fonction définie, localement intégrable et positive sur $[a, +\infty[$, et si f admet une limite λ non nulle en $+\infty$, f n'est pas intégrable sur $[a, +\infty[$. En effet, la limite λ est alors strictement positive, donc, en choisissant un réel μ tel que $0 < \mu < \lambda$, il existe un réel $b \geq a$ tel que $f(x) \geq \mu$ pour tout $x \geq b$, ce qui montre que, pour tout réel $x \geq b$, $\int_b^x f(t) dt \geq \mu(x - b)$, d'où l'on déduit que $\lim_{x \rightarrow +\infty} \int_a^x f(t) dt = +\infty$.

11.12. Fonction localement intégrable et positive sur $[0, +\infty[$, qui est intégrable sur $[0, +\infty[$ et qui ne tend pas vers 0 en $+\infty$.

Nous posons, pour tout entier $n \geq 2$:

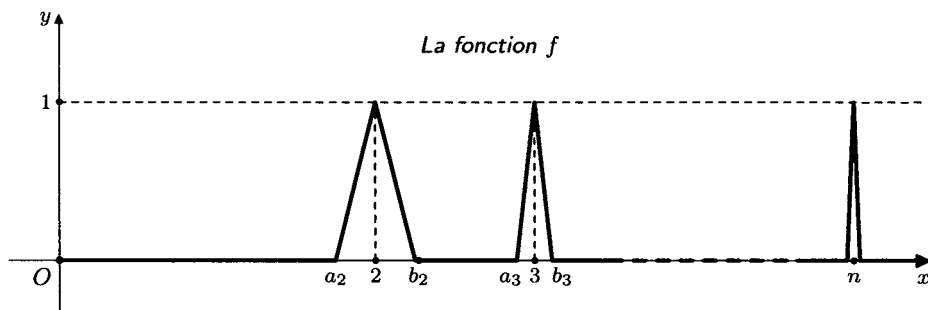
$$a_n = n - \frac{1}{n^2}, \quad b_n = n + \frac{1}{n^2} \quad \text{et} \quad S_n = [a_n, b_n],$$

et nous associons à tout entier $n \geq 2$ l'application :

$$\left| \begin{array}{l} \varphi_n : S_n \longrightarrow \mathbb{R} \\ x \longmapsto \varphi_n(x) = \begin{cases} n^2x - n^3 + 1 & \text{si } a_n \leq x < n, \\ -n^2x + 1 + n^3 & \text{si } n \leq x \leq b_n. \end{cases} \end{array} \right.$$

La famille $(S_n)_{n \geq 2}$ est une suite segments deux à deux disjoints, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : [0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} \varphi_n(x) & \text{si } x \in S_n \text{ où } n \text{ est un entier } \geq 2, \\ 0 & \text{si } x \notin \bigcup_{n \geq 2} S_n. \end{cases} \end{array} \right.$$



Pour tout entier $n \geq 2$, $\varphi_n(a_n) = 0$, $n^2 \times n - n^3 + 1 = 1 = \varphi_n(n)$ et $\varphi_n(b_n) = 0$, donc f est clairement continue sur $[0, +\infty[$; de plus f est à valeurs réelles positives ou nulles. On a, pour tout entier $n \geq 2$:

$$u_n = \int_{a_n}^{b_n} f(t) dt = \frac{1}{n^2} \quad (\text{c'est l'aire du triangle de base } [a_n, b_n]).$$

L'intégrabilité de f sur $[0, +\infty[$ découle de la convergence de la série $\sum_{n \geq 2} u_n$. De plus les suites $(a_n)_{n \geq 2}$ et $(n)_{n \geq 2}$ tendent vers $+\infty$ alors que les suites $(f(a_n))$ et $(f(n))$ convergent respectivement vers 0 et vers 1, donc la fonction f n'admet pas de limite en $+\infty$. $\square$

On peut améliorer cet exemple en imposant à f de n'être bornée sur aucun voisinage de $+\infty$.

11.13. Fonction localement intégrable et positive sur $[0, +\infty[$, qui est intégrable sur $[0, +\infty[$ et qui n'est bornée sur aucun voisinage de $+\infty$.

Nous posons, pour tout entier $n \geq 2$:

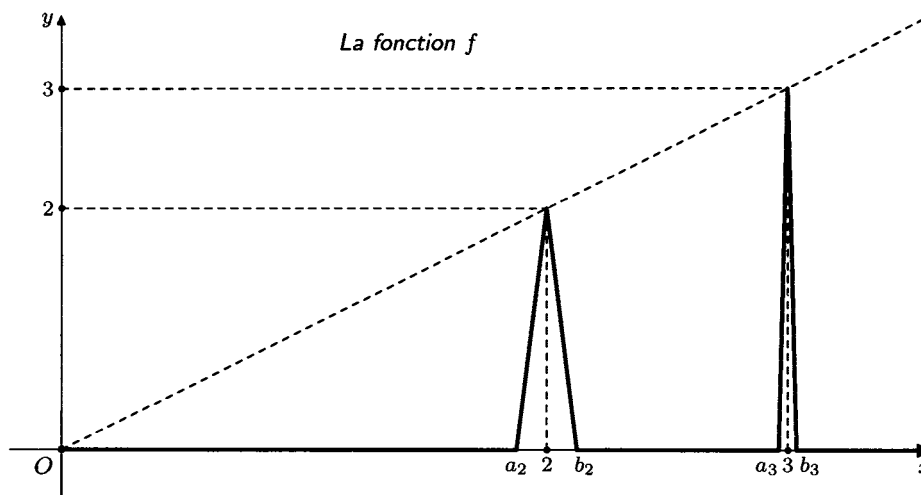
$$a_n = n - \frac{1}{n^3}, \quad b_n = n + \frac{1}{n^3} \quad \text{et} \quad S_n = [a_n, b_n],$$

et nous associons à tout entier $n \geq 2$ l'application :

$$\left| \begin{array}{l} \varphi_n : S_n \longrightarrow \mathbb{R} \\ x \longmapsto \varphi_n(x) = \begin{cases} n^4x - n^5 + n & \text{si } a_n \leq x < n, \\ -n^4x + n^5 + n & \text{si } n \leq x \leq b_n. \end{cases} \end{array} \right.$$

La famille $(S_n)_{n \geq 2}$ est une suite segments deux à deux disjoints, ce qui justifie l'existence de l'application :

$$f : [0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} \varphi_n(x) & \text{si } x \in S_n \text{ où } n \text{ est un entier } \geq 2, \\ 0 & \text{si } x \notin \bigcup_{n \geq 2} S_n. \end{cases}$$



On prouve comme dans l'exemple précédent 11.12 que f est continue sur $[0, +\infty[$ et à valeurs réelles positives ou nulles. On a, pour tout entier $n \geq 2$, $f(n) = n$ et :

$$u_n = \int_{a_n}^{b_n} f(t) dt = \frac{1}{n^2} \quad (\text{c'est l'aire du triangle de base } [a_n, b_n]).$$

L'intégrabilité de f sur $[0, +\infty[$ découle de la convergence de la série $\sum_{n \geq 2} u_n$.

Les suites $(a_n)_{n \geq 2}$ et $(b_n)_{n \geq 2}$ tendent vers $+\infty$, la suite $(f(a_n))$ converge vers 0 et la suite $(f(b_n))$ tend vers $+\infty$, donc f n'admet pas de limite dans $\mathbb{R}$ en $+\infty$ et ne tend pas vers $+\infty$ en $+\infty$, et f n'est bornée sur aucun voisinage de $+\infty$. $\square$

Le produit de deux fonctions intégrables au sens de Riemann est intégrable au sens de Riemann, ce qui est dû en particulier au fait que les fonctions intégrables au sens de Riemann sont bornées. Ceci devient faux pour l'intégrabilité sur un intervalle autre qu'un segment.

11.14. Fonction positive et intégrable dont le carré n'est pas intégrable.

Nous considérons l'application :

$$f :]0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \frac{1}{\sqrt{x}}.$$

La fonction f est continue sur $]0, 1]$ et à valeurs dans $\mathbb{R}_+$ et, pour tout $x \in]0, 1]$:

$$\int_x^1 f(t) dt = 2 [\sqrt{t}]_{t=x}^{t=1} = 2(1 - \sqrt{x}),$$

qui admet pour limite 2 dans $\mathbb{R}$ quand x tend vers 0 à droite, donc f est intégrable sur $]0, 1]$ et $\int_0^1 f(t) dt = 2$. Nous posons $g = f^2 = f \times f$. Alors $g(x) = 1/x$ pour tout point x de $]0, 1]$. Il en résulte que, pour tout $x \in]0, 1]$:

$$\int_x^1 g(t) dt = \int_x^1 \frac{1}{t} dt = [\ln t]_{t=x}^{t=1} = -\ln x,$$

qui tend vers $+\infty$ quand x tend vers 0 à droite, ce qui montre que $g = f^2$ n'est pas intégrable⁵ sur l'intervalle $]0, 1]$. $\square$

■ Primitives et intégrales

Soit I un intervalle et f une fonction définie sur I .

THÉORÈME 11.2. — Si la fonction f est continue sur I , alors, pour tout point a de I , l'application de I dans $\mathbb{R}$:

$$F_a : x \mapsto F_a(x) = \int_a^x f(t) dt$$

est définie et dérivable sur I et, pour tout point x de I , $F_a'(x) = f(x)$.

Les applications F_a ci-dessus sont appelées les intégrales indéfinies de f et, si f est continue sur I , elles sont de classe $\mathcal{C}^1$ sur I . On dispose d'une réciproque.

THÉORÈME 11.3. — Si g est une fonction définie et de classe $\mathcal{C}^1$ sur I , alors, pour tout point a de I , on a, pour tout $x \in I$:

$$g(x) = g(a) + \int_a^x g'(t) dt$$

Ainsi, pour les fonctions définies et de classe $\mathcal{C}^1$ sur I , la dérivation est, à une constante près, l'opération inverse de l'intégrale indéfinie.

DÉFINITION 11.5. — Une primitive de f sur I est une application F de I dans $\mathbb{R}$, dérivable sur I , telle que $F'(x) = f(x)$ pour tout point x de I .

Si f admet des primitives, deux primitives de f sur I diffèrent d'une constante.

11.15. Intégrale indéfinie qui n'est pas dérivable en au moins un point de l'intervalle.

Nous introduisons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x < 0, \\ 1 & \text{si } x \geq 0. \end{cases} \end{array} \right.$$

La fonction f est réglée sur $\mathbb{R}$, ce qui justifie l'existence de l'intégrale indéfinie :

$$\left| \begin{array}{l} F : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto F(x) = \int_0^x f(t) dt. \end{array} \right.$$

5. On peut aussi appliquer le critère de Riemann, qui affirme que la fonction $x \mapsto \frac{1}{x^\alpha}$ est intégrable sur $]0, 1]$ si, et seulement si, $\alpha < 1$.

Si x est un réel, on a $F(x) = \int_0^x 1 \, dt = x$ si $x \geq 0$ et $F(x) = \int_0^x 0 \, dt = 0$ si $x < 0$.

La fonction F est donc dérivable à gauche et à droite en 0, $F_g'(0) = 0$ et $F_d'(0) = 1$, ce qui montre que F n'est pas dérivable en 0. $\square$

11.16. Intégrale indéfinie qui n'est dérivable en aucun point d'un ensemble dénombrable dense dans l'intervalle.

Nous utilisons la fonction f définie dans l'exemple 10.1 (pages 190 et 191). C'est une application de $[0, 1]$ dans $\mathbb{R}$, strictement croissante sur $[0, 1]$ et discontinue en tous les rationnels appartenant à $]0, 1[$. La fonction f étant croissante sur $[0, 1]$, f est réglée sur $[0, 1]$; elle admet donc une limite à gauche $f(a^-)$ dans $\mathbb{R}$ en tout point de $]0, 1]$ et une limite à droite $f(a^+)$ dans $\mathbb{R}$ en tout point de $[0, 1[$.

La discontinuité de f en tous les points de l'ensemble $\mathcal{Q} = \mathbb{Q} \cap]0, 1[$ montre que, pour tout élément a de $\mathcal{Q}$, $f(a^-) \neq f(a^+)$. On en déduit que la fonction :

$$F : x \mapsto F(x) = \int_0^x f(t) \, dt$$

est dérivable à gauche et à droite en tout point de $\mathcal{Q}$ et que, pour tout point a de $\mathcal{Q}$, $F_g'(a) = f(a^-) \neq f(a^+) = F_d'(a)$, ce qui montre que F n'est pas dérivable en a . Enfin, $\mathcal{Q} = \mathbb{Q} \cap]0, 1[$ est dénombrable et dense dans $[0, 1]$. $\square$

11.17. Intégrale indéfinie F d'une fonction f dérivable en 0 telle que $F'(0) \neq f(0)$.

Nous choisissons une fonction g définie et continue sur le segment $[-1, 1]$ — il y en a ! — et nous définissons l'application :

$$\left| \begin{array}{l} f : [-1, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} g(x) & \text{si } x \neq 0, \\ 1 + g(0) & \text{si } x = 0. \end{cases} \end{array} \right.$$

La fonction f est réglée sur $[-1, 1]$, ce qui justifie l'existence de l'intégrale indéfinie :

$$\left| \begin{array}{l} F : [-1, 1] \longrightarrow \mathbb{R} \\ x \longmapsto F(x) = \int_0^x f(t) \, dt. \end{array} \right.$$

Comme on ne change pas l'intégrabilité ni l'intégrale d'une fonction en modifiant sa valeur en un point, on a, pour tout point x de $[-1, 1]$:

$$F(x) = \int_0^x g(t) \, dt$$

donc F est dérivable sur $[-1, 1]$ et $F'(x) = g(x)$ pour tout point x de $[-1, 1]$. En particulier $F'(0) = g(0) \neq f(0)$. $\square$

11.18. Intégrale indéfinie F d'une fonction f dérivable sur son intervalle de définition telle que $F'(x) \neq f(x)$ pour tous les points x d'un ensemble dense dans cet intervalle.

Nous reprenons l'application f de l'exemple 11.6 (pages 211 et 212). Nous avons vu que f est intégrable sur $[0, 1]$ et que son intégrale y est nulle. Si a et b sont des réels tels que $a < b$ et si une fonction est définie et intégrable au sens de Riemann

sur $[a, b]$, cette fonction est localement intégrable sur $[a, b]$. Ceci justifie l'existence de l'intégrale indéfinie :

$$\left| \begin{array}{l} F : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto F(x) = \int_0^x f(t) dt. \end{array} \right.$$

Comme f est positive ou nulle sur $[0, 1]$, F est croissante sur $[0, 1]$. Il en résulte que, pour tout point x de $[0, 1]$:

$$0 = F(0) \leq F(x) \leq F(1) = \int_0^1 f(t) dt = 0$$

donc $F(x) = 0$. Par conséquent F est identiquement nulle sur $[0, 1]$, donc F est dérivable sur cet intervalle et $F'(x) = 0$ pour tout $x \in [0, 1]$. On en déduit que $F'(x) \neq f(x)$ pour tous les points x de $\mathcal{Q} = \mathbb{Q} \cap]0, 1[$, ensemble dense dans $[0, 1]$. $\square$

Nous allons nous intéresser au problème inverse, c'est-à-dire à l'intégrabilité de la fonction dérivée d'une fonction f dérivable sur un intervalle et, dans l'affirmative, nous allons comparer, a étant un point de cet intervalle, les fonctions f et :

$$g : x \mapsto g(x) = \int_a^x f'(t) dt.$$

Si une fonction f est définie et dérivable sur un intervalle I et si f' est localement intégrable sur I , alors, quels que soient les points a et x de I , $g(x) = f(x) - f(a)$. Cependant, si f n'est dérivable que presque partout, ceci n'empêche pas de calculer g , mais alors l'égalité devient fausse.

11.19. Fonction dérivable sur $] -1, 1[$ dont la fonction dérivée n'est intégrable sur aucun segment $[-a, a]$ où a est un réel tel que $0 < a < 1$.

Nous introduisons l'application, de classe $\mathcal{C}^\infty$ sur $] -1, 0[$ et sur $] 0, 1[$:

$$\left| \begin{array}{l} f :] -1, 1[\longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ \frac{x}{\ln|x|} \cos \frac{1}{x} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

On a, pour tout réel $x \neq 0$:

$$\mathcal{Q}f(0, x) = \frac{f(x) - f(0)}{x - 0} = \frac{f(x)}{x} \text{ donc } |\mathcal{Q}f(0, x)| = \frac{\left| \cos \frac{1}{x} \right|}{|\ln|x||} \leq \frac{1}{|\ln|x||}.$$

Par suite $\mathcal{Q}f(0, \cdot)$ admet pour limite 0 en 0 donc f est dérivable en 0 et $f'(0) = 0$. La fonction f est dérivable en tout point de $] -1, 1[\setminus \{0\}$ et le calcul donne, pour tout $x \in] -1, 1[\setminus \{0\}$:

$$f'(x) = \frac{\cos \frac{1}{x}}{\ln|x|} - \frac{\cos \frac{1}{x}}{\ln^2|x|} + \frac{\sin \frac{1}{x}}{x \ln|x|}.$$

Soit a un nombre réel tel que $0 < a < 1$. Des trois termes composant l'expression de $f'(x)$ pour $x \neq 0$, les deux premiers sont, en les prolongeant par 0 en 0, des expressions continues de x sur $[-a, a]$, donc bornées et intégrables au sens de Riemann sur le segment $[-a, a]$.

Nous considérons la fonction du troisième terme prolongée en 0 ; c'est l'application :

$$\left| \begin{array}{l} g :]-1, 1[\longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} 0 & \text{si } x = 0, \\ \frac{\sin \frac{1}{x}}{x \ln|x|} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

La suite $(x_n)_{n \geq 0}$ de terme général $x_n = 2/(\pi + 4n\pi)$ est à valeurs dans $]0, 1[$ et converge vers 0. On a, pour tout entier naturel n , $g(x_n) = 1/(x_n \ln(x_n))$, terme général d'une suite qui tend vers $-\infty$, donc la fonction g n'est bornée sur aucun voisinage de 0. Par conséquent f' n'est pas bornée sur $[-a, a]$, donc la fonction dérivée f' n'est pas intégrable au sens de Riemann sur le segment $[-a, a]$.

Nous allons démontrer qu'en fait la fonction g , donc aussi f' , n'est pas non plus intégrable au sens de Lebesgue sur le segment $[0, a]$, donc sur $[-a, a]$, en constatant que la fonction $g^+ : x \mapsto g^+(x) = \text{Max}(g(x), 0)$ ne l'est pas. Nous introduisons les suites $(y_n)_{n \geq 0}$, $(z_n)_{n \geq 0}$, $(u_n)_{n \geq 0}$ et $(J_n)_{n \geq 0}$, de termes généraux :

$$y_n = \frac{1}{\frac{7\pi}{6} + 2n\pi}, \quad z_n = \frac{1}{\frac{11\pi}{6} + 2n\pi}, \quad u_n = y_n - z_n \quad \text{et} \quad J_n = \int_{z_n}^{y_n} g^+(x) dx.$$

Soit n un entier naturel. On a $0 < z_n < y_n < 1/2$ et, pour tout point x de $[z_n, y_n]$, $(7\pi)/6 + 2n\pi \leq 1/x \leq (11\pi)/6 + 2n\pi$, donc $\sin(1/x) \leq -1/2$, ce qui montre que :

$$g^+(x) = g(x) \geq \frac{1}{-2x \ln x} = \frac{1}{2x \ln \frac{1}{x}} > 0.$$

On remarque que $y_n < \frac{1}{e}$ et que la fonction $x \mapsto x \ln \frac{1}{x}$ est croissante sur $]0, \frac{1}{e}]$, d'où l'inégalité :

$$J_n \geq \int_{z_n}^{y_n} \frac{1}{2y_n \ln \frac{1}{y_n}} dx = \frac{u_n}{2y_n \ln \frac{1}{y_n}}.$$

Or :

$$u_n = \frac{6}{7\pi + 12n\pi} - \frac{6}{11\pi + 12n\pi} = \frac{24\pi}{(7\pi + 12n\pi)(11\pi + 12n\pi)} \geq \frac{24\pi}{(12(n+1)\pi)^2} = \frac{1}{6\pi(n+1)^2}$$

donc :

$$\begin{aligned} J_n &\geq \frac{\frac{7\pi}{6} + 2n\pi}{12\pi(n+1)^2 \ln \frac{1}{y_n}} \geq \frac{2n\pi}{12\pi(n+1)^2 \ln \frac{1}{y_n}} \\ &\geq \frac{n}{6(n+1)^2 \ln \frac{1}{y_n}} \geq \frac{n}{6(n+1)^2} \times \frac{1}{\ln(2(n+1)\pi)} = v_n. \end{aligned}$$

On a $\ln(2(n+1)\pi) = \ln(2\pi) + \ln(n+1) \underset{n \rightarrow \infty}{\sim} \ln n$, donc $v_n \underset{n \rightarrow \infty}{\sim} \frac{1}{6} \times \frac{1}{n \ln n}$.

Par conséquent la série $\sum v_n$ diverge⁶, donc aussi la série $\sum J_n$. Nous choisissons $n_0 \in \mathbb{N}$ tel que $y_{n_0} \leq a$. On a alors, pour tout entier $n \geq n_0$, $0 < z_n < y_{n_0} \leq a$ et :

$$\int_{z_n}^a g^+(x) dx \geq \int_{z_n}^{y_{n_0}} g^+(x) dx \geq \sum_{p=n_0}^n J_p,$$

donc la suite de terme général $\int_{z_n}^a g^+(x) dx$ tend vers $+\infty$. En conclusion, g^+ n'est pas intégrable au sens de Lebesgue sur $[-a, a]$, donc ni g ni f' ne sont intégrables au sens de Lebesgue sur $[-a, a]$. $\square$

6. En effet son terme général est équivalent à celui d'une série à termes positifs, qui diverge par le critère de Bertrand ou en utilisant le résultat de l'exemple 7.12 (pages 117 et 118).

11.20. Fonction dérivable presque partout sur $[0, 1]$ telle que⁷ :

$$\int_0^1 f'(t) dt \neq f(1) - f(0).$$

Nous reprenons la fonction f définie comme limite uniforme sur $[0, 1]$ d'une suite $(f_n)_{n \geq 0}$ d'applications de $[0, 1]$ dans $\mathbb{R}$, affines par morceaux sur le segment $[0, 1]$, dans l'exemple 10.3 (pages 193 à 196) dont nous utilisons toutes les notations.

Soit $n \in \mathbb{N}^*$. Nous avons montré dans l'exemple 10.3 que $f(y_{k,n}) - f(x_{k,n}) = \frac{1}{2^n}$ pour tout $k \in A_n$ et que :

$$\sum_{k \in A_n} (f(y_{k,n}) - f(x_{k,n})) = \frac{2^n}{2^n} = 1.$$

Comme f est croissante et que $f(0) = 0$ et $f(1) = 1$, on a $f(x_{k+1,n}) = f(x_{k,n}) = 0$ si $k \notin A_n$. On en déduit que pour ces valeurs de k , f est constante sur $]x_{k,n}, x_{k+1,n}[$, donc dérivable et de dérivée nulle sur cet intervalle. Nous posons :

$$J_n = \bigcup_{k \notin A_n}]x_{k,n}, x_{k+1,n}[.$$

L'ensemble J_n est ouvert comme réunion d'intervalles ouverts. La fonction f est donc dérivable et de dérivée nulle sur J_n . De plus J_n est la réunion de $3^n - 2^n$ intervalles disjoints de mesure commune égale à $\frac{1}{3^n}$, donc la mesure de J_n est :

$$\mu_n = \frac{3^n - 2^n}{3^n} = 1 - \left(\frac{2}{3}\right)^n.$$

On déduit de tout ceci que la fonction f est donc dérivable et de dérivée nulle sur

$$J = \bigcup_{n \in \mathbb{N}^*} J_n$$

et, comme la suite (μ_n) converge vers 1, la mesure de J est égale à 1. Il en résulte que f est dérivable presque partout et que sa dérivée est nulle presque partout. Nous avons donc finalement :

$$\int_0^1 f'(t) dt = 0 \neq 1 = f(1) - f(0). \quad \square$$

■ Intégrales dépendant d'un paramètre

Soit I et J des intervalles d'intérieur non vide, les bornes inférieure et supérieure de J dans la droite numérique achevée $\overline{\mathbb{R}}$ étant respectivement a et b .

Soit $f : (x, t) \mapsto f(x, t)$ une fonction à valeurs réelles définie sur $I \times J$ telle que, pour tout $x \in I$, la fonction partielle $f_x : t \mapsto f_x(t) = f(x, t)$ est intégrable sur J . On s'intéresse alors aux propriétés de l'application de I dans $\mathbb{R}$:

$$F : x \mapsto F(x) = \int_a^b f(x, t) dt.$$

7. Ce contre-exemple a été introduit par Cantor pour prouver que l'égalité :

$$\int_0^1 f'(t) dt = f(1) - f(0),$$

valable pour toute fonction continue sur $[0, 1]$, dérivable sur $]0, 1[$ et de dérivée intégrable, ne se généralise pas à des fonctions dérivables presque partout.

THÉORÈME 11.4. — Si, pour tout point x de I , la fonction $f_x : t \mapsto f(x, t) = f(x, t)$ est continue sur J et s'il existe une fonction φ à valeurs dans $\mathbb{R}_+$, intégrable sur J , telle que $|f(x, t)| \leq \varphi(t)$ pour tout $x \in I$ et tout $t \in J$, alors la fonction f_x est, pour tout $x \in J$, intégrable sur J et la fonction :

$$F : x \mapsto F(x) = \int_a^b f(x, t) dt.$$

est définie et continue sur I .

11.21. Intégrale dépendant d'un paramètre qui n'est pas une fonction continue de ce paramètre.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \times [0, 1] \longrightarrow \mathbb{R} \\ (x, t) \longmapsto f(x, t) = \begin{cases} 1 & \text{si } x \text{ est rationnel,} \\ 0 & \text{si } x \text{ est irrationnel.} \end{cases} \end{array} \right.$$

Pour tout nombre réel x , la fonction $f_x : t \mapsto f(x, t) = f(x, t)$ est constante donc intégrable sur le segment $[0, 1]$, et :

$$F(x) = \int_0^1 f(x, t) dt = 1 \text{ si } x \in \mathbb{Q} \text{ et } F(x) = \int_0^1 f(x, t) dt = 0 \text{ si } x \in \mathbb{R} \setminus \mathbb{Q}.$$

Par conséquent l'application :

$$F : x \mapsto F(x) = \int_0^1 f(x, t) dt$$

de $\mathbb{R}$ dans $\mathbb{R}$ est la fonction de Dirichlet, discontinue en tout point de $\mathbb{R}$. $\square$

Cet exemple est élémentaire puisque f n'est pas continue. Donnons deux exemples dans lesquels la fonction est continue, sa valeur absolue n'admettant pas de fonction majorante intégrable.

11.22. Intégrale d'une fonction dépendant d'un paramètre qui n'est pas une fonction continue de ce paramètre alors que la fonction de départ est continue.

Nous introduisons l'application, clairement continue sur $[0, +\infty[\times [0, +\infty[$:

$$\left| \begin{array}{l} f : [0, +\infty[\times [0, +\infty[\longrightarrow \mathbb{R} \\ (x, t) \longmapsto f(x, t) = x e^{-xt}. \end{array} \right.$$

On a $f(0, t) = 0$ pour tout point t de $[0, +\infty[$ et, pour tout réel $x > 0$, la fonction $f_x : t \mapsto f(x, t) = f(x, t)$ est négligeable au voisinage de $+\infty$ devant $t \mapsto 1/t^2$, fonction intégrable sur $[1, +\infty[$, donc f_x est intégrable sur $[0, +\infty[$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} F : [0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto F(x) = \int_0^{+\infty} f(x, t) dt. \end{array} \right.$$

On voit que $F(0) = 0$ et on a, pour tout nombre réel $x > 0$:

$$F(x) = \int_0^{+\infty} x e^{-xt} dt = [-e^{-xt}]_{t=0}^{+\infty} = 0 - (-1) = 1,$$

donc la fonction F n'est pas continue en 0. $\square$

11.23. Exemple du même type que le précédent.

Nous considérons l'ouvert $\Omega = \mathbb{R} \times]0, +\infty[$ de $\mathbb{R}^2$ et l'application :

$$\left| \begin{array}{l} f : \Omega \longrightarrow \mathbb{R} \\ (x, t) \longmapsto f(x, t) = \begin{cases} 0 & \text{si } x = 0, \\ \frac{t}{x^2} e^{-\frac{t}{|x|}} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

Compte tenu des théorèmes concernant les opérations sur les fonctions continues à valeurs dans $\mathbb{R}$, la fonction f est clairement continue en tout point de l'ouvert $(\mathbb{R} \setminus \{0\}) \times]0, +\infty[=]-\infty, 0[\times]0, +\infty[\cup]0, +\infty[\times]0, +\infty[$.

La fonction $\varphi : u \mapsto \varphi(u) = u^3 e^{-u}$ est continue sur $[0, +\infty[$ et à valeurs dans $\mathbb{R}_+$, elle s'annule en 0 et tend vers 0 en $+\infty$, donc elle est bornée sur l'intervalle $[0, +\infty[$, ce qui justifie l'existence d'une constante réelle $K > 0$ telle que $0 \leq \varphi(u) \leq K$ pour tout nombre réel $u \geq 0$.

Soit t_0 un point de $]0, +\infty[$. Nous étudions la continuité de f en $(0, t_0)$. Pour ceci nous choisissons un réel a tel que $0 < a < t_0$. Alors $U = \mathbb{R} \times]a, +\infty[$ est un ouvert de $\mathbb{R}^2$, donc un voisinage du point $(0, t_0)$ dans $\mathbb{R}^2$. Si (x, t) est un point de U , alors $0 < a < t$, $f(x, t) = 0$ pour $x = 0$ et, dans le cas où $x \neq 0$:

$$0 < f(x, t) = \frac{t}{x^2} e^{-\frac{t}{|x|}} = \frac{|x|}{t^2} \times \frac{t^3}{|x|^3} e^{-\frac{t}{|x|}} = \frac{|x|}{t^2} \varphi\left(\frac{t}{|x|}\right) \leq \frac{|x|}{t^2} K \leq K \frac{|x|}{a^2}.$$

On en déduit que, pour tout point (x, t) de U , $0 \leq f(x, t) \leq K \frac{|x|}{a^2}$.

La majoration obtenue étant indépendante de la variable t , $f(x, t)$ admet pour limite $0 = f(0, t_0)$ quand (x, t) tend vers $(0, t_0)$. Par suite f est continue en $(0, t_0)$. Il en résulte que f est continue sur l'ouvert $\Omega = \mathbb{R} \times]0, +\infty[$.

On a $f(0, t) = 0$ pour tout $t > 0$. Si $x \neq 0$, la fonction $f_x : t \mapsto f_x(t) = f(x, t)$ se prolonge par continuité à droite en 0 par $f_x(0) = 0$; de plus f_x est négligeable au voisinage de $+\infty$ devant $t \mapsto 1/t^2$, fonction intégrable sur $[1, +\infty[$, donc f_x est intégrable sur $[0, +\infty[$. Tout ceci justifie l'existence de l'application :

$$\left| \begin{array}{l} F : [0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto F(x) = \int_0^{+\infty} f(x, t) dt. \end{array} \right.$$

On a $F(0) = \int_0^{+\infty} f(0, t) dt = \int_0^{+\infty} 0 dt = 0$. Une intégration par parties donne :

$$\int_0^{+\infty} u e^{-u} du = 1.$$

Soit x un réel non nul. A l'aide du changement de variable $t = |x| \times u$ on obtient :

$$F(x) = \int_0^{+\infty} \frac{t}{x^2} e^{-\frac{t}{|x|}} dt = \frac{1}{|x|} \int_0^{+\infty} u e^{-u} du = \frac{1}{|x|}.$$

Ainsi la fonction F n'est pas continue en 0. $\square$

THÉORÈME 11.5. — Théorème de dérivation sous le signe somme.

Soit f une fonction vérifiant les hypothèses du théorème précédent 11.4, dont on utilise les notations. Si la dérivée partielle de f par rapport à x existe en tout point (x, t) de $I \times J$, si, pour tout point x de I , la fonction dérivée partielle :

$$t \mapsto \frac{\partial}{\partial x}(f(x, t))$$

est continue sur J et s'il existe une fonction ψ à valeurs dans $\mathbb{R}_+$ définie et intégrable sur J telle que, pour tout $x \in I$ et tout $t \in J$:

$$\left| \frac{\partial}{\partial x}(f(x, t)) \right| \leq \psi(t),$$

alors F est de classe $\mathcal{C}^1$ sur I et, pour tout $x \in I$, $F'(x) = \int_a^b \frac{\partial}{\partial x}(f(x, t)) dt$.

11.24. Fonction f de classe $\mathcal{C}^1$ telle que $F : x \mapsto F(x) = \int_a^b f(x, t) dt$ est définie mais n'est pas dérivable.

Nous considérons l'application, clairement continue sur $\mathbb{R} \times [0, +\infty[$:

$$\begin{cases} f : \mathbb{R} \times [0, +\infty[\longrightarrow \mathbb{R} \\ (x, t) \longmapsto f(x, t) = x^2 e^{-t|x|}, \end{cases}$$

Soit t un point de $[0, +\infty[$. On a, pour tout réel $x > 0$, $\varphi_t(x) = f(x, t) = x^2 e^{-tx}$, donc φ_t est dérivable en tout point de $]0, +\infty[$ et $\varphi_t'(x) = (2 - tx)x e^{-tx}$ pour tout réel $x > 0$. De même, pour tout nombre réel $x < 0$, $\varphi_t(x) = f(x, t) = x^2 e^{tx}$, donc φ_t est dérivable en tout point de $]0, +\infty[$ et $\varphi_t'(x) = (2 + tx)x e^{tx}$ pour tout nombre réel $x < 0$. Enfin, pour tout réel $x \neq 0$:

$$\frac{f(x, t) - f(0, t)}{x - 0} = \frac{f(x, t)}{x} = x e^{-t|x|},$$

qui admet pour limite 0 quand x tend vers 0, donc la fonction $\varphi_t : x \mapsto f(x, t)$ est dérivable en 0 et $\varphi_t'(0) = 0$. La dérivée partielle de f par rapport à x existe donc en tout point (x, t) de $\mathbb{R} \times [0, +\infty[$ et, pour tout point (x, t) de $\mathbb{R} \times [0, +\infty[$:

$$g(x, t) = \frac{\partial}{\partial x}(f(x, t)) = (2 - t|x|)x e^{-t|x|}.$$

Clairement, la fonction g est continue sur $\mathbb{R} \times [0, +\infty[$, donc f est de classe $\mathcal{C}^1$ sur $\mathbb{R} \times [0, +\infty[$. On a, pour tout réel $x \neq 0$:

$$F(x) = \int_0^{+\infty} f(x, t) dt = x^2 \int_0^{+\infty} e^{-t|x|} dt = x^2 \left[-\frac{1}{|x|} e^{-t|x|} \right]_{t=0}^{+\infty} = |x|,$$

égalité valable aussi pour $x = 0$. La fonction F n'est donc pas dérivable sur $\mathbb{R}$. $\square$

Remarquons que dans l'exemple précédent 11.24, on obtient, pour tout réel x :

$$\int_0^{+\infty} \frac{\partial}{\partial x}(f(x, t)) dt = \begin{cases} 1 & \text{si } x > 0, \\ 0 & \text{si } x = 0, \\ -1 & \text{si } x < 0, \end{cases}$$

soit par le calcul, soit, pour $x \neq 0$, en appliquant, pour tout $\varepsilon > 0$, le théorème 11.5 de dérivation sous le signe somme sur $[\varepsilon, +\infty[\times [0, +\infty[$ et sur $]-\infty, -\varepsilon] \times [0, +\infty[$.

■ Sommes de Riemann

Soit a et b des nombres réels tels que $a < b$.

DÉFINITION 11.6. — Si $\sigma = (x_0, x_1, x_2, \dots, x_n)$ est une subdivision du segment $[a, b]$, le pas — ou le module — de σ est le réel strictement positif :

$$|\sigma| = \max_{i \in [0, n-1]} (x_{i+1} - x_i).$$

DÉFINITION 11.7. — Une subdivision pointée du segment $[a, b]$ est un couple (σ, τ) où $\sigma = (x_0, x_1, \dots, x_n)$ est une subdivision de $[a, b]$ et $\tau = (t_0, t_1, \dots, t_{n-1})$ une famille de points de $[a, b]$ telle que $x_i \leq t_i \leq x_{i+1}$ pour tout $i \in \llbracket 0, n-1 \rrbracket$.

DÉFINITION 11.8. — Si f est une fonction définie sur le segment $[a, b]$ et si (σ, τ) , où $\sigma = (x_0, x_1, x_2, \dots, x_n)$ et $\tau = (t_0, t_1, \dots, t_{n-1})$, est une subdivision pointée de $[a, b]$, la somme de Riemann associée à f et à (σ, τ) est le nombre réel :

$$R(f, (\sigma, \tau)) = \sum_{i=0}^{n-1} (x_{i+1} - x_i) f(t_i).$$

THÉORÈME 11.6. — Si f est une fonction définie sur $[a, b]$, f est intégrable au sens de Riemann sur le segment $[a, b]$ si, et seulement si, f est bornée sur $[a, b]$ et les sommes de Riemann $R(f, (\sigma, \tau))$ admettent une limite⁸ dans $\mathbb{R}$ quand le pas $|\sigma|$ de σ tend vers 0, et, si f est intégrable au sens de Riemann sur $[a, b]$, l'intégrale de f sur $[a, b]$ est la limite⁹ de $R(f, (\sigma, \tau))$ quand $|\sigma|$ tend vers 0.

11.25. Fonction définie sur un segment dont les sommes de Riemann n'admettent pas de limite dans $\mathbb{R}$.

Il faut évidemment choisir une fonction qui n'est pas intégrable au sens de Riemann. Nous reprenons la fonction f de Dirichlet : $f(x) = 1$ si x est rationnel et $f(x) = 0$ si x est irrationnel. Nous considérons, pour tout entier $n \geq 1$, les subdivisions pointées (σ_n, ρ_n) et (σ_n, τ_n) de $[0, 1]$ obtenues de la manière suivante :

$$\sigma_n = \left(\frac{0}{n}, \frac{1}{n}, \frac{2}{n}, \dots, \frac{i}{n}, \dots, \frac{n}{n} \right), \rho_n = (r_0, r_1, \dots, r_{n-1}) \text{ et } \tau_n = (t_0, t_1, \dots, t_{n-1})$$

où, pour tout $i \in \llbracket 0, n-1 \rrbracket$, r_i est un rationnel et t_i un irrationnel choisis dans le segment $[i/n, (i+1)/n]$. Pour tout entier $n \geq 1$, le pas de la subdivision σ_n est égal à $1/n$, donc la suite $(|\sigma_n|)$ tend vers 0, alors que, pour tout entier $n \geq 1$, $R(f, (\sigma_n, \rho_n)) = 1$ et $R(f, (\sigma_n, \tau_n)) = 0$. Si $R(f, (\sigma, \tau))$ admettait une limite ℓ dans $\mathbb{R}$ quand $|\sigma|$ tend vers 0, on aurait à la fois $\ell = 1$ et $\ell = 0$, ce qui est faux. En conclusion, $R(f, (\sigma, \tau))$ n'admet pas de limite dans $\mathbb{R}$ quand $|\sigma|$ tend vers 0. $\square$

La convergence des sommes de Riemann s'applique à des fonctions intégrables sur un segment. Ceci devient faux sur un intervalle quelconque.

11.26. Fonction intégrable sur un intervalle autre qu'un segment dont les sommes de Riemann n'admettent pas de limite dans $\mathbb{R}$.

Nous avons vu dans l'exemple 11.14 (pages 218 et 219) que l'application f de $]0, 1]$ dans $\mathbb{R}$, définie par $f(x) = 1/\sqrt{x}$, est intégrable sur l'intervalle $]0, 1]$.

8. Si $G : (\sigma, \tau) \mapsto G(\sigma, \tau)$ est une fonction à valeurs réelles définie sur l'ensemble des subdivisions pointées de $[a, b]$ et ℓ un nombre réel, $G(\sigma, \tau)$ admet pour limite ℓ dans $\mathbb{R}$ quand le pas $|\sigma|$ de σ tend vers 0 si, pour tout réel $\varepsilon > 0$, il existe un nombre réel $\delta > 0$ tel que, pour toute subdivision pointée (σ, τ) de $[a, b]$, $|\sigma| < \delta$ entraîne $|G(\sigma, \tau) - \ell| < \varepsilon$.

9. Cauchy définit en 1823 ([CAU2], VINGT-UNIÈME LEÇON, *Intégrale définie*, pages 81 à 84), l'intégrale sur un segment d'une fonction f , définie et continue sur ce segment, comme la limite dans $\mathbb{R}$, quand le pas de la subdivision tend vers 0, de ce que nous appelons maintenant les sommes de Riemann de f , mais sa « preuve » de l'existence de cette limite est très insuffisante ; elle devient correcte si l'on utilise la notion de convergence uniforme, introduite par Eduard Heine en 1872 — quarante-neuf ans plus tard — et le théorème de Heine (théorème 8.3, page 144).

Nous associons à tout entier $n \geq 1$ la subdivision pointée (σ_n, τ_n) de $[0, 1]$, définie par $\sigma_n = (x_0, x_1, x_2, \dots, x_n)$ et $\tau_n = (t_0, t_1, \dots, t_{n-1})$ où :

$$x_i = \frac{i}{n} \text{ pour tout } i \in \llbracket 0, n \rrbracket, \quad t_0 = \frac{1}{n^4} \text{ et } t_i = \frac{i+1}{n} \text{ pour tout } i \in \llbracket 1, n-1 \rrbracket.$$

On a, pour tout entier $n \geq 1$:

$$R(f, (\sigma_n, \tau_n)) = \sum_{i=0}^{n-1} (x_{i+1} - x_i) f(t_i) = \frac{1}{n} \left(n^2 + \sum_{i=1}^{n-1} \sqrt{\frac{n}{i+1}} \right) \geq n,$$

donc la somme de Riemann $R(f, (\sigma_n, \tau_n))$ tend vers $+\infty$, alors que la suite $(|\sigma_n|)$, de terme général $1/n$, tend vers 0. En conclusion, les sommes de Riemann associées à la fonction f sur l'intervalle $]0, 1]$ n'admettent pas de limite dans $\mathbb{R}$ quand le pas de la subdivision tend vers 0. $\square$

■ Intégration des relations de comparaison

THÉORÈME 11.7. — Soit a et b des points de la droite numérique achevée $\overline{\mathbb{R}}$ tels que $-\infty < a < b \leq +\infty$, et f et g des fonctions définies et localement intégrables sur l'intervalle $I = [a, b[$. On suppose que g est à valeurs réelles positives et que :

$$f(x) = \underset{x \rightarrow b^-}{o}(g(x)) \quad (\text{resp. } f(x) \underset{x \rightarrow b^-}{\sim} g(x)).$$

a) Si la fonction g n'est pas intégrable sur $[a, b[$, alors :

$$\int_a^x f(t) dt = \underset{x \rightarrow b^-}{o} \left(\int_a^x g(t) dt \right) \quad \left(\text{resp. } \int_a^x f(t) dt \underset{x \rightarrow b^-}{\sim} \int_a^x g(t) dt \right).$$

b) Si la fonction g est intégrable sur $[a, b[$, alors :

$$\int_x^b f(t) dt = \underset{x \rightarrow b^-}{o} \left(\int_x^b g(t) dt \right) \quad \left(\text{resp. } \int_x^b f(t) dt \underset{x \rightarrow b^-}{\sim} \int_x^b g(t) dt \right).$$

On dispose du même théorème, dans le cas où $-\infty \leq a < b < +\infty$, en remplaçant $I = [a, b[$ par $I =]a, b]$ et $x \rightarrow b^-$ par $x \rightarrow a^+$.

11.27. Fonctions f et g telles que $f(x) = \underset{x \rightarrow +\infty}{o}(g(x))$ alors que :

$$\int_0^x g(t) dt = \underset{x \rightarrow +\infty}{o} \left(\int_0^x f(t) dt \right).$$

Nous notons E la fonction «partie entière» et nous introduisons les applications continues :

$$f : x \mapsto f(x) = |\sin x| \quad \text{et} \quad g : x \mapsto g(x) = \sqrt{x} \sin x$$

de $[0, +\infty[$ dans $\mathbb{R}$. Il est clair que f n'est pas intégrable sur $[0, +\infty[$ et que :

$$f(x) = \frac{|g(x)|}{\sqrt{x}} = \underset{x \rightarrow +\infty}{o}(g(x)).$$

Nous considérons les applications :

$$F : x \mapsto F(x) = \int_0^x f(t) dt \text{ et } G : x \mapsto G(x) = \int_0^x g(t) dt$$

de $[0, +\infty[$ dans $\mathbb{R}$. Comme f est positive, F est croissante sur $\mathbb{R}_+$.

Soit x un nombre réel > 0 . Nous posons $n = E(x/\pi)$. On a $n\pi \leq x < (n+1)\pi$ donc $F(n\pi) \leq F(x) \leq F((n+1)\pi)$. Comme la fonction f admet π pour période, on a $F(n\pi) = nF(\pi) = 2n$.

Par suite $2E\left(\frac{x}{\pi}\right) \leq F(x) \leq 2E\left(\frac{x+\pi}{\pi}\right)$. De plus $2E\left(\frac{x}{\pi}\right) \leq \frac{2x}{\pi} \leq 2E\left(\frac{x+\pi}{\pi}\right)$.

Il en résulte que :

$$F(x) \underset{x \rightarrow +\infty}{\sim} \frac{2x}{\pi}.$$

Une intégration par parties donne, pour tout nombre réel $x > 0$:

$$G(x) = - \underbrace{\sqrt{x} \cos x}_{= G_1(x)} + \underbrace{\int_0^x \frac{\cos t}{2\sqrt{t}} dt}_{= G_2(x)}.$$

On a, pour tout réel $x > 0$:

$$|G_1(x)| = |\sqrt{x} \cos x| \leq \sqrt{x} \text{ et } |G_2(x)| \leq \int_0^x \left| \frac{\cos t}{2\sqrt{t}} \right| dt \leq \int_0^x \frac{1}{2\sqrt{t}} dt = \sqrt{x}$$

donc $|G(x)| \leq 2\sqrt{x} = \frac{2}{\sqrt{x}} x$. Ainsi $G(x) = o_{x \rightarrow +\infty}(x)$ donc $G(x) = o_{x \rightarrow +\infty}(F(x))$. $\square$

11.28. Fonctions h et g équivalentes au voisinage de $+\infty$ alors que $x \mapsto \int_0^x h(t) dt$ n'est pas équivalente à $x \mapsto \int_0^x g(t) dt$ au voisinage de $+\infty$.

Nous reprenons les fonctions f et g de l'exemple précédent 11.27, dont nous utilisons les notations et les résultats, et nous posons $h = f + g$.

On a $f(x) = o_{x \rightarrow +\infty}(g(x))$, donc $h(x) \underset{x \rightarrow +\infty}{\sim} g(x)$.

Nous considérons l'application $H : x \mapsto H(x) = \int_0^x h(t) dt$ de $[0, +\infty[$ dans $\mathbb{R}$. Par linéarité de l'intégrale, $H = F + G$.

Or $G(x) = o_{x \rightarrow +\infty}(F(x))$ donc $H(x) \underset{x \rightarrow +\infty}{\sim} F(x)$; par suite $G(x) = o_{x \rightarrow +\infty}(H(x))$.

Ainsi les fonctions H et G ne sont pas équivalentes au voisinage de $+\infty$. $\square$

11.29. Fonctions f et g intégrables sur l'intervalle $[1, +\infty[$ telles que $f(x) = o_{x \rightarrow +\infty}(g(x))$ alors que :

$$\int_x^{+\infty} g(t) dt = o_{x \rightarrow +\infty} \left(\int_x^{+\infty} f(t) dt \right).$$

Nous introduisons les applications continues de $[1, +\infty[$ dans $\mathbb{R}$:

$$f : x \mapsto f(x) = \frac{\sin^2 x}{x^2} \text{ et } g : x \mapsto g(x) = \frac{\sin x}{x\sqrt{x}}.$$

On a, pour tout nombre réel $x \geq 1$, $|f(x)| \leq \frac{1}{x^2}$ et $|g(x)| \leq \frac{1}{x^{3/2}}$.

Les fonctions $x \mapsto \frac{1}{x^2}$ et $x \mapsto \frac{1}{x^{3/2}}$ étant intégrables sur $[1, +\infty[$, il en est de même de f et de g .

On a, pour tout x , $0 \leq f(x) = \frac{g(x) \sin x}{\sqrt{x}} \leq \frac{|g(x)|}{\sqrt{x}}$. Or $\frac{|g(x)|}{\sqrt{x}} = o_{x \rightarrow +\infty}(g(x))$, donc :

$$f(x) = o_{x \rightarrow +\infty}(g(x)).$$

Nous posons, pour tout réel $x \geq 1$:

$$F(x) = \int_x^{+\infty} f(t) dt = \int_x^{+\infty} \frac{\sin^2 t}{t^2} dt \text{ et } G(x) = \int_x^{+\infty} g(t) dt = \int_x^{+\infty} \frac{\sin t}{t\sqrt{t}} dt.$$

Soit $x \in [1, +\infty[$. Une intégration par parties donne $G(x) = \frac{\cos x}{x\sqrt{x}} - \frac{3}{2} \int_x^{+\infty} \frac{\cos t}{t^2\sqrt{t}} dt$.
On a :

$$\left| \frac{\cos x}{x\sqrt{x}} \right| \leq \frac{1}{x\sqrt{x}} \text{ et } \left| \int_x^{+\infty} \frac{\cos t}{t^2\sqrt{t}} dt \right| \leq \int_x^{+\infty} \left| \frac{\cos t}{t^2\sqrt{t}} \right| dt \leq \int_x^{+\infty} \frac{1}{t^2\sqrt{t}} dt = \frac{2}{3} \frac{1}{x\sqrt{x}}$$

donc :

$$|G(x)| \leq \frac{1}{x\sqrt{x}} + \frac{3}{2} \frac{2}{3} \frac{1}{x\sqrt{x}} = \frac{2}{x\sqrt{x}} = \frac{2}{\sqrt{x}} \frac{1}{x}.$$

Par conséquent $G(x) = o_{x \rightarrow +\infty}\left(\frac{1}{x}\right)$. Soit x un point de l'intervalle $[1, +\infty[$. On a :

$$F(x) = \int_x^{+\infty} \frac{1 - \cos(2t)}{2t^2} dt = \int_x^{+\infty} \frac{1}{2t^2} dt - \underbrace{\int_x^{+\infty} \frac{\cos(2t)}{2t^2} dt}_{= K(x)} = \frac{1}{2x} - K(x).$$

Par une intégration par parties, il vient : $K(x) = -\frac{\sin(2x)}{4x^2} + \int_x^{+\infty} \frac{\sin(2t)}{2t^3} dt$. Or :

$$\left| \frac{\sin(2x)}{4x^2} \right| \leq \frac{1}{4x^2} \text{ et } \left| \int_x^{+\infty} \frac{\sin(2t)}{2t^3} dt \right| \leq \int_x^{+\infty} \frac{|\sin(2t)|}{2t^3} dt \leq \int_x^{+\infty} \frac{1}{2t^3} dt = \frac{1}{4x^2},$$

donc $|K(x)| \leq \frac{1}{2x^2}$. Par conséquent $K(x) = o_{x \rightarrow +\infty}\left(\frac{1}{2x}\right)$, donc $F(x) \sim_{x \rightarrow +\infty} \frac{1}{2x}$.

Il en résulte que $G(x) = o_{x \rightarrow +\infty}(F(x))$. $\square$

11.30. Fonctions h et g intégrables sur $[1, +\infty[$ et équivalentes au voisinage de $+\infty$, alors que la fonction $x \mapsto \int_x^{+\infty} h(t) dt$ n'est pas équivalente à $x \mapsto \int_x^{+\infty} g(t) dt$ au voisinage de $+\infty$.

Nous reprenons les fonctions f et g de l'exemple précédent 11.29, dont nous utilisons les notations et les résultats, et nous posons $h = f + g$.

Comme $f(x) = o_{x \rightarrow +\infty}(g(x))$, on a $h(x) \sim_{x \rightarrow +\infty} g(x)$. Nous considérons l'application :

$$H : x \mapsto H(x) = \int_x^{+\infty} h(t) dt$$

de $[1, +\infty[$ dans $\mathbb{R}$. Par linéarité de l'intégrale, on a $H = F + G$.

De $G(x) = o_{x \rightarrow +\infty}(F(x))$, on déduit que $H(x) \sim_{x \rightarrow +\infty} F(x)$, ce qui montre que :

$$G(x) = o_{x \rightarrow +\infty}(H(x)).$$

Ainsi les fonctions H et G ne sont pas équivalentes au voisinage de $+\infty$. $\square$

Chapitre 12

Suites de fonctions

La définition de la convergence d'une suite de nombres réels découle de la notion de valeur absolue : c'est la seule qui soit féconde en mathématiques. La notion de limite d'une suite de fonctions est par contre plus délicate car elle peut varier selon les propriétés que l'on désire voir conserver à la limite. Il s'agit d'un problème de topologie sur l'ensemble des fonctions considérées, mais les modes de convergence les plus utilisés peuvent être définis plus directement. Nous nous intéresserons aux convergences simple, uniforme, en moyenne et en moyenne quadratique.

Dans tout le chapitre, les fonctions sont des fonctions de $\mathbb{R}$ dans $\mathbb{R}$.

■ Convergence simple

DÉFINITION 12.1. — Une suite (f_n) de fonctions définies sur une partie $\mathcal{A}$ de $\mathbb{R}$ converge simplement sur $\mathcal{A}$ vers une fonction f définie sur $\mathcal{A}$ si, pour tout point x de $\mathcal{A}$, la suite $(f_n(x))$ converge vers le nombre réel $f(x)$.

Cette notion ne fait pas entrer en ligne de compte l'interaction que peuvent avoir les points de l'ensemble de définition et l'on ne peut s'attendre pour cette notion de limite à la conservation de propriétés topologiques intéressantes.

12.1. Limite simple d'une suite de fonctions continues qui est discontinue.

Nous considérons la suite $(f_n)_{n \geq 0}$ d'applications de $[0, 1]$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = x^n. \end{array} \right.$$

Pour tout point x de $[0, 1[$, la suite (x^n) converge vers 0, et la suite (1^n) converge vers 1, donc la suite (f_n) converge simplement vers l'application :

$$\left| \begin{array}{l} f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x \in [0, 1[, \\ 1 & \text{si } x = 1, \end{cases} \end{array} \right.$$

et, puisque $\lim_{x \rightarrow 1^-} f(x) = 0 \neq 1 = f(1)$, la fonction f est discontinue en 1. $\square$

12.2. Limite simple d'une suite de fonctions continues qui est discontinue sur un ensemble dense.

Nous associons à tout entier $n \geq 1$ l'ensemble $A_n = \left\{ \frac{k}{2^n} \mid k \in [0, 2^n] \text{ et } k \text{ impair} \right\}$.

On a en particulier $A_1 = \left\{ \frac{1}{2} \right\}$, $A_2 = \left\{ \frac{1}{4}, \frac{3}{4} \right\}$ et $A_3 = \left\{ \frac{1}{8}, \frac{3}{8}, \frac{5}{8}, \frac{7}{8} \right\}$.

Remarquons que $(A_n)_{n \in \mathbb{N}^*}$ est une famille d'ensembles deux à deux disjoints. En effet, si $p, q \in \mathbb{N}^*$ et si $p < q$, alors, pour tout entier k impair appartenant à $[0, 2^p]$, $k/2^p = (2^{q-p}k)/2^q$ et l'entier $2^{q-p}k$ est pair, donc $k/2^p$ n'appartient pas à A_q .

Nous posons $B = \bigcup_{n \in \mathbb{N}^*} A_n$ et, pour tout entier $n \geq 1$, $B_n = \bigcup_{1 \leq p \leq n} A_p$.

Nous notons, pour tout élément b de B , $m(b)$ l'unique entier $n \geq 1$ tel que $b \in A_n$.

Soit $n \in \mathbb{N}^*$. Si $1 \leq p, q \leq n$, si $k \in [0, 2^p]$ et $\ell \in [0, 2^q]$, si k et ℓ sont impairs et si :

$$\frac{k}{2^p} < \frac{\ell}{2^q},$$

on a $2^q k < 2^p \ell$ et :

- si $p \neq q$, alors $p+q \leq 2n-1$ et, comme l'entier $2^p \ell - 2^q k$ est pair et strictement positif, $2^p \ell - 2^q k \geq 2$ donc :

$$\frac{\ell}{2^q} - \frac{k}{2^p} = \frac{2^p \ell - 2^q k}{2^{p+q}} \geq \frac{2}{2^{p+q}} \geq \frac{2}{2^{2n-1}} = \frac{4}{2^{2n}};$$

- si $p = q$, alors $k < \ell$ donc, $\ell - k$ étant pair et strictement positif, $\ell - k \geq 2$, ce qui donne :

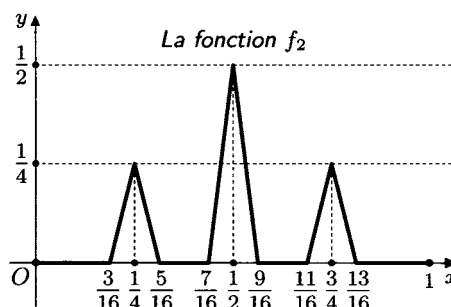
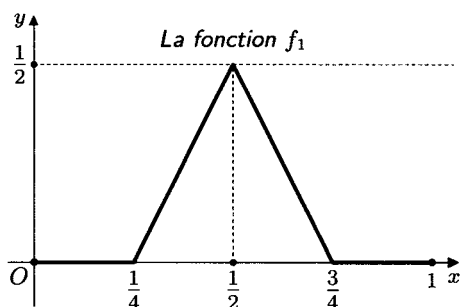
$$\frac{\ell}{2^q} - \frac{k}{2^p} = \frac{\ell - k}{2^p} \geq \frac{2}{2^p} \geq \frac{2}{2^n} = \frac{2 \times 2^n}{2^{2n}} \geq \frac{4}{2^{2n}},$$

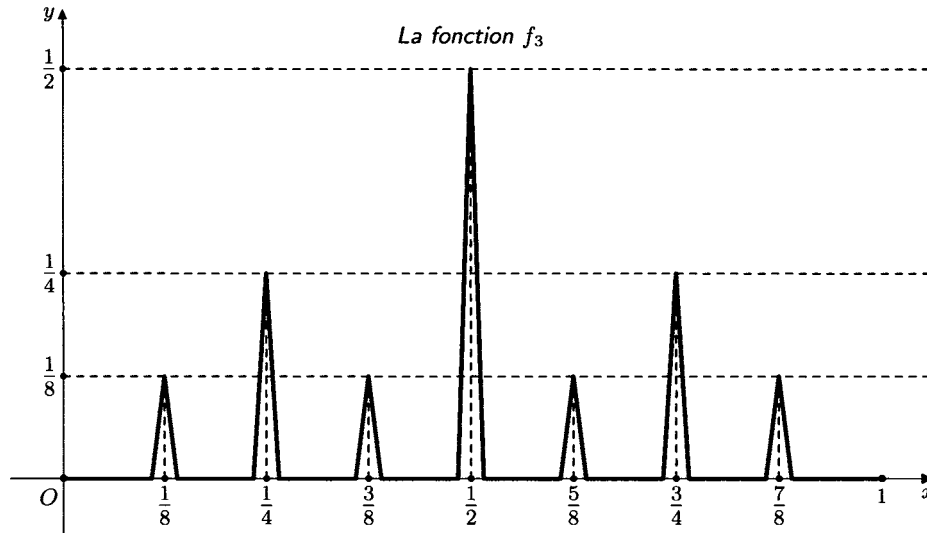
d'où l'on déduit, dans les deux cas, que $\frac{k}{2^p} + \frac{1}{2^{2n}} < \frac{\ell}{2^q} - \frac{1}{2^{2n}}$. Il en résulte que :

$$\left(\left[b - \frac{1}{2^{2n}}, b + \frac{1}{2^{2n}} \right] \right)_{b \in B_n}$$

est une famille de segments deux à deux disjoints, clairement inclus dans $[0, 1]$, ce qui justifie l'existence de l'application f_n de $[0, 1]$ dans $\mathbb{R}$ définie ainsi :

- pour tout point b de B_n , $f_n(b - (1/2^{2n})) = 0$, la restriction de f_n au segment $[b - (1/2^{2n}), b]$ est affine, $f_n(b) = 1/2^{2n}$, la restriction de f_n à $[b, b + (1/2^{2n})]$ est affine et $f_n(b + (1/2^{2n})) = 0$;
- $f_n(x) = 0$ pour tout point x de $[0, 1]$ n'appartenant pas à la réunion de la famille $([b - (1/2^{2n}), b + (1/2^{2n})])_{b \in B_n}$.





Si b appartient à B , alors $f_n(b) = 1/2^{m(b)}$ pour tout entier $n \geq m(b)$, donc la suite $(f_n(b))$ converge dans $\mathbb{R}$ vers $1/2^{m(b)}$.

Soit x un point de $[0, 1]$ n'appartenant pas à B . Soit ε un réel strictement positif. Nous choisissons un entier $n_0 \geq 1$ tel que $1/2^{n_0} < \varepsilon$. Le point x n'appartient pas à B_{n_0} et B_{n_0} est une partie finie non vide non vide de $\mathbb{R}$, donc la distance α de x à B_{n_0} est un réel strictement positif. Nous choisissons enfin un entier $N \geq n_0$ tel que $1/2^{2N} < \alpha$. Soit n un entier $\geq N$. S'il n'existe aucun point b de B_n tel que $x \in [b - 1/2^{2n}, b + 1/2^{2n}]$, alors $f_n(x) = 0$. Nous supposons donc qu'il existe $b \in B_n$ tel que $x \in [b - 1/2^{2n}, b + 1/2^{2n}]$. La distance de x à b étant inférieure ou égale à $1/2^{2n}$, elle est strictement inférieure à α , donc $b \notin B_{n_0}$. Par suite, pour tout $p \in \llbracket 1, n_0 \rrbracket$, $b \notin A_p$ donc $m(b) \neq p$, d'où l'on déduit que $m(b) > n_0$. Il en résulte que $0 \leq f_n(x) \leq f_n(b) = 1/2^{m(b)} < 1/2^{n_0} < \varepsilon$. Nous avons ainsi établi que, pour tout entier $n \geq N$, $0 \leq f_n(x) < \varepsilon$. La suite $(f_n(x))$ converge donc vers 0.

Finalement la suite de fonctions (f_n) converge simplement vers l'application :

$$\left| \begin{array}{l} f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} \frac{1}{2^{m(x)}} & \text{si } x \in B, \\ 0 & \text{si } x \notin B. \end{cases} \end{array} \right.$$

Il nous reste à démontrer que f est discontinue en tous les points de B et que B est dense dans $[0, 1]$ pour conclure.

Soit b un point de B . Le réel $\varepsilon_0 = 1/2^{m(b)}$ est strictement positif. Soit η un nombre réel > 0 . Alors $c = \text{Max}(0, b - \eta) < d = \text{Min}(1, b + \eta)$, donc l'intervalle $]c, d[$ contient des irrationnels. Nous choisissons un nombre irrationnel x dans $]c, d[$. Alors $x \in [0, 1]$ et $x \notin B$ — en effet, les éléments de B sont des rationnels — donc $f(x) = 0$, ce qui montre que $|f(x) - f(b)| = |f(b)| = 1/2^{m(b)} \geq \varepsilon_0$. Par conséquent f est discontinue en b .

Soit u et v des réels tels que $0 \leq u < v \leq 1$. Nous choisissons un entier naturel n tel que $1/2^n < v - u$ et nous notons p la partie entière de $2^n u$. Alors p est un entier relatif et $p/2^n \leq u < (p+1)/2^n$, ce qui, en ajoutant $1/2^n$, donne :

$$b = \frac{(p+1)}{2^n} \leq u + \frac{1}{2^n} < u + (v - u) = v.$$

De plus, $b > u \geq 0$ donc $p + 1 > 0$, ce qui, puisque $p + 1$ est entier, montre que $p + 1 \in \mathbb{N}^*$. Nous notons n_0 l'exposant de 2 dans la décomposition de $p + 1$ en produit de facteurs premiers. On a $p + 1 = 2^{n_0}k$ où k est un entier naturel impair. Alors $k \geq 1$ et $b = (2^{n_0}k)/2^n$, donc $(2^{n_0}k)/2^n < v < 1$, ce qui montre que $2^{n_0}/2^n < 1/k \leq 1$. Par suite $2^{n_0} < 2^n$ donc $n_0 < n$, d'où l'on déduit que :

$$n - n_0 \in \mathbb{N}^* \text{ et } b = \frac{k}{2^{n-n_0}}.$$

Il en résulte que $b \in B$. Comme $u < b < v$, le segment $[u, v]$ rencontre B .

En conclusion, B est dense dans $[0, 1]$. $\square$

12.3. Limite simple d'une suite double de fonctions continues qui n'est continue en aucun point.

Soit x un nombre réel irrationnel. Soit $n \in \mathbb{N}$. Comme $n!x$ n'appartient pas à $\mathbb{Z}$, $|\cos(n!\pi x)| < 1$, donc la suite $(u_{n,m})_{m \geq 0}$, de terme général $u_{n,m} = (\cos(n!\pi x))^{2m}$, converge vers 0.

Soit x un nombre rationnel. Nous notons (p, q) le représentant irréductible de x de dénominateur positif. Soit n un entier $\geq q$. Alors q divise $n!$, donc $n!x = (n!p)/q$ appartient à $\mathbb{Z}$, d'où l'on déduit que $|\cos(n!\pi x)| = 1$. Par conséquent la suite $(u_{n,m})_{m \geq 0}$, de terme général $u_{n,m} = (\cos(n!\pi x))^{2m}$, converge vers 1.

Il résulte de cette étude que l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \lim_{n \rightarrow \infty} \lim_{m \rightarrow \infty} (\cos(n!\pi x))^{2m} \end{array} \right.$$

est la fonction de Dirichlet, qui vaut 1 sur les rationnels et 0 sur les irrationnels¹. Quels que soient les entiers naturels n et m , la fonction $x \mapsto (\cos(n!\pi x))^{2m}$ est continue sur $\mathbb{R}$, alors que la fonction f est, comme nous l'avons vu en particulier dans l'exemple 8.1 (page 134), discontinue en tous les points de $\mathbb{R}$. $\square$

12.4. Limite simple d'une suite de fonctions bornées qui n'est pas bornée.

Nous associons à tout entier naturel n l'application :

$$\left| \begin{array}{l} f_n :]0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \text{Min}\left(\frac{1}{x}, n\right). \end{array} \right.$$

On a, pour tout entier naturel n , $0 \leq f_n(x) \leq n$ pour tout point x de $]0, +\infty[$, donc la fonction f_n est bornée sur $]0, +\infty[$.

Soit x un point de $]0, +\infty[$. La propriété d'Archimède nous fournit un entier naturel $N > 1/x$. Alors $N \in \mathbb{N}^*$ et, pour tout entier $n \geq N$, $n \geq 1/x$ donc $f_n(x) = 1/x$. Il découle de cette étude que la suite de fonctions (f_n) converge simplement vers l'application :

$$\left| \begin{array}{l} f :]0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \frac{1}{x}, \end{array} \right.$$

qui n'est pas bornée sur $]0, +\infty[$. $\square$

1. On doit au mathématicien allemand Alfred Pringsheim (1850-1941), à la fin du XIX^e siècle, l'expression sous cette forme de la fonction de Dirichlet.

La convergence simple est insuffisante pour obtenir l'interversion des limites.

12.5. Suite de fonctions (f_n) telle que :

$$\lim_{n \rightarrow \infty} \lim_{x \rightarrow 1} f_n(x) \neq \lim_{x \rightarrow 1} \lim_{n \rightarrow \infty} f_n(x).$$

Nous reprenons la suite $(f_n)_{n \geq 0}$ d'applications de $[0, 1]$ dans $\mathbb{R}$ de l'exemple 12.1 (page 231), de terme général :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = x^n. \end{array} \right.$$

On a, pour tout entier naturel n , $\lim_{x \rightarrow 1} f_n(x) = 1$, donc $\lim_{n \rightarrow \infty} \lim_{x \rightarrow 1} f_n(x) = 1$.

Pour tout $x \in [0, 1[$, $0 \leq x < 1$ et $f_n(x) = x^n$, donc $\lim_{n \rightarrow \infty} f_n(x) = 0$. Par suite :

$$\lim_{x \rightarrow 1} \lim_{n \rightarrow \infty} f_n(x) = 0 \neq 1 = \lim_{n \rightarrow \infty} \lim_{x \rightarrow 1} f_n(x). \quad \square$$

Voici un exemple analogue pour une limite en $+\infty$.

12.6. Suite de fonctions (f_n) telle que :

$$\lim_{n \rightarrow \infty} \lim_{x \rightarrow +\infty} f_n(x) \neq \lim_{x \rightarrow +\infty} \lim_{n \rightarrow \infty} f_n(x).$$

Nous considérons la suite $(f_n)_{n \geq 1}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \arctan \frac{x}{n}. \end{array} \right.$$

On a, pour tout entier $n \geq 1$, $\lim_{x \rightarrow +\infty} f_n(x) = \frac{\pi}{2}$, donc $\lim_{n \rightarrow \infty} \lim_{x \rightarrow +\infty} f_n(x) = \frac{\pi}{2}$.

Par ailleurs, $\lim_{n \rightarrow \infty} f_n(x) = 0$ pour tout réel x , donc $\lim_{x \rightarrow +\infty} \lim_{n \rightarrow \infty} f_n(x) = 0 \neq \frac{\pi}{2}$. $\square$

■ Convergence uniforme

La convergence simple est une propriété trop faible pour la conservation à la limite de certaines propriétés comme par exemple la continuité. Aussi on introduit une notion plus forte de convergence, appelée la convergence uniforme².

Soit $\mathcal{A}$ une partie de $\mathbb{R}$, (f_n) une suite de fonctions définies sur $\mathcal{A}$ et f une fonction définie sur $\mathcal{A}$.

DÉFINITION 12.2. — La suite (f_n) converge uniformément sur $\mathcal{A}$ vers la fonction f si la fonction $f_n - f$ est, à partir d'un certain rang n_0 , bornée sur $\mathcal{A}$ et si la suite $(\sigma_n)_{n \geq n_0}$, de terme général $\sigma_n = \sup_{x \in \mathcal{A}} |f_n(x) - f(x)|$, converge vers 0.

La suite (f_n) converge uniformément vers f sur $\mathcal{A}$ si, et seulement si, pour tout réel $\varepsilon > 0$, il existe un entier naturel N tel que, pour tout entier $n \geq N$ et tout point x de $\mathcal{A}$, $|f_n(x) - f(x)| \leq \varepsilon$.

2. La convergence uniforme est introduite en 1841 par Karl Weierstrass mais il ne la publie pas. George Stokes en 1847 et Philipp Seidel en 1848 la redéfinissent indépendamment l'un de l'autre.

THÉORÈME 12.1. — La limite uniforme d'une suite de fonctions continues est une fonction continue.

La convergence uniforme de la suite (f_n) vers f sur $\mathcal{A}$ entraîne clairement la convergence simple de (f_n) vers f sur $\mathcal{A}$. La réciproque est fausse comme nous le montre l'exemple 12.1 (page 231) puisque, dans cet exemple, la fonction limite f n'est pas continue. Donnons un autre exemple du même type, pour une suite de fonctions définies sur un intervalle non compact.

12.7. Suite de fonctions qui converge simplement mais pas uniformément.

Nous associons à tout entier naturel n l'application :

$$\left| \begin{array}{l} f_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \frac{1}{1 + (x - n)^2} \end{array} \right.$$

Pour tout réel x , la suite de terme général $1 + (x - n)^2$ tend vers $+\infty$, donc la suite $(f_n(x))$ converge vers 0. La suite de fonctions (f_n) converge donc simplement vers l'application nulle $\mathbf{0} : x \mapsto \mathbf{0}(x) = 0$ de $\mathbb{R}$ dans $\mathbb{R}$. Or, pour tout entier naturel n , $f_n(n) = 1$ donc la borne supérieure σ_n de $\{|f_n(x) - \mathbf{0}(x)| \mid x \in \mathbb{R}\}$ est supérieure ou égale à 1, ce qui montre que la suite (σ_n) ne converge pas vers 0. En conclusion, la convergence de la suite de fonctions (f_n) vers $\mathbf{0}$ n'est pas uniforme. $\square$

Après l'introduction de la notion de convergence uniforme, différents mathématiciens se sont interrogés sur la réciproque du théorème 12.1 : si la limite simple d'une suite de fonctions continues est continue, la convergence de la suite est-elle uniforme ? Certains pensaient même l'avoir démontrée. Georg Cantor mit fin à ce débat en 1880 grâce au contre-exemple qui suit.

12.8. Suite de fonctions continues qui converge simplement vers une fonction continue, la convergence de cette suite n'étant pas uniforme.

Nous introduisons la suite $(f_n)_{n \geq 0}$ d'applications de $[0, +\infty[$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : [0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \frac{2nx}{1 + n^2 x^2} \end{array} \right.$$

Si x est un réel strictement positif, alors :

$$f_n(x) \underset{n \rightarrow \infty}{\sim} \frac{2}{nx}$$

donc la suite $(f_n(x))$ converge vers 0. De plus $f_n(0) = 0$ pour tout entier naturel n , donc la suite de fonctions (f_n) converge simplement vers l'application nulle $\mathbf{0} : x \mapsto \mathbf{0}(x) = 0$ de $[0, +\infty[$ dans $\mathbb{R}$, qui est évidemment continue sur $[0, +\infty[$. Pour tout entier $n \geq 1$, $f_n(1/n) = 1$, donc la borne supérieure σ_n de l'ensemble $\{|f_n(x) - \mathbf{0}(x)| \mid x \in [0, +\infty[\}$ est supérieure ou égale à 1, ce qui montre que la suite (σ_n) ne converge pas vers 0. En conclusion, la convergence de la suite de fonctions (f_n) vers $\mathbf{0}$ n'est pas uniforme. $\square$

La convergence simple n'entraîne pas, en général, la convergence uniforme. Le théorème suivant donne des conditions suffisantes pour l'obtenir.

THÉORÈME 12.2. — Théorème de Dini³.

Soit $\mathcal{A}$ une partie de $\mathbb{R}$, (f_n) une suite de fonctions définies sur $\mathcal{A}$ et f une fonction définie sur $\mathcal{A}$. Si (f_n) converge simplement vers f sur $\mathcal{A}$ et si les trois assertions suivantes sont vraies :

- (1) pour tout n , f_n est continue sur $\mathcal{A}$,
 - (2) $\mathcal{A}$ est un compact de $\mathbb{R}$,
 - (3) pour tout point x de $\mathcal{A}$, la suite $(|f_n(x) - f(x)|)$ est décroissante,
- alors la suite (f_n) converge uniformément vers f sur $\mathcal{A}$.

Nous démontrons que chacune de ces trois conditions est nécessaire à l'application du théorème de Dini.

12.9. Cas où l'assertion (2) du théorème de Dini n'est pas vérifiée.

Nous considérons la suite $(f_n)_{n \geq 1}$ d'applications de $]0, +\infty[$ dans $\mathbb{R}$, de terme général :

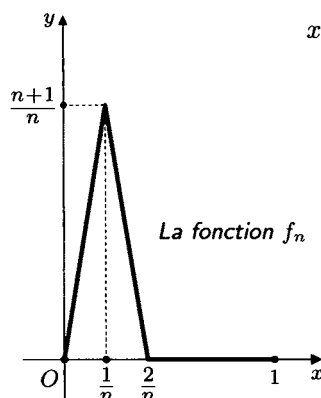
$$\left| \begin{array}{l} f_n :]0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \frac{1}{nx}. \end{array} \right.$$

Pour tout nombre réel $x > 0$, la suite $(f_n(x))$ est décroissante et converge vers 0. Par conséquent, la suite (f_n) converge simplement sur $]0, +\infty[$ vers l'application nulle $0 : x \mapsto 0(x) = 0$ de $]0, +\infty[$ dans $\mathbb{R}$. Pour tout entier $n \geq 1$, $f_n(1/n) = 1$, donc la borne supérieure σ_n de $\{|f_n(x) - 0(x)| \mid x \in]0, +\infty[\}$ est supérieure ou égale à 1, ce qui montre que la suite (σ_n) ne converge pas vers 0. En conclusion, les assertions (1) et (3) du théorème de Dini sont vraies, mais la convergence de la suite de fonctions (f_n) vers l'application nulle n'est pas uniforme. $\square$

12.10. Cas où l'assertion (3) du théorème de Dini n'est pas vraie.

Nous introduisons la suite $(f_n)_{n \geq 2}$ d'applications de $[0, 1]$ dans $\mathbb{R}$, de terme général :

$$f_n : [0, 1] \longrightarrow \mathbb{R} \quad x \longmapsto f_n(x) = \begin{cases} (n+1)x & \text{si } 0 \leq x \leq \frac{1}{n}, \\ -(n+1)x + 2\frac{n+1}{n} & \text{si } \frac{1}{n} < x < \frac{2}{n}, \\ 0 & \text{si } \frac{2}{n} \leq x \leq 1. \end{cases}$$



Soit n un entier supérieur ou égal à 2.

Si l'on substitue $\frac{1}{n}$ (resp. $\frac{2}{n}$) à x dans l'expression :

$$-(n+1)x + 2\frac{n+1}{n},$$

on trouve $\frac{n+1}{n} = f_n\left(\frac{1}{n}\right)$ (resp. $0 = f_n\left(\frac{2}{n}\right)$).

Il en résulte que la fonction f_n est affine par morceaux continue sur le segment $[0, 1]$.

3. Du nom du mathématicien italien Ulisse Dini (1845-1918).

Soit x un réel tel que $0 < x \leq 1$. La suite $(2/n)$ converge vers 0, donc il existe un entier $N \geq 2$ tel que $2/N < x$. Nous choisissons un tel N . Alors, pour tout entier $n \geq N$, on a $2/n \leq 2/N < x \leq 1$ donc $f_n(x) = 0$. Par conséquent la suite $(f_n(x))$ converge vers 0. Comme $f_n(0) = 0$ pour tout entier $n \geq 2$, la suite $(f_n(0))$ converge vers 0. La suite de fonctions (f_n) converge donc simplement sur le compact $[0, 1]$ vers l'application nulle $f : x \mapsto f(x) = 0$ de $[0, 1]$ dans $\mathbb{R}$. D'autre part on a, pour tout entier $n \geq 2$:

$$\sigma_n = \sup_{x \in [0,1]} |f_n(x) - f(x)| = \sup_{x \in [0,1]} f_n(x) = f_n\left(\frac{1}{n}\right) = \frac{n+1}{n} \geq 1$$

donc la suite (σ_n) ne converge pas vers 0. Par conséquent la convergence de la suite de fonctions (f_n) vers f n'est pas uniforme, bien que les conditions (1) et (2) du théorème Dini soient vérifiées. On peut remarquer qu'une condition légèrement différente de l'assertion (3) est cependant vraie : la suite (σ_n) , de terme général :

$$\sigma_n = \sup_{x \in [0,1]} |f_n(x) - f(x)|,$$

est strictement décroissante — mais de limite $\ell > 0$. $\square$

12.11. Cas où l'assertion (1) du théorème de Dini n'est pas vraie.

L'ensemble $\Omega = \mathbb{Q} \cap [0, 1]$ étant dénombrable, $\Omega = \{r_n \mid n \in \mathbb{N}\} = \{r_0, r_1, r_2, \dots\}$ où $(r_n)_{n \geq 0}$ est une suite d'éléments deux à deux distincts de Ω .

Nous associons à tout entier naturel n l'application :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \begin{cases} 0 & \text{si } x \in [0, 1] \setminus \Omega, \\ 0 & \text{si } x = r_k \text{ où } k \in \llbracket 0, n \rrbracket, \\ 1 & \text{si } x = r_k \text{ pour un entier } k > n. \end{cases} \end{array} \right.$$

Si $x \in [0, 1] \setminus \Omega$, la suite $(f_n(x))$ est constante égale à 0, donc elle converge vers 0. Si $x \in \Omega$, il existe un entier naturel k et un seul tel que $x = r_k$, et alors $f_n(x) = 0$ pour tout entier $n \geq k$, donc la suite $(f_n(x))$ converge également vers 0. Par conséquent, la suite (f_n) converge simplement sur $[0, 1]$ vers l'application nulle $f : x \mapsto f(x) = 0$ de $[0, 1]$ dans $\mathbb{R}$. Pour tout entier naturel n , $f_n(r_{n+1}) = 1$, donc la borne supérieure σ_n de $\{|f_n(x) - f(x)| \mid x \in [0, 1]\}$ est égale à 1, ce qui montre que la suite (σ_n) ne converge pas vers 0. La convergence de la suite (f_n) vers f sur le compact $[0, 1]$ n'est donc pas uniforme, bien que les conditions (2) et (3) du théorème de Dini soient vérifiées. $\square$

Le produit de deux suites convergentes de nombre réels est une suite convergente. On en déduit que ce résultat s'applique au produit de suites de fonctions pour la convergence simple ; il devient faux pour la convergence uniforme.

12.12. Deux suites de fonctions qui convergent uniformément mais dont le produit ne converge pas uniformément.

Nous considérons les suites $(f_n)_{n \geq 0}$ et $(g_n)_{n \geq 0}$ d'applications de $[1, 2]$ dans $\mathbb{R}$ de termes généraux :

$$\left| \begin{array}{l} f_n : [1, 2] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = x \left(1 + \frac{1}{n+1}\right) \end{array} \right.$$

et :

$$g_n : [1, 2] \longrightarrow \mathbb{R} \\ x \longmapsto g_n(x) = \begin{cases} \frac{1}{n+1} & \text{si } x \text{ est irrationnel,} \\ q + \frac{1}{n+1} & \text{si } x \in \mathbb{Q} \cap [1, 2] \text{ et si } (p, q) \text{ est} \\ & \text{le représentant irréductible de} \\ & \text{x de dénominateur positif.} \end{cases}$$

On a, pour tout point x de $[1, 2]$, $f_n(x) - x = \frac{x}{n+1}$, donc $|f_n(x) - x| \leq \frac{2}{n+1}$.

La majoration de $|f_n(x) - x|$ obtenue ci-dessus ne dépend pas de la variable x et la suite de terme général $2/(n+1)$ converge vers 0, donc la suite (f_n) converge uniformément sur $[1, 2]$ vers l'application $f : x \mapsto f(x) = x$ de $[1, 2]$ dans $\mathbb{R}$.

Soit x un point de $[1, 2]$. Si x est irrationnel, $|g_n(x) - 0| = 1/(n+1)$; si x est rationnel et si (p, q) est le représentant irréductible de x de dénominateur positif, alors $|g_n(x) - q| = 1/(n+1)$. On en déduit, en raisonnant comme pour la suite (f_n) , que la suite (g_n) converge uniformément sur $[1, 2]$ vers l'application :

$$g : [1, 2] \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} 0 & \text{si } x \text{ est irrationnel,} \\ q & \text{si } x \in \mathbb{Q} \cap [1, 2] \text{ et si } (p, q) \text{ est} \\ & \text{le représentant irréductible de} \\ & \text{x de dénominateur positif.} \end{cases}$$

Pour tout $x \in [1, 2]$, la suite $(f_n(x))$ converge vers $f(x)$ et la suite $(g_n(x))$ vers $g(x)$, donc la suite $(f_n(x)g_n(x))$ converge vers $f(x)g(x)$: la suite $(f_n g_n)$ converge simplement sur $[1, 2]$ vers la fonction produit fg .

Pour tout $x \in \mathbb{Q} \cap [1, 2]$, on a, en notant (p, q) le représentant irréductible de x de dénominateur positif :

$$f_n(x)g_n(x) - f(x)g(x) = \frac{xq}{n+1} + \frac{x}{(n+1)^2} + \frac{x}{n+1} > 0.$$

Si $n \in \mathbb{N}^*$, $a_n = \frac{n+1}{n}$ appartient à $\mathbb{Q} \cap [1, 2]$ et, n et $n+1$ étant premiers entre eux, le représentant irréductible de a_n de dénominateur positif est $(n+1, n)$, donc :

$$|f_n(a_n)g_n(a_n) - f(a_n)g(a_n)| = 1 + \frac{1}{n(n+1)} + \frac{1}{n} \geq 1$$

ce qui montre que $\sigma_n = \sup \{ |f_n(x)g_n(x) - f(x)g(x)| \mid x \in [1, 2] \} \geq 1$. La suite $(f_n g_n)$ ne converge donc pas uniformément sur $[1, 2]$ vers la fonction fg . $\square$

Si I et J sont des intervalles, (f_n) une suite de fonctions définies sur I et à valeurs dans J , f une fonction définie sur I à valeurs dans J et g une fonction définie sur J , si la suite (f_n) converge uniformément vers f sur I et si g est uniformément continue sur J , alors la suite $(g \circ f_n)$ converge uniformément vers $g \circ f$ sur I ; si l'on suppose seulement que, pour tout point x de I , g est continue en $f(x)$, la suite $(g \circ f_n)$ converge simplement vers $g \circ f$ sur I , mais rien ne permet d'affirmer que cette convergence est uniforme.

12.13. Suite de fonctions (f_n) , à valeurs dans un intervalle J et qui converge uniformément, et fonction g continue sur J telle que la suite $(g \circ f_n)$ ne converge pas uniformément.

Nous considérons l'application :

$$g : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = x^2$$

et la suite $(f_n)_{n \geq 1}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R} \longrightarrow \mathbb{R} \\ t \longmapsto f_n(t) = t + \frac{1}{n}. \end{array} \right.$$

On a, pour tout entier $n \geq 1$, $|f_n(t) - t| = 1/n$ pour tout réel t , $1/n$ ne dépend pas de la variable t et la suite $(1/n)$ converge vers 0, donc la suite de fonctions (f_n) converge uniformément sur $\mathbb{R}$ vers l'application $f : t \mapsto f(t) = t$ de $\mathbb{R}$ dans $\mathbb{R}$. On a, pour tout nombre réel t :

$$(1) \quad g \circ f_n(t) - g \circ f(t) = (f_n(t))^2 - (f(t))^2 = t^2 + \frac{2t}{n} + \frac{1}{n^2} - t^2 = \frac{2t}{n} + \frac{1}{n^2}.$$

Par conséquent, pour tout réel t , la suite $(g \circ f_n(t))$ converge vers $g \circ f(t)$, donc la suite de fonctions $(g \circ f_n)$ converge simplement sur $\mathbb{R}$ vers la fonction $g \circ f$. On déduit de (1) que, pour tout $n \in \mathbb{N}^*$, $g \circ f_n(n) - g \circ f(n) = 2 + (1/n^2) \geq 2$ donc $\sigma_n = \sup \{|f_n(x)g_n(x) - f(x)g(x)| \mid x \in [1, 2]\} \geq 2$. Il en résulte que la convergence sur $\mathbb{R}$ de la suite de fonctions $(g \circ f_n)$ vers $g \circ f$ n'est pas uniforme. $\square$

■ Dérivation

Soit (f_n) une suite de fonctions qui converge uniformément vers une fonction f . La dérivabilité de la fonction f_n pour tout n n'entraîne pas la dérivabilité de f , ni même, si f est dérivable, la convergence de la suite de fonctions (f_n') vers f' . Pour obtenir un résultat intéressant, il faut des hypothèses plus fortes sur la convergence de la suite de fonctions (f_n') .

THÉORÈME 12.3. — Soit Ω un ouvert de $\mathbb{R}$ et (f_n) une suite de fonctions définies et dérivables sur Ω , qui converge simplement sur Ω vers une fonction f . Si g est une fonction définie sur Ω et si la suite (f_n') des fonctions dérivées converge uniformément⁴ vers g , la fonction f est dérivable sur Ω et $f' = g$.

On peut, dans le texte du théorème 12.3, remplacer l'ouvert Ω par un intervalle I .

12.14. Suite de fonctions (f_n) de classe $\mathcal{C}^1$ sur $\mathbb{R}$ ayant pour limite uniforme une fonction dérivable, alors que la suite (f_n') des fonctions dérivées ne converge pas.

Nous considérons la suite $(f_n)_{n \geq 1}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \frac{\sin(nx)}{\sqrt{n}}. \end{array} \right.$$

Pour tout entier $n \geq 1$, f_n est de classe $\mathcal{C}^1$ sur $\mathbb{R}$ et, pour tout nombre réel x , $|f_n(x)| \leq (1/\sqrt{n})$. Cette majoration ne dépend pas de la variable x et la suite de terme général $1/\sqrt{n}$ converge vers 0, donc la suite (f_n) converge uniformément sur $\mathbb{R}$ vers l'application nulle $0 : x \mapsto 0(x) = 0$ de $\mathbb{R}$ dans $\mathbb{R}$, fonction constante donc dérivable sur $\mathbb{R}$.

4. Il suffit en fait que, pour tout point a de Ω , (f_n') converge uniformément vers g sur au moins un voisinage de a . En particulier, la convergence uniforme sur tout compact de la suite des fonctions dérivées est une condition suffisante.

Pour tout entier $n \geq 1$, on a $f_n'(x) = \sqrt{n} \cos(nx)$ pour tout nombre réel x , donc $f_n'(\pi) = (-1)^n \sqrt{n}$, terme général d'une suite divergente. Par conséquent la suite (f_n') d'applications de $\mathbb{R}$ dans $\mathbb{R}$ ne converge pas simplement sur $\mathbb{R}$. $\square$

12.15. Suite de fonctions (f_n) de classe $\mathcal{C}^1$ sur $\mathbb{R}$ ayant pour limite uniforme une fonction dérivable f et telle que la suite (f_n') converge vers une fonction $g \neq f'$.

Nous introduisons la suite $(f_n)_{n \geq 0}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \frac{x}{1+n^2x^2}. \end{array} \right.$$

C'est clairement une suite de fonctions de classe $\mathcal{C}^1$ sur $\mathbb{R}$.

Soit n un entier ≥ 1 . On a, si x un nombre réel différent de 0 :

$$|f_n(x)| = \frac{|x|}{1+n^2x^2} \leq \begin{cases} |x| & \text{car } 1+n^2x^2 > 1 > 0, \\ \frac{|x|}{n^2x^2} = \frac{1}{n^2|x|} & \text{car } 1+n^2x^2 > n^2x^2 > 0, \end{cases}$$

donc :

- si $|x| \leq \frac{1}{n}$, $|f_n(x)| \leq |x| \leq \frac{1}{n}$;
- si $|x| > \frac{1}{n}$, $n^2|x| > n > 0$, donc $|f_n(x)| \leq \frac{1}{n^2|x|} < \frac{1}{n}$.

De plus $f_n(0) = 0 < \frac{1}{n}$, donc $|f_n(x)| \leq \frac{1}{n}$ pour tout nombre réel x .

La majoration obtenue ne dépend pas de la variable x et la suite $(1/n)$ converge vers 0, donc la suite (f_n) converge uniformément sur $\mathbb{R}$ vers l'application nulle $f : x \mapsto f(x) = 0$ de $\mathbb{R}$ dans $\mathbb{R}$, fonction constante donc dérivable sur $\mathbb{R}$, et $f'(x) = 0$ pour tout nombre réel x .

On a, pour tout entier naturel n et tout nombre réel x , $f_n'(x) = \frac{1-n^2x^2}{(1+n^2x^2)^2}$.

Pour tout $n \in \mathbb{N}$, $f_n'(0) = 1$, donc la suite $(f_n'(0))$ converge vers 1. Si x est un réel et si $x \neq 0$, alors $x^2 > 0$, donc :

$$f_n'(x) \underset{n \rightarrow \infty}{\sim} \frac{-n^2x^2}{n^4x^4} = -\frac{1}{n^2x^2}$$

ce qui montre que la suite $(f_n'(x))$ converge vers 0. Il en résulte que la suite de fonctions (f_n') converge simplement sur $\mathbb{R}$ vers l'application :

$$\left| \begin{array}{l} g : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} 1 & \text{si } x = 0, \\ 0 & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

La fonction limite g de la suite (f_n') est donc différente de la fonction dérivée f' de la limite uniforme f de la suite (f_n) . $\square$

12.16. Suite de fonctions (f_n) de classe $\mathcal{C}^1$ sur $\mathbb{R}$ ayant pour limite uniforme une fonction f qui n'est pas dérivable sur $\mathbb{R}$.

Nous considérons la suite $(f_n)_{n \geq 1}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \sqrt{x^2 + \frac{1}{n^2}}. \end{array} \right.$$

Pour tout entier $n \geq 1$ on a, pour tout nombre réel x :

$$0 \leq |x|^2 = x^2 \leq x^2 + \frac{1}{n^2} \leq |x|^2 + \frac{2|x|}{n} + \frac{1}{n^2} = \left(|x| + \frac{1}{n}\right)^2$$

donc $|x| \leq f_n(x) \leq |x| + \frac{1}{n}$, ce qui donne l'inégalité $|f_n(x) - |x|| \leq \frac{1}{n}$.

La majoration obtenue ne dépend pas de la variable x et la suite de terme général $1/n$ converge vers 0, donc la suite (f_n) converge uniformément sur $\mathbb{R}$ vers l'application $f : x \mapsto f(x) = |x|$ de $\mathbb{R}$ dans $\mathbb{R}$.

La fonction limite f de la suite (f_n) n'est donc pas dérivable en 0.

Pour tout entier $n \geq 1$, $x^2 + (1/n^2) > 0$ pour tout $x \in \mathbb{R}$, donc f_n est de classe $\mathcal{C}^1$ sur $\mathbb{R}$ et, pour tout nombre réel x :

$$f_n'(x) = \frac{x}{\sqrt{x^2 + \frac{1}{n^2}}}.$$

Pour tout $n \in \mathbb{N}^*$, $f_n'(0) = 0$, donc la suite $(f_n'(0))$ converge vers 0. Si x est un réel différent de 0, alors, comme $\lim_{n \rightarrow \infty} (1/n^2) = 0$, on obtient :

$$\lim_{n \rightarrow \infty} f_n'(x) = \frac{x}{\sqrt{x^2}} = \frac{x}{|x|} = \begin{cases} 1 & \text{si } x > 0, \\ -1 & \text{si } x < 0. \end{cases}$$

Nous remarquons que la suite (f_n') converge simplement vers l'application :

$$\left| \begin{array}{l} g : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} -1 & \text{si } x < 0, \\ 0 & \text{si } x = 0, \\ 1 & \text{si } x > 0, \end{cases} \end{array} \right.$$

et que, pour tout réel $x \neq 0$, $g(x) = f'(x)$. $\square$

12.17. Suite de fonctions (f_n) de classe $\mathcal{C}^1$ sur $\mathbb{R}$ ayant pour limite uniforme une fonction f qui n'est dérivable en aucun point.

Nous utilisons dans cet exemple les propriétés de la convergence uniforme des séries de fonctions, rappelées au chapitre 13.

Nous considérons l'application :

$$\left| \begin{array}{l} g : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \frac{1}{4} \left(1 + \sin \left(\frac{\pi}{2} (4x - 1) \right) \right) \end{array} \right.$$

La fonction g est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$, 1 est une période de g , g s'annule en tous les entiers relatifs, et, pour tout nombre réel x , $0 \leq g(x) \leq 1/2$ et :

$$g'(x) = \frac{\pi}{2} \cos \left(\frac{\pi}{2} (4x - 1) \right),$$

donc $|g'(x)| \leq \frac{\pi}{2}$. La fonction g est donc lipschitzienne de rapport $\frac{\pi}{2}$ sur $\mathbb{R}$.

Nous associons à g la suite $(u_n)_{n \geq 0}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} u_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto u_n(x) = \frac{g(4^n x)}{4^n} \end{array} \right.$$

Pour tout entier naturel n et tout nombre réel x , $|u_n(x)| \leq 1/(2 \times 4^n)$, donc la série de fonctions $\sum_n u_n$ converge normalement, donc uniformément, sur $\mathbb{R}$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \sum_{n=0}^{+\infty} u_n(x) \end{array} \right.$$

et, u_n étant continue sur $\mathbb{R}$ pour tout $n \in \mathbb{N}$, montre que f est continue sur $\mathbb{R}$. On prouve alors, comme dans l'exemple 9.6 (pages 161 et 162), que la fonction f n'est dérivable en aucun point de $\mathbb{R}$.

De plus la suite $(f_n)_{n \geq 0}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \sum_{p=0}^n u_p(x) \end{array} \right.$$

converge uniformément vers f sur $\mathbb{R}$ et, pour tout entier naturel n , la fonction f_n est de classe $\mathcal{C}^1$ — en fait de classe $\mathcal{C}^\infty$ — sur $\mathbb{R}$. $\square$

■ Intégration

THÉORÈME 12.4. — Si a et b sont des réels tels que $a < b$, (f_n) une suite de fonctions définies et intégrables au sens de Riemann sur le segment $[a, b]$ et f une fonction définie sur $[a, b]$, et si (f_n) converge uniformément vers f sur $[a, b]$, alors f est intégrable au sens de Riemann sur $[a, b]$, la suite de terme général $\int_a^b f_n(t) dt$ converge dans $\mathbb{R}$ et :

$$\lim_{n \rightarrow \infty} \int_a^b f_n(t) dt = \int_a^b f(t) dt.$$

Ce résultat est encore vrai pour l'intégrale de Lebesgue, mais on dispose d'un résultat plus général.

THÉORÈME 12.5. — **Théorème de convergence dominée de Lebesgue.**

Si I est un intervalle, (f_n) une suite de fonctions intégrables au sens de Lebesgue sur I et f une fonction définie sur I , si la suite (f_n) converge presque partout (en particulier simplement) vers f et s'il existe une fonction g définie, positive et intégrable au sens de Lebesgue sur I telle que, pour tout n et pour presque tout point x de I , $|f_n(x)| \leq g(x)$, alors f est intégrable au sens de Lebesgue sur l'intervalle I , la suite de terme général $\int_I f_n(t) dt$ converge dans $\mathbb{R}$ et :

$$\lim_{n \rightarrow \infty} \int_I f_n(t) dt = \int_I f(t) dt.$$

12.18. Suite de fonctions intégrables au sens de Riemann sur un segment, qui converge simplement sur ce segment vers une fonction qui n'est pas intégrable.

L'ensemble $\Omega = \mathbb{Q} \cap [0, 1]$ étant dénombrable, $\Omega = \{r_n \mid n \in \mathbb{N}\} = \{r_0, r_1, r_2, \dots\}$ où $(r_n)_{n \geq 0}$ est une suite d'éléments deux à deux distincts de Ω . Nous associons à tout entier naturel n l'ensemble $\Omega_n = \{r_p \mid p \in \llbracket 0, n \rrbracket\} = \{r_0, r_1, \dots, r_n\}$ et

l'application :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \begin{cases} 1 & \text{si } x \in \mathcal{Q}_n, \\ 0 & \text{si } x \in [0, 1] \setminus \mathcal{Q}_n. \end{cases} \end{array} \right.$$

Soit n un entier naturel. La fonction f_n est nulle sur $[0, 1]$ sauf en un nombre fini de points, donc elle est intégrable au sens de Riemann sur $[0, 1]$. Clairement, la suite (f_n) converge simplement sur $[0, 1]$ vers l'application :

$$\left| \begin{array}{l} f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 1 & \text{si } x \in \mathcal{Q}, \\ 0 & \text{si } x \in [0, 1] \setminus \mathcal{Q}, \end{cases} \end{array} \right.$$

restriction à $[0, 1]$ de la fonction de Dirichlet, dont nous savons (exemple 11.1, page 208) qu'elle n'est pas intégrable au sens de Riemann sur le segment $[0, 1]$. $\square$

Le théorème de convergence dominée de Lebesgue permet d'affirmer que la fonction f de l'exemple précédent 12.18 est intégrable au sens de Lebesgue et que son intégrale sur $[0, 1]$ est nulle. On retrouve ce résultat en remarquant que $\mathcal{Q}$, ensemble dénombrable, est de mesure nulle.

12.19. Suite (f_n) de fonctions intégrables au sens de Riemann sur $[0, 1]$, de limite simple sur $[0, 1]$ une fonction f intégrable sur $[0, 1]$, mais telle que $\int_0^1 f(t) dt \neq \lim_{n \rightarrow \infty} \int_0^1 f_n(t) dt$.

Nous introduisons la suite $(f_n)_{n \geq 2}$ d'applications de $[0, 1]$ dans $\mathbb{R}$, de terme général :

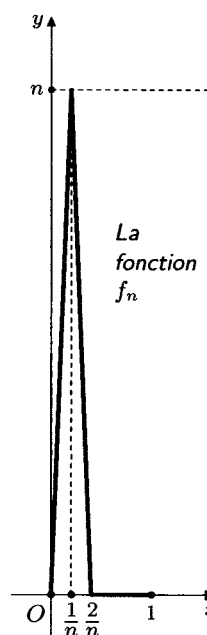
$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \begin{cases} n^2 x & \text{si } 0 \leq x \leq \frac{1}{n}, \\ -n^2 x + 2n & \text{si } \frac{1}{n} < x < \frac{2}{n}, \\ 0 & \text{si } \frac{2}{n} \leq x \leq 1. \end{cases} \end{array} \right.$$

Pour tout entier $n \geq 2$, $f_n(0) = 0$, donc la suite $(f_n(0))$ converge vers 0. Si x appartient à $]0, 1]$, alors, en choisissant un entier naturel $N > 2/x$, on obtient, pour tout entier $n \geq N$, $2/n < x \leq 1$ donc $f_n(x) = 0$, ce qui montre que la suite $(f_n(x))$ converge vers 0.

Il en résulte que la suite (f_n) converge simplement sur $[0, 1]$ vers l'application nulle $f : x \mapsto f(x) = 0$ de $[0, 1]$ dans $\mathbb{R}$.

On a, pour tout entier $n \geq 2$, $\int_0^1 f_n(t) dt = 1$ puisque c'est l'aire d'un triangle de base $2/n$ et de hauteur n , donc :

$$\lim_{n \rightarrow \infty} \int_0^1 f_n(t) dt = 1 \neq 0 = \int_0^1 f(t) dt. \quad \square$$

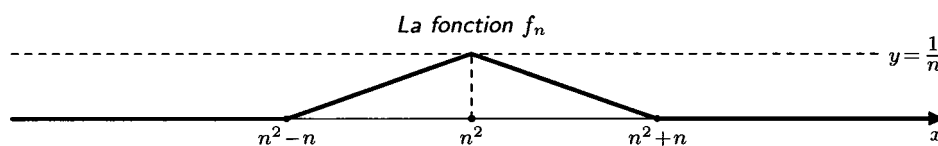


Si I est un intervalle d'intérieur non vide, si I n'est pas compact et si (f_n) est une suite de fonctions définies et intégrables sur I , la convergence uniforme de la suite (f_n) vers une fonction f définie et intégrable sur I n'est pas suffisante pour obtenir l'égalité $\int_I f(t) dt = \lim_{n \rightarrow \infty} \int_I f_n(t) dt$.

12.20. Suite (f_n) de fonction intégrables sur l'intervalle $[0, +\infty[$, convergeant uniformément vers une fonction f intégrable sur I , mais telle que $\int_0^{+\infty} f(t) dt \neq \lim_{n \rightarrow \infty} \int_0^{+\infty} f_n(t) dt$.

Nous considérons la suite $(f_n)_{n \geq 1}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \begin{cases} \frac{x}{n^2} - 1 + \frac{1}{n} & \text{si } n^2 - n < x \leq n^2, \\ -\frac{x}{n^2} + 1 + \frac{1}{n} & \text{si } n^2 < x < n^2 + n, \\ 0 & \text{si } x < n^2 - n \text{ ou } x > n^2 + n. \end{cases} \end{array} \right.$$



Pour tout entier $n \geq 1$, la substitution de $n^2 - n$ (resp. de $n^2 + n$) à x dans l'expression $(x/n^2) - 1 + (1/n)$ (resp. $-(x/n^2) + 1 + (1/n)$) donne $0 = f_n(n^2 - n)$ (resp. $0 = f_n(n^2 + n)$), donc f_n est continue sur $\mathbb{R}$. Pour tout entier $n \geq 1$ et tout nombre réel $x \geq 0$, $|f_n(x)| = f_n(x) \leq 1/n$. Cette majoration est indépendante de la variable x et la suite $(1/n)$ converge vers 0, donc la suite (f_n) converge uniformément sur $[0, +\infty[$ vers l'application nulle $f : x \mapsto f(x) = 0$ de $[0, +\infty[$ dans $\mathbb{R}$, intégrable et d'intégrale nulle sur $[0, +\infty[$.

Soit n un entier ≥ 1 . On a, pour tout nombre réel $x \geq n^2 + n$:

$$\int_0^x f_n(t) dt = \int_{n^2-n}^{n^2+n} f_n(t) dt = 1$$

(c'est l'aire d'un triangle de base $2n$ et de hauteur $1/n$), donc $\int_0^x f_n(t) dt$ admet pour limite 1 quand x tend vers $+\infty$, ce qui montre que la fonction f_n est intégrable sur $[0, +\infty[$ et que son intégrale sur $[0, +\infty[$ est égale à 1. Par conséquent :

$$\lim_{n \rightarrow \infty} \int_0^1 f_n(t) dt = 1 \neq 0 = \int_0^1 f(t) dt. \quad \square$$

■ Convergence en moyenne

Soit a et b des nombres réels tels que $a < b$. Nous avons vu que si une suite (f_n) d'applications continues de $[a, b]$ dans $\mathbb{R}$ converge uniformément sur $[a, b]$ vers une application f de $[a, b]$ dans $\mathbb{R}$, alors f est continue sur $[a, b]$ (théorème 12.1, page 236). Nous munissons l'ensemble $\mathcal{C}([a, b])$ des applications continues de $[a, b]$ dans $\mathbb{R}$ de sa structure d'espace vectoriel sur $\mathbb{R}$ et nous considérons sur $\mathcal{C}([a, b])$ la norme uniforme, c'est-à-dire l'application :

$$\left| \begin{array}{l} \|\cdot\|_\infty : \mathcal{C}([a, b]) \longrightarrow \mathbb{R}_+ \\ f \longmapsto \|f\|_\infty = \sup_{t \in [a, b]} |f(t)|. \end{array} \right.$$

Alors $\|\cdot\|_\infty$ est une norme sur l'espace vectoriel réel $\mathcal{C}([a, b])$ et, si (f_n) est une

suite d'applications continues de $[a, b]$ dans $\mathbb{R}$ et f une application continue de $[a, b]$ dans $\mathbb{R}$, (f_n) converge uniformément vers f sur $[a, b]$ si, et seulement si, la suite (f_n) converge vers f dans l'espace vectoriel normé $(\mathcal{C}([a, b]), \|\cdot\|_\infty)$ — en effet, pour tout n , la borne supérieure σ_n de l'ensemble $\{|f_n(t) - f(t)| \mid t \in [a, b]\}$ est égale à $\|f_n - f\|_\infty$.

L'intégrale nous permet d'introduire d'autres types de convergence sur l'ensemble des applications continues de $[a, b]$ dans $\mathbb{R}$.

THÉORÈME 12.6. — Les applications :

$$\text{et : } \left\{ \begin{array}{l} \|\cdot\|_1 : \mathcal{C}([a, b]) \longrightarrow \mathbb{R}_+ \\ f \longmapsto \|f\|_1 = \int_a^b |f(t)| dt \\ \|\cdot\|_2 : \mathcal{C}([a, b]) \longrightarrow \mathbb{R}_+ \\ f \longmapsto \|f\|_2 = \sqrt{\int_a^b f(t)^2 dt} \end{array} \right.$$

sont des normes sur l'espace vectoriel réel $\mathcal{C}([a, b])$ des applications continues de $[a, b]$ dans $\mathbb{R}$ et on a, sur $\mathcal{C}([a, b])$:

$$\|\cdot\|_1 \leq \sqrt{b-a} \|\cdot\|_2, \quad \|\cdot\|_2 \leq \sqrt{b-a} \|\cdot\|_\infty \quad \text{et} \quad \|\cdot\|_1 \leq (b-a) \|\cdot\|_\infty.$$

La convergence pour la norme $\|\cdot\|_1$ s'appelle la convergence en moyenne et celle pour $\|\cdot\|_2$ est appelée la convergence en moyenne quadratique. Nous utilisons les abréviations suivantes : CVS pour «converge simplement», CVU pour «converge uniformément», CVM pour «converge en moyenne» et CVQ pour «converge en moyenne quadratique». Comme la convergence uniforme vers une fonction entraîne la convergence simple vers cette fonction, on déduit des inégalités du théorème 12.6 le schéma d'implications suivant, qui s'applique à une suite (f_n) d'applications continues de $[a, b]$ dans $\mathbb{R}$ et à une application continue f de $[a, b]$ dans $\mathbb{R}$:

$$\begin{array}{c} (f_n) \text{ CVU vers } f \implies (f_n) \text{ CVQ vers } f \implies (f_n) \text{ CVM vers } f \\ \downarrow \\ (f_n) \text{ CVS vers } f \end{array}$$

Nous démontrons que l'on ne peut obtenir d'autres implications, hormis celles qui découlent de la transitivité. Pour le rapport entre les convergences simple et uniforme, nous avons déjà vu plusieurs exemples.

12.21. Suite de fonctions convergeant simplement sur le segment $[0, 1]$ vers une fonction f , qui ne converge ni en moyenne, ni en moyenne quadratique vers f sur $[0, 1]$.

Nous reprenons la suite de fonctions $(f_n)_{n \geq 2}$, définies sur $[0, 1]$, de l'exemple 12.19 (page 244). Nous avons vu que la suite (f_n) converge simplement sur $[0, 1]$ vers l'application nulle $f : t \mapsto f(t) = 0$ de $[0, 1]$ dans $\mathbb{R}$ et que, pour tout entier $n \geq 2$:

$$\int_0^1 |f_n(t)| dt = \int_0^1 f_n(t) dt = 1.$$

D'autre part, pour tout entier $n \geq 2$:

$$\int_0^1 (f_n(t))^2 dt = 2 \int_0^{1/n} (f_n(t))^2 dt = 2 \int_0^{1/n} n^4 t^2 dt = 2 \left[n^4 \frac{t^3}{3} \right]_{t=0}^{t=1/n} = \frac{2n}{3}.$$

Tout ceci montre que la suite (f_n) ne converge sur le segment $[0, 1]$ ni en moyenne ni en moyenne quadratique vers sa limite simple f . $\square$

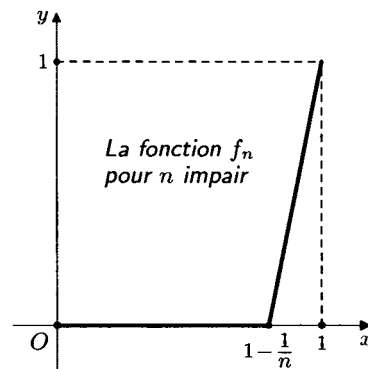
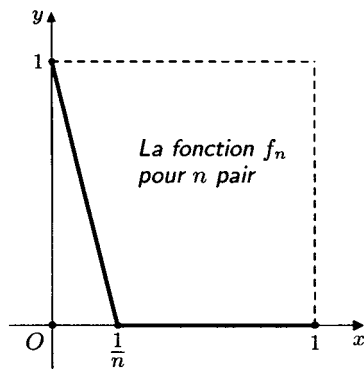
12.22. Suite de fonctions convergeant en moyenne et en moyenne quadratique, mais qui ne converge pas simplement.

On associe à tout entier naturel n pair et différent de 0 l'application :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \begin{cases} -nx + 1 & \text{si } 0 \leq x \leq \frac{1}{n}, \\ 0 & \text{si } \frac{1}{n} < x \leq 1, \end{cases} \end{array} \right.$$

et à tout entier naturel n impair l'application :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \begin{cases} 0 & \text{si } 0 \leq x < 1 - \frac{1}{n}, \\ nx - n + 1 & \text{si } 1 - \frac{1}{n} \leq x \leq 1. \end{cases} \end{array} \right.$$



Pour tout entier $n \geq 1$, on a $0 \leq f_n(t) \leq 1$ donc $0 \leq (f_n(t))^2 \leq f_n(t) = |f_n(t)|$ pour tout point t de $[0, 1]$ et l'intégrale de f_n sur le segment $[0, 1]$ est l'aire du triangle mis en évidence sur chacun des deux dessins ci-dessus, donc :

$$0 \leq \int_0^1 (f_n(t))^2 dt \leq \int_0^1 |f_n(t)| dt = \frac{1}{2n}.$$

La suite $(1/(2n))$ converge vers 0, donc (f_n) converge en moyenne et en moyenne quadratique sur $[0, 1]$ vers l'application nulle $\mathbf{0} : x \mapsto \mathbf{0}(x) = 0$ de $[0, 1]$ dans $\mathbb{R}$. D'autre part $f_n(0)$ est égal à 1 ou à 0 selon que n est pair ou impair, donc la suite (f_n) ne converge simplement sur $[0, 1]$ vers aucune application de $[0, 1]$ dans $\mathbb{R}$. $\square$

12.23. Suite de fonctions continues convergeant en moyenne et en moyenne quadratique et qui converge simplement vers une fonction discontinue.

Nous considérons la suite de fonctions $(f_n)_{n \geq 0}$, déjà utilisée dans l'exemple 12.1 (page 231), dont le terme général est l'application $f_n : x \mapsto f_n(x) = x^n$ de $[0, 1]$ dans $\mathbb{R}$. C'est une suite de fonctions continues sur $[0, 1]$, qui converge simplement sur $[0, 1]$ vers l'application f de $[0, 1]$ dans $\mathbb{R}$ définie par $f(1) = 1$ et $f(x) = 0$ pour tout $x \in [0, 1[$. La fonction limite f est donc discontinue en 1.

On a, pour tout entier naturel n :

$$\|f_n\|_1 = \int_0^1 |f_n(t)| dt = \int_0^1 t^n dt = \frac{1}{n+1}$$

et :

$$\|f_n\|_2 = \sqrt{\int_0^1 (f_n(t))^2 dt} = \sqrt{\int_0^1 t^{2n} dt} = \frac{1}{\sqrt{2n+1}}.$$

Comme les suites $(1/(n+1))$ et $(1/\sqrt{2n+1})$ convergent toutes deux vers 0, la suite de fonctions (f_n) converge en moyenne et en moyenne quadratique vers l'application nulle $\mathbf{0} : 0 \mapsto \mathbf{0}(x) = 0$ de $[0, 1]$ dans $\mathbb{R}$. $\square$

Remarquons que si l'on se place dans l'ensemble $\mathcal{C}([0, 1])$ des applications continues de $[0, 1]$ dans $\mathbb{R}$, la suite de fonctions (f_n) de l'exemple précédent 12.23 admet des limites différentes pour la convergence simple et pour la convergence en moyenne ; dans cet exemple, l'application nulle $\mathbf{0}$ et f sont toutes les deux des limites de la suite (f_n) et $f \neq \mathbf{0}$. Sur un ensemble plus large, la convergence en moyenne n'est plus définie par une norme, mais par une semi-norme⁵ et une suite de fonctions peut alors avoir plusieurs limites en moyenne.

12.24. Suite de fonctions convergeant en moyenne et en moyenne quadratique et qui ne converge pas uniformément.

Nous reprenons la suite $(f_n)_{n \geq 0}$ de l'exemple 12.23, qui converge en moyenne et en moyenne quadratique sur $[0, 1]$ vers l'application nulle de $[0, 1]$ dans $\mathbb{R}$. Pour tout $n \in \mathbb{N}$, la fonction f_n est continue sur $[0, 1]$. Si la suite (f_n) convergeait uniformément sur le segment $[0, 1]$, ce serait vers la fonction f rappelée dans l'exemple 12.23 et, en application du théorème 12.1 (page 236), f serait continue sur $[0, 1]$. Or la fonction f est discontinue en 1, donc la suite (f_n) ne converge pas uniformément sur le segment $[0, 1]$. $\square$

12.25. Suite de fonctions qui converge en moyenne mais qui ne converge pas en moyenne quadratique.

Nous considérons la suite $(f_n)_{n \geq 2}$ d'applications de $[0, 1]$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \frac{1}{\sqrt{\ln n} \sqrt{x + \frac{1}{n}}}. \end{array} \right.$$

Pour tout entier $n \geq 2$, f_n est continue et positive sur $[0, 1]$ et :

$$\int_0^1 |f_n(t)| dt = \frac{1}{\sqrt{\ln n}} \left[2\sqrt{t + \frac{1}{n}} \right]_{t=0}^{t=1} = \frac{2}{\sqrt{\ln n}} \left(\sqrt{1 + \frac{1}{n}} - \sqrt{\frac{1}{n}} \right) = \frac{\frac{2}{\sqrt{\ln n}}}{\sqrt{1 + \frac{1}{n}} + \sqrt{\frac{1}{n}}}$$

donc la suite $(\|f_n\|_1)$ converge vers 0, ce qui montre que (f_n) converge en moyenne vers l'application nulle $\mathbf{0} : x \mapsto \mathbf{0}(x) = 0$ de $[0, 1]$ dans $\mathbb{R}$. Supposons que la suite (f_n) converge en moyenne quadratique vers une fonction f sur $[0, 1]$. Alors (f_n)

5. Voir dans le chapitre 17 le dernier paragraphe de la page 319.

converge en moyenne vers f sur $[0, 1]$ et $f = \mathbf{0}$ — schéma d'implications de la page 246 et unicité de la limite dans un espace vectoriel normé —, ce qui prouve que la suite $(\|f_n\|_2)$ converge vers 0. Or, pour tout entier $n \geq 2$:

$$\int_0^1 (f_n(t))^2 dt = \frac{1}{\ln n} \int_0^1 \frac{1}{t + \frac{1}{n}} dt = \frac{1}{\ln n} \left[\ln \left(t + \frac{1}{n} \right) \right]_{t=0}^{t=1} = \frac{\ln \left(1 + \frac{1}{n} \right)}{\ln n} + 1,$$

donc la suite $(\|f_n\|_2)$ converge vers 1. En conclusion, la suite (f_n) ne converge pas en moyenne quadratique sur $[0, 1]$. $\square$

Lorsque l'intervalle I de départ n'est pas un intervalle compact, on peut encore définir la norme $\|\cdot\|_1$ sur l'espace vectoriel des applications continues intégrables sur I et la norme $\|\cdot\|_2$ sur l'espace vectoriel des applications continues de carré intégrable sur I , et l'exemple 12.20 (page 245) montre que la convergence uniforme n'entraîne pas la convergence en moyenne ; de même la convergence en moyenne quadratique n'entraîne pas la convergence en moyenne comme le montre l'exemple suivant.

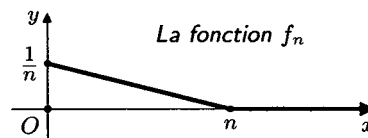
12.26. Suite de fonctions convergeant en moyenne quadratique mais pas en moyenne sur un intervalle non compact.

Nous considérons la suite $(f_n)_{n \geq 1}$ d'applications de $[0, +\infty[$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : [0, +\infty[\rightarrow \mathbb{R} \\ x \mapsto f_n(x) = \begin{cases} \frac{1}{n} \left(1 - \frac{x}{n} \right) & \text{si } x \leq n, \\ 0 & \text{si } x > n. \end{cases} \end{array} \right.$$

Soit $n \in \mathbb{N}^*$. Comme $f_n(n) = 0$, la fonction f_n est continue et positive sur $[0, +\infty[$. On a, pour tout nombre réel $x \geq n$:

$$\int_0^x |f_n(t)| dt = \int_0^n f_n(t) dt = \frac{1}{2}$$



(c'est l'aire du triangle mis en évidence sur le dessin ci-contre), donc $\int_0^x |f_n(t)| dt$ admet pour limite $1/2$ quand x tend vers $+\infty$, ce qui montre que la fonction f_n est intégrable sur $[0, +\infty[$ et que :

$$\int_0^{+\infty} |f_n(t)| dt = \frac{1}{2}.$$

De plus, pour tout nombre réel $x \geq n$:

$$\int_0^x (f_n(t))^2 dt = \frac{1}{n^2} \int_0^n \left(1 - \frac{t}{n} \right)^2 dt = \frac{1}{n^2} \left[-\frac{n}{3} \left(1 - \frac{t}{n} \right)^3 \right]_{t=0}^{t=n} = \frac{1}{3n},$$

donc $\int_0^x (f_n(t))^2 dt$ admet pour limite $1/(3n)$ quand x tend vers $+\infty$, d'où l'on déduit que la fonction $(f_n)^2$ est intégrable sur $[0, +\infty[$ et que :

$$\int_0^{+\infty} (f_n(t))^2 dt = \frac{1}{3n}.$$

Comme la suite $(1/(3n))$ converge vers 0, la suite (f_n) converge en moyenne quadratique sur $\mathbb{R}_+$ vers l'application nulle $\mathbf{0} : x \mapsto \mathbf{0}(x) = 0$ de $[0, +\infty[$ dans $\mathbb{R}$. Or, pour tout entier $n \geq 1$, $\int_0^{+\infty} |f_n(t)| dt = 1/2$, donc la suite (f_n) ne converge pas en moyenne vers $\mathbf{0}$ sur $[0, +\infty[$. $\square$

Chapitre 13

Séries de fonctions

Les séries de fonctions sont apparues à la fin du XVII^e siècle, lorsque Isaac Newton puis Brook Taylor décomposent des fonctions en séries, en fait en séries entières, pour calculer des intégrales. Il faut attendre Cauchy en 1821 pour avoir des critères précis de convergence, et ce n'est que dans les années 1840 que l'on introduit différents types de convergence. La richesse de cette théorie résulte de ce que l'on peut ainsi étudier des fonctions qui ne s'exprime pas à l'aide des fonctions connues. C'est le cas en particulier des solutions de certaines équations différentielles. L'intérêt est bien sûr de trouver des conditions pour que les propriétés des sommes partielles se retrouvent à la limite.

Dans tout le chapitre, les fonctions sont des fonctions de $\mathbb{R}$ dans $\mathbb{R}$.

■ Différents types de convergence

Pour les définitions et les théorèmes, $\mathcal{A}$ est une partie de $\mathbb{R}$ et (f_n) une suite de fonctions définies sur $\mathcal{A}$.

DÉFINITION 13.1. — La série de fonctions $\sum f_n$ converge simplement (resp. uniformément) sur $\mathcal{A}$ si la suite de fonctions (S_n) , de terme général :

$$S_n = \sum_{p=0}^n f_p,$$

converge simplement (resp. uniformément) sur $\mathcal{A}$ et, dans ce cas, la somme S de la série de fonctions $\sum f_n$ est la limite simple (resp. uniforme) de la suite (S_n) .

DÉFINITION 13.2. — Si la série de fonctions $\sum f_n$ converge simplement sur $\mathcal{A}$, (S_n) est appelée la suite des sommes partielles de la série et la suite des restes de cette série est la suite $(R_n)_{n \geq 0}$ de fonctions dont le terme général est la fonction, définie sur $\mathcal{A}$:

$$R_n : x \mapsto R_n(x) = \sum_{p=n+1}^{+\infty} f_p(x).$$

Par conséquent, si la série de fonctions $\sum f_n$ converge simplement sur $\mathcal{A}$, si (R_n) est sa suite des restes et S sa somme, on a, pour tout n et pour tout $x \in \mathcal{A}$:

$$R_n(x) = S(x) - S_n(x).$$

On en déduit le théorème suivant.

THÉORÈME 13.1. — Si la série de fonctions $\sum f_n$ converge simplement sur $\mathcal{A}$ et si sa somme est la fonction S , la série $\sum f_n$ converge uniformément sur $\mathcal{A}$ si, et seulement si, la suite de fonctions (R_n) converge uniformément sur $\mathcal{A}$ vers l'application nulle $\mathbf{0} : x \mapsto \mathbf{0}(x) = 0$ de $\mathcal{A}$ dans $\mathbb{R}$.

Comme pour les suites de fonctions, la convergence uniforme entraîne la convergence simple. La réciproque est fautive comme le montre l'exemple suivant.

13.1. Série de fonctions qui converge simplement mais qui ne converge pas uniformément.

Nous considérons la suite $(f_n)_{n \geq 0}$ d'applications de $]0, +\infty[$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n :]0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = e^{-nx}. \end{array} \right.$$

Pour tout nombre réel $x > 0$, on a $0 < e^{-x} < 1$ et la série $\sum_{n \geq 0} f_n(x)$ est une série géométrique de raison e^{-x} . On en déduit que la série de fonction $\sum f_n$ converge simplement sur $]0, +\infty[$ et que sa somme est l'application :

$$\left| \begin{array}{l} S :]0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto S(x) = \frac{1}{1 - e^{-x}} = \frac{e^x}{e^x - 1}. \end{array} \right.$$

Nous utilisons la suite (R_n) des restes de la série $\sum f_n$. On a, pour tout réel $x > 0$ et tout entier naturel n :

$$R_n(x) = \sum_{p=n+1}^{+\infty} f_p(x) = \sum_{p=n+1}^{+\infty} (e^{-x})^p = \frac{e^{-(n+1)x}}{1 - e^{-x}}.$$

Il en résulte que, pour tout entier naturel n :

$$R_n\left(\frac{1}{n+1}\right) = \frac{e^{-1}}{1 - e^{-1/(n+1)}},$$

terme général d'une suite qui diverge vers $+\infty$. Si la suite de fonctions (R_n) convergerait uniformément sur $]0, +\infty[$ vers l'application nulle, il existerait un entier naturel N tel que, pour tout entier $n \geq N$, $|R_n(x)| \leq 1$ pour tout réel $x > 0$, donc $|R_n(\frac{1}{n+1})| \leq 1$, en contradiction avec $\lim_{(n \rightarrow \infty)} R_n(\frac{1}{n+1}) = +\infty$. Par conséquent la suite (R_n) ne converge pas uniformément sur $]0, +\infty[$ vers l'application nulle, donc la série de fonctions $\sum f_n$ ne converge pas uniformément sur $]0, +\infty[$. $\square$

La convergence uniforme sur $\mathcal{A}$ est celle correspondant à la norme uniforme :

$$\left| \begin{array}{l} \|\cdot\|_\infty : \mathcal{B}(\mathcal{A}, \mathbb{R}) \longrightarrow \mathbb{R}_+ \\ f \longmapsto \|f\|_\infty = \sup_{x \in \mathcal{A}} |f(x)| \end{array} \right.$$

sur l'espace vectoriel réel $\mathcal{B}(\mathcal{A}, \mathbb{R})$ des applications bornées de $\mathcal{A}$ dans $\mathbb{R}$. Si la série $\sum f_n$ converge uniformément sur $\mathcal{A}$, alors, à partir d'un certain rang, f_n est bornée sur $\mathcal{A}$, et la suite $(\|f_n\|_\infty)$ converge vers 0.

Dans l'exemple précédent 13.1, $\mathcal{A} =]0, +\infty[$ et, pour tout entier naturel n , on a $0 < e^{-nx} \leq 1$ pour tout $x \in]0, +\infty[$ et e^{-nx} admet pour limite 1 quand x tend vers 0 à droite, donc la fonction f_n est bornée sur $]0, +\infty[$ et $\|f_n\|_\infty = 1$: la suite $(\|f_n\|_\infty)$ ne converge pas vers 0.

13.2. Série de fonctions $\sum g_n$ qui converge simplement mais qui ne converge pas uniformément, alors que la suite $(\|g_n\|_\infty)$ converge vers 0.

Nous reprenons la suite de fonctions (f_n) de l'exemple précédent 13.1, à laquelle nous associons la suite $(g_n)_{n \geq 1}$ d'applications de $]0, +\infty[$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} g_n :]0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto g_n(x) = \frac{1}{n} f_n(x) = \frac{1}{n} e^{-nx}. \end{array} \right.$$

Pour tout nombre réel $x > 0$, on a, pour tout entier $n \geq 1$, $0 \leq g_n(x) \leq f_n(x)$ et la série $\sum f_n(x)$ converge, donc la série $\sum g_n(x)$ converge. La série de fonctions $\sum g_n$ converge donc simplement sur $]0, +\infty[$ et sa somme est l'application :

$$T : x \mapsto T(x) = \sum_{n=1}^{+\infty} g_n(x)$$

de $]0, +\infty[$ dans $\mathbb{R}$, ce qui signifie que la suite de fonctions $(T_n)_{n \geq 1}$, dont le terme général est l'application :

$$T_n : x \mapsto T_n(x) = \sum_{p=1}^n g_p(x)$$

de l'intervalle $]0, +\infty[$ dans $\mathbb{R}$, converge simplement vers T sur $]0, +\infty[$.

Soit n un entier ≥ 1 . Soit x un point de $]0, +\infty[$. On a $g_p(x) > 0$ pour tout entier $p \geq 1$, donc $T(x) \geq T_{2n}(x)$, ce qui montre que :

$$T(x) - T_n(x) \geq T_{2n}(x) - T_n(x) = \sum_{p=n+1}^{2n} g_p(x).$$

Pour tout $p \in [n+1, 2n]$, $-px \geq -2nx$ donc $e^{-px} \geq e^{-2nx} > 0$, et $\frac{1}{p} \geq \frac{1}{2n} > 0$ donc $g_p(x) \geq g_{2n}(x)$. Il en résulte que :

$$T(x) - T_n(x) \geq n g_{2n}(x) = \frac{1}{2} e^{-2nx}.$$

Comme $\frac{1}{2} e^{-2nx}$ admet pour limite 1/2 quand x tend vers 0 à droite, il existe un réel $\delta_n > 0$ tel que, pour tout $x \in]0, \delta_n[$, $\frac{1}{2} e^{-2nx} > 1/4$ donc $T(x) - T_n(x) > 1/4$.

Si la série de fonctions $\sum g_n$ convergeait uniformément sur $]0, +\infty[$, il existerait un entier $N \geq 1$ tel que, pour tout entier $n \geq N$, $|T(x) - T_n(x)| \leq 1/4$ pour tout $x \in]0, +\infty[$, donc pour tout $x \in]0, \delta_n[$, en contradiction avec ce qui précède. Par conséquent la série de fonctions $\sum g_n$ ne converge pas uniformément sur $]0, +\infty[$.

Soit n un entier ≥ 1 . Nous avons vu que la fonction f_n est bornée sur $]0, +\infty[$ et que $\|f_n\|_\infty = 1$, d'où l'on déduit que la fonction g_n est bornée sur $]0, +\infty[$ et que :

$$\|g_n\|_\infty = \frac{1}{n}.$$

Par conséquent la suite $(\|g_n\|_\infty)$ converge vers 0. $\square$

Montrer la convergence uniforme d'une série de fonctions est parfois délicat car, en général, nous n'en connaissons pas la somme. De plus, l'intérêt des séries est de tirer des conclusions sur la série et sa somme par la simple étude du terme général. Aussi, il est utile de définir un nouveau type de convergence plus fort que la convergence uniforme.

DÉFINITION 13.3. — La série $\sum f_n$ de fonctions définies et bornées sur $\mathcal{A}$ converge normalement sur $\mathcal{A}$ si la série $\sum \|f_n\|_\infty$ converge¹

THÉORÈME 13.2. — Si une série $\sum f_n$ de fonctions définies et bornées sur $\mathcal{A}$ converge normalement sur $\mathcal{A}$, elle converge uniformément sur $\mathcal{A}$.

En effet, si la série $\sum \|f_n\|_\infty$ converge, alors, pour tout point x de l'ensemble $\mathcal{A}$, on a $|f_n(x)| \leq \|f_n\|_\infty$ pour tout n , donc la série $\sum f_n(x)$ converge, ce qui montre que la série $\sum f_n$ converge simplement sur $\mathcal{A}$ et, en notant (R_n) sa suite des restes, on a, pour tout point x de $\mathcal{A}$ et pour tout n :

$$|R_n(x)| \leq \sum_{k=n+1}^{+\infty} |f_k(x)| \leq \sum_{k=n+1}^{+\infty} \|f_k\|_\infty = \rho_n,$$

cette majoration ne dépend pas de la variable x et la suite (ρ_n) converge vers 0, donc la suite de fonctions (R_n) converge uniformément sur $\mathcal{A}$ vers l'application nulle de $\mathcal{A}$ dans $\mathbb{R}$. La réciproque est fausse.

13.3. Série de fonctions qui converge uniformément mais qui ne converge pas normalement.

Nous considérons la suite $(f_n)_{n \geq 1}$ d'applications de $[0, 1]$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \frac{(-1)^{n+1} x^n}{n}. \end{array} \right.$$

Si x est un point du segment $[0, 1]$, la suite (x^n/n) tend vers 0 en décroissant, donc on déduit du critère des séries alternées² que la série $\sum f_n(x)$ converge. La série de fonctions $\sum f_n$ converge donc simplement sur $[0, 1]$; nous notons $(R_n)_{n \geq 1}$ sa suite des restes. La convergence, pour tout $x \in [0, 1]$, de la série $\sum f_n(x)$ découlant de l'application du critère des séries alternées, on a, pour tout point x de $[0, 1]$ et tout entier $n \geq 1$:

$$|R_n(x)| = \left| \sum_{p=n+1}^{+\infty} f_p(x) \right| \leq |f_{n+1}(x)| = \frac{x^{n+1}}{n+1} \leq \frac{1}{n+1} < \frac{1}{n}.$$

Cette majoration ne dépend pas de la variable x et la suite $(1/n)$ converge vers 0, donc la suite de fonctions (R_n) converge uniformément vers l'application nulle sur le segment $[0, 1]$. La série de fonctions $\sum f_n$ converge donc uniformément sur $[0, 1]$. De plus, pour tout entier $n \geq 1$, f_n est bornée sur $[0, 1]$ et $\|f_n\|_\infty = |f_n(1)| = 1/n$, terme général d'une série divergente, donc la série de fonctions $\sum f_n$ ne converge pas normalement³ sur le segment $[0, 1]$. $\square$

1. C'est Karl Weierstrass qui a introduit cette notion comme critère suffisant de convergence uniforme d'une série de fonctions; voir [WALT], §7.4 et §7.5.

2. Théorème 7.3, page 112.

3. Elle converge uniformément sur le segment $[0, 1]$ vers la fonction $x \mapsto \ln(1+x)$.

DÉFINITION 13.4. — La série $\sum f_n$ converge absolument sur $\mathcal{A}$ si la série $\sum |f_n|$ converge simplement sur $\mathcal{A}$.

La convergence normale entraîne la convergence absolue, qui entraîne la convergence simple.

13.4. Série de fonctions qui converge absolument mais qui ne converge ni uniformément ni normalement.

Dans les exemples 13.1 (page 251) et 13.2 (page 252), la série de fonctions étudiée est une série de fonctions à valeurs positives, donc les notions de convergence simple et de convergence absolue coïncident pour ces séries. Ceci nous fournit donc des exemples de séries de fonctions qui convergent absolument mais qui ne convergent pas uniformément ; dans les deux cas, la convergence n'est donc pas normale. $\square$

Inversement, la convergence uniforme et, *a fortiori*, la convergence simple, n'entraîne pas la convergence absolue.

13.5. Série de fonctions qui converge uniformément mais qui ne converge pas absolument.

Nous reprenons la série $\sum f_n$ de l'exemple 13.3 (page 253). Nous avons démontré la convergence simple puis la convergence uniforme de cette série de fonctions sur le segment $[0, 1]$. Si la série $\sum f_n$ convergeait absolument sur $[0, 1]$, alors, pour tout $x \in [0, 1]$, $\sum |f_n(x)|$ serait une série convergente, en particulier la série $\sum |f_n(1)|$. Or, pour tout entier $n \geq 1$, $|f_n(1)| = 1/n$, terme général d'une série divergente, donc la série de fonctions $\sum f_n$ ne converge pas absolument sur $[0, 1]$. $\square$

Dans les exemples précédents, la convergence n'était pas normale, mais la convergence uniforme ou bien la convergence absolue faisait défaut. Cependant, les deux simultanément n'entraînent pas la convergence normale.

13.6. Série de fonctions convergeant absolument et uniformément mais qui ne converge pas normalement.

Nous introduisons la suite $(f_n)_{n \geq 1}$ d'applications de $\mathbb{R}_+$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R}_+ \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \begin{cases} \frac{1}{x} & \text{si } n \leq x < n+1, \\ 0 & \text{si } 0 \leq x < n \text{ ou } x \geq n+1. \end{cases} \end{array} \right.$$

Alors, pour tout nombre réel $x > 0$ et tout entier $n \geq 1$, $f_n(x) = 1/x$ si n est la partie entière de x et $f_n(x) = 0$ dans le cas où n n'est pas la partie entière de x . Pour tout réel $x \geq 1$, la partie entière m de x est un entier supérieur ou égal à 1, donc $f_m(x) = 1/x$ et, pour tout $n \in \mathbb{N}^* \setminus \{m\}$, $f_n(x) = 0$, donc la série $\sum_{n \geq 1} f_n(x)$ converge et sa somme vaut $1/x$. La série de fonctions $\sum f_n$ converge donc simplement sur l'intervalle $[1, +\infty[$ et sa somme est l'application :

$$\left| \begin{array}{l} S : [1, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto S(x) = \frac{1}{x}. \end{array} \right.$$

Par ailleurs, les fonctions f_n sont à valeurs réelles positives, ce qui entraîne la convergence absolue de la série $\sum f_n$ sur l'intervalle $[1, +\infty[$.

Nous notons $(R_n)_{n \geq 1}$ la suite des restes de la série $\sum f_n$. Soit $n \in \mathbb{N}^*$ et x un point de $[1, +\infty[$. Notons m la partie entière de x , qui appartient à $\mathbb{N}^*$. On a :

$$R_n(x) = \sum_{p=n+1}^{+\infty} f_p(x).$$

- Si $1 \leq x < n+1$, alors $m < n+1$ donc $m \leq n$, ce qui montre que $R_n(x) = 0$.
- Si $x \geq n+1$, alors $m \geq n+1$ donc $R_n(x) = 1/x \leq 1/(n+1) < 1/n$.

Par conséquent, pour tout $x \in [1, +\infty[$ et tout entier $n \geq 1$, $0 \leq R_n(x) \leq 1/n$. La majoration obtenue est indépendante de la variable x et la suite de terme général $1/n$ converge vers 0, donc la suite de fonctions (R_n) converge uniformément sur $\mathbb{R}$ vers l'application nulle, ce qui prouve la convergence uniforme sur $[1, +\infty[$ de la série de fonctions $\sum f_n$.

D'autre part, pour tout entier $n \geq 1$, la fonction f_n est positive et décroissante sur $[n, n+1[$ et nulle sur son complémentaire dans $[1, +\infty[$, donc f_n est bornée sur $[1, +\infty[$ et $\|f_n\|_\infty = f_n(n) = 1/n$, terme général d'une série divergente. En conclusion, la série de fonction $\sum f_n$ ne converge pas normalement sur $[1, +\infty[$. $\square$

13.7. Autre série de fonctions qui converge absolument et uniformément, mais pas normalement.

Nous considérons la suite $(f_n)_{n \geq 1}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \frac{1}{n} e^{-(x-n)^2}. \end{array} \right.$$

Clairement, pour tout nombre réel x et tout entier $n \geq 1$, $0 \leq f_n(x) \leq 1/n$. Par suite, pour tout $n \in \mathbb{N}^*$, la fonction f_n est bornée sur $\mathbb{R}$ et, comme $f_n(n) = 1/n$, $\|f_n\|_\infty = 1/n$, terme général d'une série divergente. Il en résulte que la série $\sum f_n$ ne converge pas normalement sur $\mathbb{R}$.

Soit x un nombre réel. On a, pour tout entier $n \geq 1$:

$$\ln(n^2 f_n(x)) = \ln(n e^{-(x-n)^2}) = \ln n - (x-n)^2 = -n^2 \left(\left(1 - \frac{x}{n}\right)^2 - \frac{\ln n}{n^2} \right),$$

donc la suite $(n^2 f_n(x))$ diverge vers $-\infty$, ce qui montre que la suite $(n^2 f_n(x))$ converge vers 0, d'où l'on déduit que la série $\sum f_n(x)$ converge.

Par conséquent la série de fonctions $\sum f_n$ converge simplement sur $\mathbb{R}$. Comme $f_n(x) > 0$ pour tout nombre réel x et tout entier $n \geq 1$, la série $\sum f_n$ converge absolument sur $\mathbb{R}$. La suite $(n^2 e^{-n^2})$ convergeant vers 0, la série $\sum e^{-n^2}$ converge, ce qui justifie l'existence de la constante :

$$C = \sum_{n=0}^{+\infty} e^{-n^2}.$$

Nous notons (R_n) la suite des restes de la série $\sum f_n$.

Soit x un nombre réel et n un entier ≥ 1 . Nous notons m la partie entière de x . Supposons d'abord que $x < n$. On a $m \leq x < n$, donc $n-m$ est un entier naturel. De plus $x < m+1$ donc, pour tout entier $p \geq n+1$, $0 \leq p - (m+1) \leq p - x$, ce qui montre que $(p - (m+1))^2 < (p-x)^2 = (x-p)^2$ donc $e^{-(x-p)^2} \leq e^{-(p-(m+1))^2}$.

Comme $\frac{1}{p} \leq \frac{1}{n+1} < \frac{1}{n}$ pour tout entier $p \geq n+1$, on obtient la majoration :

$$\begin{aligned} R_n(x) &= \sum_{p=n+1}^{+\infty} f_p(x) = \sum_{p=n+1}^{+\infty} \frac{1}{p} e^{-(x-p)^2} \leq \sum_{p=n+1}^{+\infty} \frac{1}{p} e^{-(p-(m+1))^2} \\ &\leq \frac{1}{n+1} \sum_{p=n+1}^{+\infty} e^{-(p-(m+1))^2} = \frac{1}{n+1} \sum_{q=n-m}^{+\infty} e^{-q^2} \leq \frac{1}{n} \sum_{q=0}^{+\infty} e^{-q^2} = \frac{C}{n}. \end{aligned}$$

Supposons maintenant que $x \geq n$. Comme $x < m+1$, on a, pour tout entier $p \geq m+1$, $0 \leq p-(m+1) < p-x$ donc $(p-(m+1))^2 < (p-x)^2 = (x-p)^2$, ce qui nous donne : $e^{-(x-p)^2} < e^{-(p-(m+1))^2}$, d'où, puisque $1/p \leq 1/(m+1) < 1/x \leq 1/n$:

$$\sum_{p=m+1}^{+\infty} f_p(x) \leq \sum_{p=m+1}^{+\infty} \frac{1}{p} e^{-(p-(m+1))^2} \leq \frac{1}{n} \sum_{p=m+1}^{+\infty} e^{-(p-(m+1))^2} = \frac{1}{n} \sum_{q=0}^{+\infty} e^{-q^2} = \frac{C}{n}.$$

Pour tout $p \in [n+1, m]$, $x-p \geq m-p \geq 0$, donc $(x-p)^2 \geq (m-p)^2$, d'où l'on déduit l'inégalité $e^{-(x-p)^2} \leq e^{-(m-p)^2}$, et $1/p \leq 1/(n+1) < 1/n$. Par conséquent :

$$\sum_{p=n+1}^m f_p(x) = \sum_{p=n+1}^m \frac{1}{p} e^{-(x-p)^2} \leq \sum_{p=n+1}^m \frac{1}{n} e^{-(m-p)^2} = \frac{1}{n} \sum_{q=0}^{m-n-1} e^{-q^2} \leq \frac{1}{n} \sum_{q=0}^{+\infty} e^{-q^2} = \frac{C}{n},$$

d'où finalement :

$$R_n(x) = \sum_{p=n+1}^m f_p(x) + \sum_{p=m+1}^{+\infty} f_p(x) \leq \frac{2C}{n}.$$

Nous avons établi que $0 \leq R_n(x) \leq (2C)/n$ pour tout $x \in \mathbb{R}$ et tout $n \in \mathbb{N}^*$. La majoration obtenue est indépendante de x et la suite $((2C)/n)$ tend vers 0, donc la suite de fonctions (R_n) converge uniformément sur $\mathbb{R}$ vers l'application nulle, ce qui prouve la convergence uniforme sur $\mathbb{R}$ de la série de fonctions $\sum f_n$. $\square$

■ Discontinuité de la somme d'une série de fonctions

THÉORÈME 13.3. — Si (f_n) est une suite de fonctions définies et continues sur une partie $\mathcal{A}$ de $\mathbb{R}$ et si la série $\sum f_n$ converge uniformément sur $\mathcal{A}$, sa somme est une fonction continue sur $\mathcal{A}$.

Ceci devient en général faux pour la convergence simple⁴.

13.8. Série de fonctions continues dont la somme est discontinue.

Nous considérons la suite $(f_n)_{n \geq 1}$ d'applications continues de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \frac{(-1)^n}{n} \sin(nx). \end{array} \right.$$

Soit x un réel. Nous introduisons les suites $(\varepsilon_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$ de termes généraux :

$$\varepsilon_n = (-1)^n \sin(nx) = \sin(n(x + \pi)) \quad \text{et} \quad v_n = \frac{1}{n}.$$

4. Cauchy affirme, en 1821, que la somme d'une série de fonctions continues est toujours continue ([CAU1], chapitre VI, § 1^{er}, 1^{er} théorème, pages 131 et 132). Cinq ans plus tard, Niels Abel fournit le contre-exemple 13.8.

Ces notations sont choisies dans le but d'appliquer la règle d'Abel (théorème 7.6, page 115) pour prouver que la série $\sum f_n(x)$ converge. On a, pour tout $n \in \mathbb{N}^*$:

$$\sum_{p=1}^n |v_{p+1} - v_p| = \sum_{p=1}^n \left| \frac{1}{p+1} - \frac{1}{p} \right| = \sum_{p=1}^n \left(\frac{1}{p} - \frac{1}{p+1} \right) = 1 - \frac{1}{n+1}$$

donc la série $\sum |v_{n+1} - v_n|$ converge.

Supposons d'abord que $x \in \pi\mathbb{Z}$. Soit $n \in \mathbb{N}^*$. Pour tout $p \in \llbracket 1, n \rrbracket$, $p(x + \pi) \in \pi\mathbb{Z}$ donc $\varepsilon_p = \sin(py) = 0$, ce qui montre que $|\sum_{p=1}^n \varepsilon_p| = 0$.

Supposons maintenant que x n'appartienne pas à $\pi\mathbb{Z}$. Alors $y = x + \pi \notin \pi\mathbb{Z}$ donc $|e^{iy}| \neq 1$, et $(y/2) \notin \pi\mathbb{Z}$ donc $\sin(y/2) \neq 0$. Pour tout entier $n \geq 1$, la somme $\sum_{p=1}^n \varepsilon_p = \sum_{p=1}^n \sin(py)$ est la partie imaginaire de :

$$\sum_{p=1}^n e^{ipy} = e^{iy} \sum_{k=0}^{n-1} e^{iky} = \frac{e^{iny} - 1}{e^{iy} - 1} e^{iy} = \frac{2ie^{i\frac{ny}{2}} \sin \frac{ny}{2}}{2ie^{i\frac{y}{2}} \sin \frac{y}{2}} e^{iy} = \frac{\sin \frac{ny}{2}}{\sin \frac{y}{2}} e^{i\frac{(n+1)y}{2}}$$

donc :

$$\left| \sum_{p=1}^n \varepsilon_p \right| \leq \left| \sum_{p=1}^n e^{ipy} \right| = \frac{\left| \sin \frac{ny}{2} \right|}{\left| \sin \frac{y}{2} \right|} \leq \frac{1}{\left| \sin \frac{y}{2} \right|}.$$

Par conséquent, dans les deux cas, la suite de terme général $|\sum_{p=1}^n \varepsilon_p|$ est bornée dans $\mathbb{R}$. Finalement la règle d'Abel (voir ci-dessus) montre que la série de fonctions $\sum f_n$ converge simplement sur $\mathbb{R}$. Nous notons g sa somme.

Nous cherchons les suites $(a_n)_{n \geq 0}$ et $(b_n)_{n \geq 1}$ des coefficients de Fourier⁵ de la fonction f de période 2π telle que $f(x) = x$ pour tout point x de $]-\pi, \pi[$ et $f(\pi) = 0$. Comme f est impaire, $a_n = 0$ pour tout $n \in \mathbb{N}$. On a, pour tout $n \in \mathbb{N}^*$:

$$\begin{aligned} b_n &= \frac{2}{\pi} \int_0^\pi x \sin(nx) dx = \frac{2}{\pi} \left[-\frac{x \cos(nx)}{n} \right]_{x=0}^{x=\pi} + \frac{2}{\pi} \int_0^\pi \frac{\cos(nx)}{n} dx \\ &= -\frac{2 \cos(n\pi)}{n} + \frac{2}{\pi} \left[\frac{\sin(nx)}{n^2} \right]_{x=0}^{x=\pi} = -2 \frac{(-1)^n}{n}. \end{aligned}$$

La fonction f est continue par morceaux et sa valeur en π est la demi-somme des limites à droite et à gauche, donc f est égale à sa régularisée. Par conséquent, pour tout réel x :

$$f(x) = -2 \sum_{n=1}^{+\infty} (-1)^n \frac{\sin(nx)}{n} = -2g(x),$$

donc $g = -\frac{1}{2}f$, ce qui montre que la fonction g est discontinue en π . $\square$

L'exemple précédent 13.8 est historique, mais il en existe de plus simples.

13.9. Autre série de fonctions continues de somme discontinue.

Nous considérons la suite $(f_n)_{n \geq 0}$ d'applications continues de $[0, 1]$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = (1-x)x^n. \end{array} \right.$$

5. Voir dans la suite de ce chapitre la définition 13.7, au bas de la page 264.

Pour tout point x de $[0, 1[$, la série $\sum f_n(x)$ est le produit par $(1-x)$ d'une série géométrique de raison x où $0 \leq x < 1$, donc elle converge et sa somme est égale à $(1-x)/(1-x) = 1$. De plus $f_n(1) = 0$ pour tout entier naturel n . La série de fonctions $\sum f_n$ converge donc simplement sur le segment $[0, 1]$ et sa somme est l'application :

$$\left| \begin{array}{l} S : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto S(x) = \begin{cases} 1 & \text{si } x \in [0, 1[, \\ 0 & \text{si } x = 1, \end{cases} \end{array} \right.$$

qui est discontinue en 1. $\square$

Dans l'exemple précédent 13.9, la fonction somme est bornée.

13.10. Série de fonctions continues dont la somme est discontinue et n'est pas bornée.

Nous introduisons la suite $(f_n)_{n \geq 0}$ d'applications continues de $\left[0, \frac{\pi}{2}\right]$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : [0, \pi/2] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = (\sin x)^n \cos x. \end{array} \right.$$

Comme $f_n(\pi/2) = 0$ pour tout $n \in \mathbb{N}$, la série $\sum f_n(\pi/2)$ converge et sa somme vaut 0. Pour tout $x \in [0, \pi/2[$, on a $0 \leq \sin x < 1$, donc la série $\sum f_n(x)$ converge et sa somme est égale à $(\cos x)/(1 - \sin x)$. La série de fonctions $\sum f_n$ converge donc simplement sur le segment $[0, \pi/2]$ et sa somme est l'application :

$$\left| \begin{array}{l} S : [0, \pi/2] \longrightarrow \mathbb{R} \\ x \longmapsto S(x) = \begin{cases} \frac{\cos x}{1 - \sin x} & \text{si } x \in \left[0, \frac{\pi}{2}\right[, \\ 0 & \text{si } x = \frac{\pi}{2}. \end{cases} \end{array} \right.$$

On a, pour tout point x de l'intervalle $[0, \pi/2[$, $S(x) = \frac{\sin\left(\frac{\pi}{2} - x\right)}{1 - \cos\left(\frac{\pi}{2} - x\right)}$.

Or $\lim_{x \rightarrow (\pi/2)^-} \left(\frac{\pi}{2} - x\right) = 0$, $\sin h \underset{h \rightarrow 0}{\sim} h$ et $1 - \cos h \underset{h \rightarrow 0}{\sim} \frac{h^2}{2}$, donc :

$$S(x) \underset{x \rightarrow (\pi/2)^-}{\sim} \frac{2}{\frac{\pi}{2} - x}.$$

Par conséquent $\lim_{(x \rightarrow (\pi/2)^-)} S(x) = +\infty$, ce qui montre que S est discontinue en $\pi/2$ et que S n'est pas bornée sur le segment $[0, \pi/2]$. $\square$

■ Interversions de sommations et de limites

THÉORÈME 13.4. — Soit a et b des points de $\overline{\mathbb{R}}$ tels que $-\infty < a < b \leq +\infty$ et (f_n) une suite de fonctions définies sur l'intervalle $I = [a, b[$. Si la série $\sum f_n$ converge simplement sur I , si F est sa somme, si la convergence est uniforme sur I — ou sur un voisinage à gauche de b — et si, pour tout n , $f_n(x)$ admet une limite ℓ_n dans $\mathbb{R}$ quand x tend vers b^- (vers $+\infty$ si $b = +\infty$), alors F admet une limite dans $\mathbb{R}$ quand x tend vers b^- , la série $\sum \ell_n$ converge et on a :

$$\lim_{x \rightarrow b^-} F(x) = \sum_{n=0}^{+\infty} \ell_n.$$

Bien entendu, ce théorème s'applique si $-\infty \leq a < b < +\infty$ et $I =]a, b]$, en remplaçant la limite en b à gauche par la limite en a à droite.

Ce résultat devient faux si la convergence de $\sum f_n$ n'est pas uniforme.

13.11. Série $\sum f_n$ de somme F sur $[0, +\infty[$ pour laquelle :

$$\lim_{x \rightarrow +\infty} F(x) \neq \sum_{n=0}^{+\infty} \lim_{x \rightarrow +\infty} f_n(x).$$

Soit F une application de $[0, +\infty[$ dans $\mathbb{R}$. Nous lui associons la suite $(f_n)_{n \geq 0}$ d'applications de $[0, +\infty[$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R}_+ \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \begin{cases} F(x) & \text{si } n \leq x < n+1, \\ 0 & \text{si } 0 \leq x < n \text{ ou } x \geq n+1. \end{cases} \end{array} \right.$$

Alors, pour tout nombre réel $x > 0$ et tout entier naturel n , $f_n(x) = F(x)$ si n est la partie entière de x et $f_n(x) = 0$ dans le cas où n n'est pas la partie entière de x . Pour tout nombre réel $x \geq 0$, la partie entière m de x est un entier naturel, donc $f_m(x) = F(x)$ et, pour tout $n \in \mathbb{N} \setminus \{m\}$, $f_n(x) = 0$, donc la série $\sum f_n(x)$ converge et sa somme vaut $F(x)$. La série de fonctions $\sum f_n$ converge donc simplement sur l'intervalle $[0, +\infty[$ et sa somme est la fonction F .

Pour tout entier naturel n , $f_n(x) = 0$ pour tout réel $x > n+1$, donc $\lim_{x \rightarrow +\infty} f_n(x) = 0$. On en déduit que $\sum_{n=0}^{+\infty} \lim_{x \rightarrow +\infty} f_n(x) = 0$.

Il suffit donc de choisir une application F de $\mathbb{R}_+$ dans $\mathbb{R}$ admettant en $+\infty$ une limite ℓ dans $\mathbb{R}$ différente de 0, ou n'admettant pas de limite dans $\mathbb{R}$ en $+\infty$ (il y en a !), pour obtenir une série $\sum f_n$ de somme F sur $[0, +\infty[$ telle que :

$$\lim_{x \rightarrow +\infty} F(x) \neq \sum_{n=0}^{+\infty} \lim_{x \rightarrow +\infty} f_n(x). \quad \square$$

On peut reprocher à ce contre-exemple de ne pas convenir si l'on souhaite que (f_n) soit une suite de fonctions continues.

13.12. Série $\sum f_n$ de fonctions continues sur $[0, +\infty[$ telle que :

$$\lim_{x \rightarrow +\infty} F(x) \neq \sum_{n=0}^{+\infty} \lim_{x \rightarrow +\infty} f_n(x).$$

Nous opérons d'une manière voisine de celle de l'exemple précédent 13.11 en nous donnant une application continue F de $[0, +\infty[$ dans $\mathbb{R}$ et en considérant la suite $(f_n)_{n \geq 0}$ d'applications de $\mathbb{R}_+$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R}_+ \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \begin{cases} (x-n+1)F(x) & \text{si } n \geq 1 \text{ et } n-1 \leq x < n, \\ (n+1-x)F(x) & \text{si } n \leq x < n+1, \\ 0 & \text{si } 0 \leq x < n-1 \text{ ou } x \geq n+1. \end{cases} \end{array} \right.$$

Pour tout entier naturel n , f_n est continue sur $[0, n-1[$ et $]n-1, n[$ (si $n \geq 1$) et sur $]n+1, +\infty[$, on a $f_n(n-1) = 0$ (si $n \geq 1$) et on obtient $F(n) = f_n(n)$ en substituant n à x dans l'expression $(x-n+1)F(x)$ et $0 = f_n(n+1)$ en substituant

n à x dans l'expression $(n+1-x)F(x)$, donc f_n est continue sur $\mathbb{R}_+$. Si x est un point de $\mathbb{R}_+$, alors, en notant m la partie entière de x , $f_m(x) = (m+1-x)F(x)$, $f_{m+1}(x) = (x-m)F(x)$ et $f_p(x) = 0$ pour tout $p \in \mathbb{N} \setminus \{m, m+1\}$, donc la série $\sum f_n(x)$ converge et :

$$S(x) = \sum_{p=0}^{+\infty} f_p(x) = f_m(x) + f_{m+1}(x) = F(x).$$

Par conséquent la série de fonctions $\sum f_n$ converge simplement sur $[0, +\infty[$ et sa somme est la fonction F . On conclut comme dans l'exemple précédent 13.11. $\square$

■ Séries entières

Nous remplaçons ici la droite réelle $\mathbb{R}$ par le plan complexe $\mathbb{C}$.

DÉFINITION 13.5. — Une série entière est une série $\sum_{n \geq 0} a_n z^n$ de fonctions où $(a_n)_{n \geq 0}$ est une suite de nombres réels ou complexes et z une variable complexe.

Cette définition n'est pas vraiment correcte mais elle est commode ; il s'agit en fait de la série de fonctions $\sum_{n \geq 0} u_n$ où, pour tout entier naturel n :

$$\begin{cases} u_n : \mathbb{C} \longrightarrow \mathbb{C} \\ z \longmapsto u_n(z) = a_n z^n. \end{cases}$$

THÉORÈME 13.5. — Soit $\sum a_n z^n$ une série entière. On pose, dans $\overline{\mathbb{R}}$:

$$\ell = \limsup_{n \rightarrow \infty} |a_n|^{1/n}.$$

a) Si $\ell > 0$ et si l'on pose $R = 1/\ell$, la série converge pour tous les complexes z tels que $|z| < R$, diverge pour tous les complexes z tel que $|z| > R$ et on ne peut rien affirmer de manière générale pour les complexes z de module R .

b) Si $\ell = 0$, la série converge pour tout $z \in \mathbb{C}$ et on pose $R = +\infty$.

c) Si $\ell = +\infty$, la série ne converge que pour $z = 0$ et on pose $R = 0$.

L'élément positif R de la droite numérique achevée $\overline{\mathbb{R}}$ est appelé le rayon de convergence⁶ de la série entière $\sum a_n z^n$ et, si $0 < R < +\infty$, son cercle de convergence est le cercle de centre 0 et de rayon R .

Pour déterminer le rayon de convergence, il est cependant plus simple, quand c'est possible, d'utiliser de critère de d'Alembert.

THÉORÈME 13.6. — Critère de d'Alembert⁷.

Si $a_n \neq 0$ à partir d'un certain rang et si la suite quotient, de terme général :

$$\frac{|a_{n+1}|}{|a_n|},$$

admet une limite ℓ dans la droite numérique achevée $\overline{\mathbb{R}}$, le rayon de convergence de la série entière $\sum a_n z^n$ est l'inverse R de ℓ — avec $R = +\infty$ si $\ell = 0$ et $R = 0$ si $\ell = +\infty$.

6. Cauchy démontre en 1821 que le rayon de convergence est l'inverse, si cette limite existe, de $\lim_{n \rightarrow \infty} |a_n|^{1/n}$ ([CAU1], chapitre VI, § 4^e, 1^{er} théorème, page 151). Jacques Hadamard améliore ce résultat en 1892 en utilisant la limite supérieure.

7. Ce résultat est en fait énoncé par Cauchy en 1821 ([CAU1], ch. VI, § 4^e, 2^e théorème, p. 152).

13.13. Série entière qui diverge sur tout le cercle de convergence.

Nous considérons la série entière $\sum z^n$; ici $a_n = 1$ pour tout $n \in \mathbb{N}$. On a, pour tout entier naturel n :

$$\rho_n = \frac{|a_{n+1}|}{|a_n|} = \frac{1}{1} = 1,$$

donc la suite quotient (ρ_n) converge vers 1, ce qui montre que le rayon de convergence de la série entière $\sum z^n$ est égal à 1.

Si z est un nombre complexe et si $|z| < 1$, la série $\sum z^n$ converge et $\sum_{n=0}^{+\infty} z^n = \frac{1}{1-z}$.

Pour un nombre complexe z de module 1, $|z^n| = 1$ pour tout $n \in \mathbb{N}$, donc la suite (z^n) ne converge pas vers 0, d'où l'on déduit que la série diverge. Par conséquent la série entière $\sum z^n$ diverge sur tout le cercle de convergence. Remarquons que si z est un complexe de module 1 et si $z \neq 1$, l'expression $1/(1-z)$ a un sens mais n'est bien sûr pas la somme de la série. $\square$

13.14. Série entière qui converge sur tout le cercle de convergence.

Nous introduisons la série entière $\sum \frac{1}{n^2} z^n$. On a, pour tout entier $n \geq 1$:

$$\rho_n = \frac{\frac{1}{(n+1)^2}}{\frac{1}{n^2}} = \left(\frac{n}{n+1}\right)^2,$$

donc la suite (ρ_n) converge vers 1, ce qui prouve que le rayon de convergence de la série est égal à 1. Si $z \in \mathbb{C}$ et si z est de module 1, alors, pour tout entier $n \geq 1$:

$$\left| \frac{z^n}{n^2} \right| = \frac{1}{n^2}$$

et la série $\sum 1/n^2$ converge, donc la série $\sum \frac{1}{n^2} z^n$ converge. Ainsi la série converge sur tout le cercle de convergence. $\square$

13.15. Série entière qui converge en certains points du cercle de convergence et diverge en d'autres points de ce cercle.

Nous considérons la série entière $\sum \frac{1}{n} z^n$. On a, pour tout entier $n \geq 1$:

$$\rho_n = \frac{\frac{1}{n+1}}{\frac{1}{n}} = \frac{n}{n+1},$$

donc la suite (ρ_n) converge vers 1, ce qui prouve que le rayon de convergence de la série est égal à 1. Pour $z = 1$, il s'agit de la série harmonique, donc la série diverge, alors que pour $z = -1$, elle converge par le critère des séries alternées⁸. $\square$

Si $\sum a_n z^n$ et $\sum b_n z^n$ sont des séries entières, le rayon de convergence de la somme $\sum (a_n + b_n) z^n$ comme celui du produit $\sum c_n z^n$, où $c_n = \sum_{k=0}^n a_k b_{n-k}$ pour tout n ,

8. Théorème 7.3, page 112.

est au moins égal au plus petit des rayons de convergence de ces deux séries, et il y a égalité si ces rayons sont différents. L'inégalité peut être stricte pour la somme, ce qui est clair en prenant une série et son opposée.

13.16. Rayon de convergence du produit de deux séries entières strictement plus grand que chacun des rayons de convergence des deux séries.

Comme dans l'exemple 7.26 (page 127) nous considérons les suites $(u_n)_{n \geq 0}$ et $(v_n)_{n \geq 0}$, de termes généraux :

$$u_n = \begin{cases} 2 & \text{si } n = 0, \\ 2^n & \text{si } n \geq 1 \end{cases} \quad \text{et} \quad v_n = \begin{cases} -1 & \text{si } n = 0, \\ 1 & \text{si } n \geq 1. \end{cases}$$

La série entière produit des séries entières $\sum u_n z^n$ et $\sum v_n z^n$ est la série entière $\sum w_n z^n$ où, pour tout entier naturel n :

$$w_n = \sum_{k=0}^n u_k v_{n-k}.$$

Nous avons vu dans l'exemple 7.26 que $w_0 = -2$ et $w_n = 0$ pour tout entier $n \geq 1$. Clairement, le rayon de convergence de la série entière $\sum u_n z^n$ est égal à $1/2$ et celui de $\sum v_n z^n$ à 1 , alors que le rayon de convergence de la série entière produit $\sum w_n z^n$ est infini. $\square$

■ Série de Taylor

Nous revenons à des suites à valeurs réelles et des fonctions de $\mathbb{R}$ dans $\mathbb{R}$.

THÉORÈME 13.7. — Si le rayon de convergence R de la série entière $\sum a_n z^n$ est strictement positif, la fonction :

$$f : x \mapsto f(x) = \sum_{n=0}^{+\infty} a_n x^n$$

est définie et indéfiniment dérivable sur l'intervalle ouvert $] -R, R[$ et, pour tout entier naturel n , $f^{(n)}(0) = n! a_n$.

On peut se poser le problème inverse : étant donnée une fonction f définie et de classe $\mathcal{C}^\infty$ sur un voisinage de 0 , le rayon de convergence R de la série entière :

$$\sum_{n \geq 0} \frac{f^{(n)}(0)}{n!} z^n$$

est-il strictement positif ? Dans l'affirmative a-t-on $f(x) = \sum_{n=0}^{+\infty} \frac{f^{(n)}(0)}{n!} x^n$ pour tout point x de $] -R, R[$?

DÉFINITION 13.6. — Si f est une fonction définie et de classe $\mathcal{C}^\infty$ sur un voisinage de 0 , la série entière :

$$\sum_{n \geq 0} \frac{f^{(n)}(0)}{n!} z^n$$

est appelée la série de Taylor de f en 0 .

13.17. Fonction f indéfiniment dérivable sur $\mathbb{R}$ dont la série de Taylor admet sur $\mathbb{R}$ une somme différente⁹ de f .

Nous reprenons la fonction f de l'exemple 9.25 (page 178). C'est l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ e^{-\frac{1}{x^2}} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

Nous avons vu dans l'exemple 9.25 que f est indéfiniment dérivable sur $\mathbb{R}$ et que, pour tout entier naturel n , $f^{(n)}(0) = 0$.

On en déduit que le rayon de convergence de la série de Taylor associée est infini et que sa somme est l'application nulle $\mathbf{0} : x \mapsto \mathbf{0}(x) = 0$ de $\mathbb{R}$ dans $\mathbb{R}$, alors que $f(x) > 0$ pour tout nombre réel $x \neq 0$. $\square$

Nous rappelons un théorème classique sur la dérivation d'une série de fonctions dérivables, théorème utilisé dans ce qui suit.

THÉORÈME 13.8. — Soit I un intervalle et (f_n) une suite de fonctions définies et dérivables sur I . Si la série de fonctions $\sum f_n$ converge simplement sur I et si la série de fonctions $\sum f_n'$ converge uniformément sur I , alors la série de fonctions $\sum f_n$ converge uniformément sur toute partie bornée de I , la somme f de la série de fonctions $\sum f_n$ est dérivable sur I et, pour tout point x de I :

$$f'(x) = \sum_{n=0}^{+\infty} f_n'(x).$$

13.18. Fonction indéfiniment dérivable sur $\mathbb{R}$ dont la série de Taylor possède un rayon de convergence nul.

Nous introduisons la suite $(f_n)_{n \geq 0}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} f_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = e^{-n} \cos(n^2 x). \end{array} \right.$$

On a, pour tout nombre réel x et tout entier naturel n , $|f_n(x)| \leq e^{-n}$. Comme la série $\sum e^{-n}$ converge, la série de fonctions $\sum f_n$ converge normalement donc uniformément sur $\mathbb{R}$ et sa somme :

$$f : x \mapsto f(x) = \sum_{n=0}^{+\infty} f_n(x)$$

est une application continue de $\mathbb{R}$ dans $\mathbb{R}$. Pour tout entier naturel n , la fonction f_n est dérivable sur $\mathbb{R}$ et, pour tout réel x , $f_n'(x) = -n^2 e^{-n} \sin(n^2 x)$. On a, pour tout $x \in \mathbb{R}$ et tout $n \in \mathbb{N}$, $|f_n'(x)| \leq n^2 e^{-n}$. La suite de terme général $n^2 \times (n^2 e^{-n})$ converge vers 0, donc la série $\sum n^2 e^{-n}$ converge, ce qui montre que la série de fonctions $\sum f_n'$ converge normalement donc uniformément sur $\mathbb{R}$. Par conséquent la fonction f est dérivable sur $\mathbb{R}$ et, pour tout nombre réel x :

$$f'(x) = \sum_{n=0}^{+\infty} f_n'(x).$$

9. Cet exemple est proposé par Cauchy en 1823 ([CAU2], TRENTE-HUITIÈME LEÇON, page 152).

Pour tout entier naturel n , f_n est indéfiniment dérivable sur $\mathbb{R}$ et, pour tout entier naturel k et tout nombre réel x :

$$f_n^{(2k)}(x) = (-1)^k n^{4k} e^{-n} \cos(n^2 x) \text{ et } f_n^{(2k+1)}(x) = (-1)^{k+1} n^{4k+2} e^{-n} \sin(n^2 x),$$

ce qui donne les majorations $|f_n^{(2k)}(x)| \leq n^{4k} e^{-n}$ et $|f_n^{(2k+1)}(x)| \leq n^{4k+2} e^{-n}$, qui, puisque les séries $\sum n^{4k} e^{-n}$ et $\sum n^{4k+2} e^{-n}$ convergent — même raisonnement que dans ce qui précède —, montrent que les séries de fonctions $\sum f_n^{(2k)}$ et $\sum f_n^{(2k+1)}$ convergent normalement donc uniformément sur $\mathbb{R}$. En raisonnant de proche en proche comme pour la dérivation d'ordre 1, on en déduit que la fonction f est indéfiniment dérivable sur $\mathbb{R}$ et que, pour tout entier naturel k et tout réel x :

$$f^{(2k)}(x) = (-1)^k \sum_{n=0}^{+\infty} n^{4k} e^{-n} \cos(n^2 x)$$

et :

$$f^{(2k+1)}(x) = (-1)^{k+1} \sum_{n=0}^{+\infty} n^{4k+2} e^{-n} \sin(n^2 x).$$

La série de Taylor de f en 0 est la série entière $\sum a_q z^q$ où $(a_q)_{q \geq 0}$ est la suite de terme général :

$$a_q = \frac{f^{(p)}(0)}{p!} = \begin{cases} 0 & \text{si } q \text{ est impair,} \\ (-1)^k \frac{1}{(2k)!} \sum_{n=0}^{+\infty} n^{4k} e^{-n} & \text{si } q \text{ est pair et } q = 2k \text{ où } k \in \mathbb{N}. \end{cases}$$

Pour tout entier naturel pair q , on a, en posant $q = 2k$ où $k \in \mathbb{N}$:

$$|a_q| = \sum_{n=0}^{+\infty} \frac{n^{2q} e^{-n}}{q!} \geq \frac{q^{2q} e^{-q}}{q!} \geq \frac{q^q}{e^q}$$

car la somme d'une série à termes positifs est supérieure à chacun de ses termes et que $q! \leq q^q$. On en déduit que :

$$\limsup_{q \rightarrow \infty} |a_q|^{1/q} \geq \lim_{q \rightarrow \infty} \left(\frac{q^q}{e^q} \right)^{1/q} = \lim_{q \rightarrow \infty} \frac{q}{e} = +\infty.$$

En conclusion, le rayon de convergence de la série de Taylor de f en 0 est nul¹⁰. $\square$

■ Séries de Fourier

DÉFINITION 13.7. — Soit f une application de $\mathbb{R}$ dans $\mathbb{R}$ de période 2π et intégrable sur au moins un segment d'amplitude 2π . On associe à f les suites $(a_n)_{n \geq 0}$ et $(b_n)_{n \geq 1}$ de termes généraux :

$$a_n = \int_0^{2\pi} f(t) \cos(nt) dt \text{ et } b_n = \int_0^{2\pi} f(t) \sin(nt) dt$$

et la série de fonctions $a_0/2 + \sum_{n \geq 1} \varphi_n$ où, pour tout entier $n \geq 1$, $\varphi_n : \mathbb{R} \rightarrow \mathbb{R}$ est définie par $\varphi_n(x) = a_n \cos(nx) + b_n \sin(nx)$, série de fonctions que l'on appelle la série de Fourier¹¹ de f .

10. Mathias Lerch en 1888 et Alfred Pringsheim en 1893 proposent un autre exemple du même type, à savoir $f : x \mapsto f(x) = \sum_{n=0}^{+\infty} (\cos(2nx))/n!$ ([HAIR], chapitre III, §7, exercice 7.6).

11. Du nom du mathématicien français Joseph Fourier (1768-1830) qui introduisit ces séries en affirmant que toute fonction f est égale à la somme de cette série de fonctions.

Pour une fonction f continue sur $\mathbb{R}$, la somme de sa série de Fourier n'est pas forcément égale à f , comme nous allons le voir ; cependant il y a convergence vers f au sens de Cesàro.

13.19. Fonction continue sur $\mathbb{R}$ dont la série de Fourier ne converge pas¹² en 0.

Nous considérons la suite $(q_n)_{n \geq 1}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} q_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto q_n(x) = 2 \sin(2nx) \sum_{k=1}^n \frac{\sin(kx)}{k}. \end{array} \right.$$

Nous admettons provisoirement que, si $n \in \mathbb{N}^*$ et si $x \in \mathbb{R}$, $\left| \sum_{k=1}^n \frac{\sin(kx)}{k} \right| \leq \pi + 1$ ce qui montre que $|q_n(x)| \leq 2(\pi + 1)$.

Quels que soient les réels a et b , $\sin(a)\sin(b) = \frac{1}{2}(\cos(a-b) - \cos(a+b))$ donc, pour tout $n \in \mathbb{N}^*$ et tout $x \in \mathbb{R}$:

$$q_n(x) = \begin{cases} \frac{\cos(nx)}{n} + \frac{\cos((n+1)x)}{n-1} + \dots + \frac{\cos((2n-1)x)}{1} \\ - \frac{\cos((2n+1)x)}{1} - \dots - \frac{\cos((3n)x)}{n}. \end{cases}$$

Nous posons, pour tout entier $k \geq 1$, $n_k = 2^{k^3}$. Alors, pour tout $k \in \mathbb{N}^*$, $3n_k < n_{k+1}$ et, pour tout réel x :

$$\left| \frac{1}{k^2} q_{n_k}(x) \right| \leq \frac{1}{k^2} |q_{n_k}(x)| \leq \frac{\pi + 1}{k^2}.$$

Comme la série $\sum_{k \geq 1} (\pi + 1)/k^2$ converge, la série de fonctions $\sum_{k \geq 1} (1/k^2) q_{n_k}$ converge normalement donc uniformément sur $\mathbb{R}$, ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \sum_{k=1}^{+\infty} \frac{1}{k^2} q_{n_k}(x) \end{array} \right.$$

et prouve que f est continue sur $\mathbb{R}$. De plus, pour tout entier $k \geq 1$, la fonction q_{n_k} est paire et 2π en est une période, donc la fonction f est paire et admet 2π pour période. Pour tout entier $k \geq 1$ et tout nombre réel x , $q_{n_k}(x)$ est la somme de $2n_k$ termes en $\cos(px)$ où $n_k \leq p \leq 3n_k$ donc, comme $3n_k < n_{k+1}$, on en déduit que si $k, \ell \in \mathbb{N}^*$ et $k \neq \ell$, les termes en $\cos(ix)$ de la décomposition de q_{n_k} et ceux en $\cos(jx)$ de celle de q_{n_ℓ} sont tels que $i \neq j$.

Nous posons, pour tout entier $m \geq 1$, $S_m = \sum_{k=1}^m \frac{1}{k^2} q_{n_k}$.

Soit x un nombre réel. En remplaçant chaque q_{n_k} par sa décomposition en fonction des $\cos(px)$, on obtient une série trigonométrique. De plus la suite de fonctions $(S_m)_{m \geq 1}$ converge uniformément vers f sur $\mathbb{R}$, donc, pour tout entier $p \geq 1$,

12. Cet exemple est dû au mathématicien hongrois Lipot Fejer (1880-1959).

la suite $(A_{p,m})_{m \geq 1}$, de terme général $A_{p,m} = \int_0^{2\pi} \cos(px) S_m(x) dx$, converge vers l'intégrale :

$$\int_0^{2\pi} \cos(px) f(x) dx.$$

Ceci montre que la série trigonométrique obtenue est la série de Fourier de f .

Pour $x = 0$, la somme des termes de la série trigonométrique pour les indices compris entre n_k et $2n_k - 1$ est, pour tout entier $k \geq 1$:

$$\frac{1}{k^2} \left(\frac{1}{n_k} + \frac{1}{n_k - 1} + \cdots + 1 \right) \underset{k \rightarrow \infty}{\sim} \frac{\ln(n_k)}{k^2} = \frac{\ln(2^{k^3})}{k^2} = k \ln 2.$$

Or la suite $(k \ln 2)_{k \geq 1}$ tend vers $+\infty$, donc la série de Fourier diverge en 0.

Il nous reste à montrer que, pour tout $n \in \mathbb{N}^*$ et tout $x \in \mathbb{R}$, $\left| \sum_{k=1}^n \frac{\sin(kx)}{k} \right| \leq \pi + 1$.

Pour tout entier $k \geq 1$, la fonction $x \mapsto \sin(kx)$ est paire et 2π en est une période, donc il suffit de justifier cette majoration pour tout point x de $[0, \pi]$, donc, la majoration étant évidente pour $x = 0$ et $x = \pi$, pour tout point x de $]0, \pi[$.

Soit donc x un point de $]0, \pi[$ et $n \in \mathbb{N}^*$. Nous notons q la partie entière de π/x . Alors q est un entier relatif et $q \leq \pi/x < q + 1$, donc $qx \leq \pi < (q + 1)x$; de plus $\pi/x > 1$, donc q appartient à $\mathbb{N}^*$.

• Nous supposons d'abord que $n \leq q$. Pour tout réel $y > 0$, $\sin y < y$ donc, pour tout entier $k \geq 1$, $\sin(kx) \leq kx$, et, pour tout $k \in \llbracket 1, q \rrbracket$, on a $0 < kx \leq qx \leq \pi$ donc $\sin(kx) \geq 0$. Par conséquent :

$$0 \leq \sum_{k=1}^n \frac{\sin(kx)}{k} \leq \sum_{k=1}^q \frac{\sin(kx)}{k} \leq \sum_{k=1}^q \frac{kx}{k} \leq \sum_{k=1}^q x = qx \leq \pi < \pi + 1.$$

• Nous supposons enfin que $n > q$. La même démonstration que ci-dessus donne l'encadrement $0 \leq \sum_{k=1}^q (\sin(kx))/k \leq \pi$ qui s'écrit :

$$(1) \quad \left| \sum_{k=1}^q \frac{\sin(kx)}{k} \right| \leq \pi.$$

Nous posons $u_k = \sum_{p=q+1}^k \sin(px)$ pour tout $k \in \llbracket q+1, n \rrbracket$.

Comme dans la preuve de la convergence simple de la série $\sum f_n(x)$ de l'exemple 13.8 (pages 256 et 257), $x \notin \pi\mathbb{Z}$ donc, pour tout entier $k > q$:

$$\sum_{p=q+1}^k e^{ipx} = \frac{e^{i(k-q)x} - 1}{e^{ix} - 1} e^{i(q+1)x} = \frac{2ie^{\frac{i(k-q)x}{2}} \sin \frac{(k-q)x}{2}}{2ie^{\frac{ix}{2}} \sin \frac{x}{2}} e^{i(q+1)x} = \frac{\sin \frac{(k-q)x}{2}}{\sin \frac{x}{2}} e^{iy}$$

où $y = \frac{k+q+1}{2}x$, ce qui, puisque $0 < \frac{x}{2} < \frac{\pi}{2}$, donne :

$$|u_k| \leq \left| \sum_{p=q+1}^k e^{ipx} \right| = \frac{\left| \sin \frac{(k-q)x}{2} \right|}{\left| \sin \frac{x}{2} \right|} \leq \frac{1}{\left| \sin \frac{x}{2} \right|} = \frac{1}{\sin \frac{x}{2}}.$$

On a $u_{q+1} = \sin((q+1)x)$ et, pour tout $k \in \llbracket q+2, n \rrbracket$, $\sin(kx) = u_k - u_{k-1}$, donc :

$$\sum_{k=q+1}^n \frac{\sin(kx)}{k} = \frac{u_{q+1}}{q+1} + \sum_{k=q+2}^n \frac{u_k - u_{k-1}}{k} = \frac{u_{q+1}}{q+1} + \sum_{k=q+2}^n \frac{u_k}{k} - \sum_{k=q+2}^n \frac{u_{k-1}}{k},$$

ce qui, en décalant les indices dans la dernière somme, donne :

$$\begin{aligned} \sum_{k=q+1}^n \frac{\sin(kx)}{k} &= \frac{u_{q+1}}{q+1} + \sum_{k=q+2}^n \frac{u_k}{k} - \sum_{k=q+1}^{n-1} \frac{u_k}{k+1} \\ &= \left(\frac{1}{q+1} - \frac{1}{q+2} \right) u_{q+1} + \left(\frac{1}{q+2} - \frac{1}{q+3} \right) u_{q+2} + \cdots \\ &\quad \cdots + \left(\frac{1}{n-1} - \frac{1}{n} \right) u_{n-1} + \frac{1}{n} u_n. \end{aligned}$$

Or, pour tout $k \in \llbracket q+1, n \rrbracket$, $|u_k| \leq \frac{1}{\sin(x/2)}$, donc :

$$\begin{aligned} \left| \sum_{k=q+1}^n \frac{\sin(kx)}{k} \right| &\leq \left(\left(\frac{1}{q+1} - \frac{1}{q+2} \right) + \left(\frac{1}{q+2} - \frac{1}{q+3} \right) + \cdots + \left(\frac{1}{n-1} - \frac{1}{n} \right) + \frac{1}{n} \right) \frac{1}{\sin \frac{x}{2}} \\ &\leq \frac{1}{(q+1) \sin \frac{x}{2}}. \end{aligned}$$

De plus, sinus étant concave sur $[0, \pi]$, on a, pour tout point y de $\left[0, \frac{\pi}{2}\right]$:

$$\sin y \geq \frac{2}{\pi} y,$$

ce qui, compte tenu de ce que $\pi < (q+1)x$, montre que :

$$\frac{1}{(q+1) \sin \frac{x}{2}} = \frac{\frac{x}{2}}{\sin \frac{x}{2}} \frac{2}{(q+1)x} \leq \frac{\pi}{2} \frac{2}{\pi} = 1,$$

d'où l'on déduit l'inégalité :

$$(2) \quad \left| \sum_{k=q+1}^n \frac{\sin(kx)}{k} \right| \leq 1.$$

Il découle de la conjonction des inégalités (1) et (2) que l'on a :

$$\left| \sum_{k=1}^n \frac{\sin(kx)}{k} \right| \leq \pi + 1,$$

ce qui achève la démonstration. $\square$

Chapitre 14

Fonctions de plusieurs variables

On trouve déjà chez Leibniz, en 1694, des fonctions de plusieurs variables dans des études de géométrie. Cependant, les premiers travaux intéressants sur ce sujet sont liées à la résolution de problèmes de mécanique. Ils débutent avec Euler en 1734 et D'Alembert en 1743 avec son Traité de dynamique. Ce n'est cependant qu'au début du XX^e siècle que l'approche actuelle s'élabore avec l'introduction des espaces vectoriels normés et la notion de différentielle vue comme une application linéaire.

Une fonction de plusieurs variables est une application d'une partie de $\mathbb{R}^n$ dans $\mathbb{R}^p$ où n est un entier supérieur ou égal à 2 et p un entier supérieur ou égal à 1. La continuité et la différentiabilité de telles fonctions se définissent grâce à la structure d'espace vectoriel normé de $\mathbb{R}^n$. Rappelons que sur un espace vectoriel réel de dimension finie, en particulier sur $\mathbb{R}^n$ où n est un entier ≥ 2 , les normes sont deux à deux équivalentes.

Si $p \geq 2$ et si f est une application d'une partie de $\mathbb{R}^n$ dans $\mathbb{R}^p$, on introduit les p fonctions composantes $f_1, f_2, \dots, f_p$ de f , définies, pour tout point x du domaine de f , par $f(x) = (f_1(x), f_2(x), \dots, f_p(x))$. La fonction f est continue (resp. différentiable) si, et seulement si, chacune de ses fonctions composantes l'est aussi. Ainsi la plupart des propriétés des fonctions composantes se transmettent à la fonction elle-même. Ceci explique que les contre-exemples proposés dans ce chapitre porteront essentiellement sur des fonctions à valeurs dans $\mathbb{R}$.

Il n'en est pas de même pour l'ensemble de départ. L'entier n étant supérieur ou égal à 2, de nombreuses difficultés apparaissent pour une fonction définie sur une partie de $\mathbb{R}^n$ par rapport aux fonctions définies sur une partie de $\mathbb{R}$: le point délicat est le passage de la dimension 1 à la dimension 2. Le caractère ordonné et linéaire de $\mathbb{R}$ sous-tend de nombreux résultats en dimension 1. Aussi nos exemples concerneront des fonctions définies sur une partie de $\mathbb{R}^2$, les dimensions supérieures n'amenant pas de problèmes particuliers. Nous utiliserons sur $\mathbb{R}^2$ la norme euclidienne canonique $\|\cdot\| : z = (x, y) \mapsto \|z\| = \sqrt{x^2 + y^2}$.

L'objet de ce chapitre est l'étude de problèmes de continuité, de différentiabilité, d'extremums et d'intégration.

Les contre-exemples de ce chapitre concernent, sauf avis contraire, des fonctions de $\mathbb{R}^2$ dans $\mathbb{R}$, c'est-à-dire des applications d'une partie $\mathcal{D}$ de $\mathbb{R}^2$ — leur domaine — dans $\mathbb{R}$ et, si $\mathcal{A}$ est une partie de $\mathbb{R}^2$, une telle fonction f est définie sur $\mathcal{A}$ si l'ensemble $\mathcal{A}$ est inclus dans le domaine $\mathcal{D}$ de f . Pour les théorèmes d'inversion, il s'agira de fonctions de $\mathbb{R}^2$ dans $\mathbb{R}^2$.

■ Continuité

DÉFINITION 14.1. — Une fonction f , définie sur un ouvert U de $\mathbb{R}^2$, est partiellement continue en un point $a = (a_1, a_2)$ de U si les fonctions partielles $x \mapsto f(x, a_2)$ et $y \mapsto f(a_1, y)$ sont continues respectivement en a_1 et en a_2 .

La continuité en $a = (a_1, a_2)$ entraîne la continuité partielle. Nous démontrons que la réciproque est fautive.

14.1. Fonction discontinue en $(0, 0)$ mais partiellement continue en ce point.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R}^2 \\ (x, y) \longmapsto f(x, y) = \begin{cases} 0 & \text{si } (x, y) = (0, 0), \\ \frac{xy}{x^2 + y^2} & \text{si } (x, y) \neq (0, 0). \end{cases} \end{array} \right.$$

Les fonctions partielles de f en $(0, 0)$ sont les applications $y \mapsto f(0, y) = 0$ et $x \mapsto f(x, 0) = 0$ de $\mathbb{R}$ dans $\mathbb{R}$, identiquement nulles donc continues en 0.

Nous posons $\varepsilon_0 = 1/2$. Soit α un réel > 0 . En choisissant un nombre réel x tel que $0 < x < \alpha/2$ — par exemple $x = \alpha/4$ —, on a $\|(x, x)\| = \sqrt{x^2 + x^2} = (\sqrt{2})x < \alpha$ et :

$$|f(x, x) - f(0, 0)| = |f(x, x)| = \frac{x^2}{2x^2} = \frac{1}{2} \geq \varepsilon_0.$$

Il en résulte que la fonction f est discontinue en $(0, 0)$. $\square$

On peut reprocher à la continuité partielle d'être arbitraire, puisqu'elle privilégie les axes de coordonnées : elle n'est pas invariante par changement d'axes. Si dans l'exemple précédent on remplace la base canonique de $\mathbb{R}^2$ par (e_1, e_2) où $e_1 = (1, 1)$ et $e_2 = (1, -1)$, les applications partielles ne sont plus continues en $(0, 0)$.

14.2. Fonction discontinue en $(0, 0)$ dont la restriction à toute droite passant par l'origine est continue en $(0, 0)$.

Nous identifions $\mathbb{R}^2$ au plan complexe $\mathbb{C}$ en identifiant le couple (x, y) au nombre complexe $z = x + iy$. Rappelons que, pour tout nombre complexe $z \neq 0$, il existe un couple (r, θ) et un seul tel que $z = re^{i\theta}$, $r \in \mathbb{R} \setminus \{0\}$ et $\theta \in]0, \pi]$.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{C} \longrightarrow \mathbb{R} \\ z \longmapsto f(z) = \begin{cases} 0 & \text{si } z = 0, \\ \frac{r}{\theta} & \text{où } (r, \theta) \text{ est l'unique couple tel que} \\ & r \in \mathbb{R} \setminus \{0\}, \theta \in]0, \pi] \text{ et } z = re^{i\theta}. \end{cases} \end{array} \right.$$

Soit D une droite de $\mathbb{C}$ passant par l'origine, c'est-à-dire une droite vectorielle de l'espace vectoriel réel $\mathbb{C}$. Il existe un réel $\alpha \in]0, \pi]$ et un seul tel que $(e^{i\alpha})$ est une base de D . La droite D est donc l'ensemble des $re^{i\alpha}$ pour r parcourant $\mathbb{R}$ et, pour tout nombre réel r , $f(re^{i\alpha}) = r/\alpha$, égalité vraie pour $r \neq 0$ et pour $r = 0$. La restriction de f à D est donc continue en 0.

Nous introduisons la suite $(z_n)_{n \geq 1}$, de terme général $z_n = r_n e^{i r_n}$ où $r_n = \pi/n$. Pour tout entier $n \geq 1$, $|z_n| = \pi/n$, donc la suite (z_n) converge vers 0. De plus $f(z_n) = 1$ pour tout $n \in \mathbb{N}^*$, donc la suite $(f(z_n))$ tend vers 1 $\neq 0 = f(0)$, ce qui prouve que la fonction f est discontinue en 0. $\square$

■ Différentiabilité

DÉFINITION 14.2. — Soit E et F des espaces vectoriels normés, U un ouvert non vide de E et f une fonction de E dans F définie sur U . Si a est un point de U , la fonction f est différentiable en a s'il existe une application linéaire continue ℓ de E dans F telle que $f(a+h) = f(a) + \ell(h) + o_{h \rightarrow 0}(\|h\|)$ et, dans ce cas, ℓ est la différentielle de f en a , notée $df(a)$.

Si E et F sont des espaces vectoriels réels et si $\ell \in \mathcal{L}(E, F)$, nous notons dans ce qui suit $\ell \cdot x$ au lieu de $\ell(x)$ l'image par ℓ de $x \in E$. Rappelons que, si $n \in \mathbb{N}^*$, toute application linéaire de $\mathbb{R}^n$ dans un espace vectoriel normé F est continue.

DÉFINITION 14.3. — Soit U un ouvert non vide de $\mathbb{R}^2$, f une fonction de $\mathbb{R}^2$ dans $\mathbb{R}$ définie sur U et $a = (a_1, a_2)$ un point de U .

a) La fonction f admet en a une dérivée partielle par rapport à la première variable si la fonction partielle $f_{1,a} : x \mapsto f_{1,a}(x) = f(x, a_2)$ est dérivable en a_1 et, dans ce cas, la dérivée partielle de f en a par rapport à la première variable, que l'on nomme x , est le nombre réel :

$$\frac{\partial f}{\partial x}(a) = f_{1,a}'(a_1).$$

b) La fonction f admet en a une dérivée partielle par rapport à la deuxième variable si la fonction partielle $f_{2,a} : y \mapsto f_{2,a}(y) = f(a_1, y)$ est dérivable en a_2 et, dans ce cas, la dérivée partielle de f en a par rapport à la deuxième variable, que l'on nomme y , est le nombre réel :

$$\frac{\partial f}{\partial y}(a) = f_{2,a}'(a_2).$$

THÉORÈME 14.1. — Soit U un ouvert non vide de $\mathbb{R}^2$, f une fonction de $\mathbb{R}^2$ dans $\mathbb{R}$ définie sur U et $a = (a_1, a_2)$ un point de U . Si f est différentiable en a , la différentielle $df(a)$ est une forme linéaire sur $\mathbb{R}^2$, f admet une dérivée partielle en a par rapport à chacune des deux variables et, en notant (e_1, e_2) la base canonique de $\mathbb{R}^2$, on a :

$$\frac{\partial f}{\partial x}(a) = df(a) \cdot e_1 \quad \text{et} \quad \frac{\partial f}{\partial y}(a) = df(a) \cdot e_2$$

$$\text{et, pour tout vecteur } h = (h_1, h_2) \text{ de } \mathbb{R}^2, \quad df(a) \cdot h = h_1 \frac{\partial f}{\partial x}(a) + h_2 \frac{\partial f}{\partial y}(a).$$

Soit U un ouvert non vide de $\mathbb{R}^2$, f une fonction définie sur U et $a = (a_1, a_2)$ un point de U . Nous considérons les trois assertions :

- (1) Il existe un voisinage ouvert V de a tel que f admet, en tout point de V , une dérivée partielle par rapport à chacune des deux variables, et les deux fonctions dérivées partielles :

$$\frac{\partial f}{\partial x} \quad \text{et} \quad \frac{\partial f}{\partial y},$$

définies sur V , sont continues en a .

- (2) La fonction f est différentiable en a .
- (3) La fonction f admet en a une dérivée partielle par rapport à chacune des deux variables.

Alors (1) implique (2) et (2) implique (3). Nous démontrons que les réciproques sont fausses.

14.3. Fonction différentiable en $(0, 0)$ dont les dérivées partielles ne sont pas continues en $(0, 0)$.

Nous étudions l'application :

$$\left| \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R}^2 \\ (x, y) \longmapsto f(x, y) = \begin{cases} 0 & \text{si } (x, y) = (0, 0), \\ xy \sin \frac{1}{\sqrt{x^2 + y^2}} & \text{si } (x, y) \neq (0, 0). \end{cases} \end{array} \right.$$

Soit (x, y) un point de $\mathbb{R}^2$. Comme $x^2 \leq x^2 + y^2$, $|x| \leq \|(x, y)\|$ et de même $|y| \leq \|(x, y)\|$. Que (x, y) soit ou non différent de $(0, 0)$, on a :

$$|f(x, y) - f(0, 0)| \leq |x| |y| \leq \|(x, y)\| \times \|(x, y)\|.$$

Par suite :

$$f(x, y) - f(0, 0) = o_{(x, y) \rightarrow (0, 0)}(\|(x, y)\|)$$

donc f est différentiable en $(0, 0)$ et la différentielle de f en $(0, 0)$ est nulle.

Pour tout réel y , la fonction $x \mapsto \sqrt{x^2 + y^2}$ est dérivable en tout réel x tel que $x^2 + y^2 > 0$, ce qui équivaut à $(x, y) \neq (0, 0)$, et sa dérivée en un tel x est égale à $x/\sqrt{x^2 + y^2}$, donc la fonction $x \mapsto 1/\sqrt{x^2 + y^2}$ est dérivable en tout réel x tel que $(x, y) \neq 0$ et sa dérivée en un tel x vaut :

$$-\frac{\frac{d}{dx}(\sqrt{x^2 + y^2})}{(\sqrt{x^2 + y^2})^2} = -\frac{\frac{x}{\sqrt{x^2 + y^2}}}{x^2 + y^2} = -\frac{x}{(x^2 + y^2)^{3/2}}.$$

Il en résulte que f admet une dérivée partielle par rapport à la première variable en tout point de $\mathbb{R}^2 \setminus \{(0, 0)\}$ et que, pour tout point $z = (x, y)$ de $\mathbb{R}^2 \setminus \{(0, 0)\}$:

$$\frac{\partial f}{\partial x}(z) = y \sin \frac{1}{\sqrt{x^2 + y^2}} - \frac{x^2 y}{(x^2 + y^2)^{3/2}} \cos \frac{1}{\sqrt{x^2 + y^2}}.$$

De plus, pour tout point (x, y) de $\mathbb{R}^2$, $f(x, y) = f(y, x)$, donc f admet une dérivée partielle par rapport à la deuxième variable en tout point de $\mathbb{R}^2 \setminus \{(0, 0)\}$, dont on obtient la valeur en échangeant les lettres x et y dans l'expression de la dérivée partielle par rapport à la première variable. La différentielle de f en $(0, 0)$ étant nulle, on déduit du théorème 14.1 que f admet en $(0, 0)$ une dérivée partielle égale à 0 par rapport à chacune des deux variables.

Nous prouvons que les deux fonctions dérivées partielles sont discontinues en $(0, 0)$. Nous introduisons la suite $(z_n)_{n \geq 1}$ de points de $\mathbb{R}^2 \setminus \{(0, 0)\}$, de terme général $z_n = (x_n, y_n)$ où $x_n = y_n = 1/(n\pi\sqrt{2})$. On a, pour tout entier $n \geq 1$:

$$\frac{1}{x_n^2 + y_n^2} = n^2 \pi^2, \quad \frac{1}{\sqrt{x_n^2 + y_n^2}} = n\pi, \quad \frac{1}{(x_n^2 + y_n^2)^{3/2}} = n^3 \pi^3 \text{ et } x_n^2 y_n = \frac{1}{(2\sqrt{2})n^3 \pi^3}$$

donc :

$$\frac{\partial f}{\partial x}(z_n) = \frac{\partial f}{\partial y}(z_n) = \frac{(-1)^{n+1}}{2\sqrt{2}}$$

d'où l'on déduit que les suites $\left(\frac{\partial f}{\partial x}(z_n)\right)$ et $\left(\frac{\partial f}{\partial y}(z_n)\right)$ divergent. Or la suite (z_n) converge vers $(0, 0)$ et :

$$\frac{\partial f}{\partial x}(0, 0) = \frac{\partial f}{\partial y}(0, 0) = 0$$

donc les deux fonctions dérivées partielles sont discontinues en $(0, 0)$. $\square$

14.4. Fonction qui admet en $(0,0)$ des dérivées partielles par rapport aux deux variables sans être différentiable en $(0,0)$.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R}^2 \\ (x, y) \longmapsto f(x, y) = \begin{cases} 0 & \text{si } (x, y) = (0, 0), \\ \frac{xy}{\sqrt{x^2 + y^2}} & \text{si } (x, y) \neq (0, 0). \end{cases} \end{array} \right.$$

Pour tout point (x, y) de $\mathbb{R}^2$, $x^2 \leq x^2 + y^2$ donc $|x| \leq \|(x, y)\|$, et de même $|y| \leq \|(x, y)\|$, donc $|f(x, y)| \leq \|(x, y)\|$, que (x, y) soit ou non différent de $(0, 0)$. Par suite, f est continue en $(0, 0)$. Les fonctions partielles de f en $(0, 0)$ sont les applications $y \mapsto f(0, y) = 0$ et $x \mapsto f(x, 0) = 0$, identiquement nulles donc dérivables en 0, donc f admet en $(0, 0)$ des dérivées partielles par rapport aux deux variables et :

$$\frac{\partial f}{\partial x}(0, 0) = \frac{\partial f}{\partial y}(0, 0) = 0.$$

Supposons la fonction f différentiable en $(0, 0)$. On déduit du théorème 14.1 que la différentielle de f en $(0, 0)$ est nulle, donc que $f(x, y) = \underset{(x, y) \rightarrow (0, 0)}{o}(\|(x, y)\|)$. Il en résulte qu'en posant, pour tout $(x, y) \in \mathbb{R}^2$:

$$h(x, y) = \frac{|f(x, y)|}{\|(x, y)\|} = \frac{xy}{x^2 + y^2},$$

$h(x, y)$ admet pour limite 0 quand (x, y) tend vers $(0, 0)$. Or, pour tout nombre réel $x \neq 0$, $h(x, x) = 1/2$ et $h(x, 0) = 0$, et tout voisinage de $(0, 0)$ contient des (x, x) et des $(x, 0)$ où $x \neq 0$, donc h n'admet pas de limite en $(0, 0)$. En conclusion, la fonction f n'est pas différentiable en $(0, 0)$. $\square$

Dans les deux exemples qui suivent, nous allons répondre aux mêmes problèmes qu'aux exemples précédents 14.3 et 14.4, mais avec des conditions plus difficiles à remplir.

14.5. Fonction différentiable en $(0, 0)$ n'admettant des dérivées partielles sur aucun voisinage de $(0, 0)$ privé de $(0, 0)$.

Nous introduisons la fonction de Dirichlet¹ :

$$\left| \begin{array}{l} g : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} 1 & \text{si } x \text{ est rationnel,} \\ 0 & \text{si } x \text{ est irrationnel,} \end{cases} \end{array} \right.$$

dont nous rappelons qu'elle est discontinue en tout point de $\mathbb{R}$, et nous lui associons l'application :

$$\left| \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ (x, y) \longmapsto f(x, y) = x^2 g(x) + y^2 g(y). \end{array} \right.$$

On a, pour tout réel t , $|g(t)| \leq 1$, donc, pour tout $(x, y) \in \mathbb{R}^2$, $|f(x, y)| \leq x^2 + y^2$, d'où $|f(x, y)| \leq \|(x, y)\| \times \|(x, y)\|$. Par conséquent $f(x, y) = \underset{(x, y) \rightarrow (0, 0)}{o}(\|(x, y)\|)$.

1. Voir l'exemple 8.1, page 134.

Ainsi la fonction f est différentiable en $(0,0)$ et sa différentielle en $(0,0)$ est nulle. Soit $c = (a, b)$ un point de $\mathbb{R}^2$ différent de $(0,0)$. Si $a \neq 0$, l'application partielle $f_{1,c} : x \mapsto f_{1,c}(x) = f(x, b) = a^2 g(x) + b^2 g(b)$ est discontinue en a , donc f n'admet pas en c de dérivée partielle par rapport à la première variable. Le même raisonnement montre que si $b \neq 0$, f n'admet pas en c de dérivée partielle par rapport à la deuxième variable. On voit ainsi que si V est un voisinage ouvert de $(0,0)$, f n'admet, en tout point de $V \setminus \{(0,0)\}$, ni dérivée partielle par rapport à la première variable, ni dérivée partielle par rapport à la deuxième variable. $\square$

14.6. Fonction discontinue en $(0,0)$ admettant des dérivées partielles en $(0,0)$.

Nous introduisons l'application :

$$\left| \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ (x, y) \longmapsto f(x, y) = \begin{cases} 1 & \text{si } x \neq 0 \text{ et } y \neq 0, \\ 0 & \text{si } x = 0 \text{ ou } y = 0. \end{cases} \end{array} \right.$$

Les fonctions partielles de f en $(0,0)$ sont les applications $y \mapsto f(0, y) = 0$ et $x \mapsto f(x, 0) = 0$, identiquement nulles donc dérivables en 0 ; par conséquent f admet en $(0,0)$ des dérivées partielles par rapport aux deux variables. Pour tout réel $x \neq 0$, $f(x, x) = 1$; or $f(0,0) = 0$ et tout voisinage de $(0,0)$ contient des points (x, x) où $x \neq 0$, donc f est discontinue en $(0,0)$. $\square$

Comme nous l'avons remarqué pour la continuité, le fait de ne considérer que les dérivées partielles est arbitraire puisqu'il privilégie les axes de coordonnées, ce qui conduit à les généraliser de la manière suivante.

DÉFINITION 14.4. — Si f est une fonction définie sur un ouvert U de $\mathbb{R}^2$, $c = (a, b)$ un point de U et $u = (u_1, u_2)$ un vecteur non nul de $\mathbb{R}^2$, f est dérivable en c suivant u si la fonction $f_{(c,u)} : t \mapsto f_{(c,u)}(t) = f(c + tu)$ est dérivable en 0, et dans ce cas, la dérivée $f'_u(c)$ de f en c suivant u est la dérivée en 0 de $f_{(c,u)}$.

Soit f une fonction définie sur un ouvert U de $\mathbb{R}^2$ et $c = (a, b)$ un point de U . En notant (e_1, e_2) la base canonique de $\mathbb{R}^2$, f admet en c une dérivée partielle par rapport à la première (resp. la deuxième) variable si, et seulement si, f est dérivable en c suivant le vecteur e_1 (resp. e_2) et, dans ce cas :

$$\frac{\partial f}{\partial x}(c) = f'_{e_1}(c) \quad \left(\text{resp. } \frac{\partial f}{\partial y}(c) = f'_{e_2}(c) \right).$$

Enfin, si f est différentiable en c , alors, pour tout vecteur $u \neq (0,0)$, f est dérivable en c suivant u et $f'_u(c) = df(c) \cdot u$.

14.7. Fonction dérivable en $(0,0)$ suivant tout vecteur non nul mais discontinue en $(0,0)$.

Nous considérons l'application (voir le dessin de la page suivante) :

$$\left| \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ (x, y) \longmapsto f(x, y) = \begin{cases} 1 & \text{si } 0 < y < x^2, \\ 0 & \text{sinon.} \end{cases} \end{array} \right.$$

Remarquons que $f(0,0) = 0$.

On a, pour tout réel $x > 0$, $0 < x^2/2 < x^2$ donc $f(x, x^2/2) = 1$, ce qui montre que :

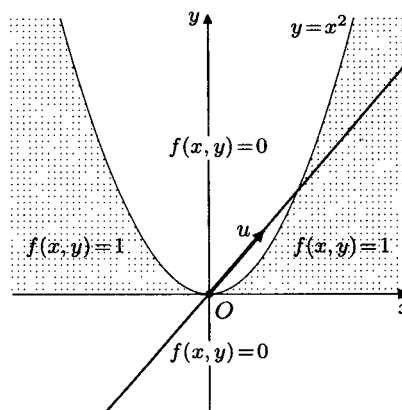
$$\lim_{x \rightarrow 0^+} f\left(x, \frac{x^2}{2}\right) = 1 \neq 0 = f(0, 0).$$

Or $(x, x^2/2)$ admet pour limite $(0, 0)$ dans $\mathbb{R}^2$ quand x tend vers 0 à droite, donc la fonction f est discontinue en $(0, 0)$.

Soit $u = (a, b)$ un vecteur de $\mathbb{R}^2$ différent de $(0, 0)$. Si $a = 0$, $f(tu) = f(0, tb) = 0$ pour tout nombre réel t et, si $a \neq 0$, $f(tu) = f(ta, tb) = 0$ pour tout nombre réel t si $b = 0$ et pour tout réel t tel que $|t| < |b|/a^2$ si $b \neq 0$. Par suite, dans tous les cas :

$$\lim_{t \rightarrow 0} \frac{f(tu)}{t} = 0,$$

ce qui montre que f est dérivable en $(0, 0)$ suivant le vecteur u et que sa dérivée en $(0, 0)$ suivant u est nulle. $\square$



14.8. Fonction continue en $(0, 0)$, dérivable en $(0, 0)$ suivant tout vecteur différent de zéro mais qui n'est pas différentiable en $(0, 0)$.

Nous étudions l'application :

$$\left| \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ (x, y) \longmapsto f(x, y) = \begin{cases} x & \text{si } y = 0, \\ 0 & \text{si } y \neq 0. \end{cases} \end{array} \right.$$

Remarquons que $f(0, 0) = 0$. Soit $u = (a, b)$ un vecteur de $\mathbb{R}^2$ différent de $(0, 0)$. Si $b = 0$, alors $a \neq 0$ donc $f(tu) = f(ta, 0) = ta$ pour tout réel $t \neq 0$, ce qui montre que f est dérivable en $(0, 0)$ suivant u et que $f'_u(0, 0) = a$. Si $b \neq 0$, on a $f(tu) = f(ta, tb) = 0$ pour tout réel $t \neq 0$, d'où l'on déduit que f est dérivable en $(0, 0)$ suivant u et que $f'_u(0, 0) = 0$.

Supposons f différentiable en $(0, 0)$. Nous notons (e_1, e_2) la base canonique de $\mathbb{R}^2$ et nous posons $u = e_1 + e_2 = (1, 1)$. On a $df(0, 0) \cdot u = f'_u(0, 0) = 0$ et, comme $e_1 = (1, 0)$ et $e_2 = (0, 1)$, on obtient :

$$df(0, 0) \cdot e_1 + df(0, 0) \cdot e_2 = f'_{e_1}(0, 0) + f'_{e_1}(0, 0) = 1 + 0 = 1,$$

ce qui contredit la linéarité de $df(0, 0)$. En conclusion, la fonction f n'est pas différentiable en $(0, 0)$. $\square$

DÉFINITION 14.5. — Soit U un ouvert de $\mathbb{R}^2$ et f une fonction définie sur U . Si a est un point de U , la fonction f est différentiable en a au sens de Gâteaux² si elle est dérivable en a suivant tout vecteur différent du vecteur zéro et si l'application Φ_a de $\mathbb{R}^2$ dans $\mathbb{R}$, définie par $\Phi_a((0, 0)) = 0$ et, pour $u \neq (0, 0)$, $\Phi_a(u) = f'_u(a)$ (dérivée de f en a suivant u), est linéaire.

2. Cette notion est introduite en 1907 par le mathématicien français René Gâteaux, qui meurt sur le front en 1914 à l'âge de vingt-cinq ans.

La fonction de l'exemple 14.8 n'est pas différentiable au sens de Gâteaux en $(0, 0)$. Si U est un ouvert de $\mathbb{R}^2$, f une fonction définie sur U et a un point de U , alors, si f est différentiable en a , f est dérivable en a suivant tout vecteur non nul et, pour tout vecteur $u \neq (0, 0)$, $f'_u(a) = df(a) \cdot u$, donc l'application Φ_a de la définition 14.5 est la différentielle $df(a)$ de f en a , qui est linéaire, ce qui montre que f est différentiable au sens de Gâteaux. Cependant la réciproque est fautive.

14.9. Fonction différentiable au sens de Gâteaux en $(0, 0)$ qui n'est pas différentiable en $(0, 0)$.

Nous identifions $\mathbb{R}^2$ au plan complexe $\mathbb{C}$ en identifiant le couple (x, y) au complexe $z = x + iy$. Rappelons que, pour tout complexe $z \neq 0$, il existe un couple (r, θ) et un seul tel que $z = re^{i\theta}$, $r \in \mathbb{R} \setminus \{0\}$ et $\theta \in]0, \pi]$. Nous considérons l'application :

$$f : \mathbb{C} \longrightarrow \mathbb{R} \\ z \longmapsto f(z) = \begin{cases} 0 & \text{si } z = 0, \\ \frac{r^2}{\theta} & \text{où } (r, \theta) \text{ est l'unique couple tel que} \\ & r \in \mathbb{R} \setminus \{0\}, \theta \in]0, \pi] \text{ et } z = re^{i\theta}. \end{cases}$$

Soit u un vecteur de $\mathbb{C}$ différent du vecteur zéro. On a $u = \rho e^{i\alpha}$ où $\rho \in \mathbb{R} \setminus \{0\}$ et $\alpha \in]0, \pi]$. Pour tout réel $t \neq 0$, on a $tu = (t\rho)e^{i\alpha}$, $t\rho \neq 0$ et $\alpha \in]0, \pi]$, donc :

$$\frac{1}{t} (f(tu) - f(0)) = \frac{1}{t} f(tu) = \frac{1}{t} \frac{t^2 \rho^2}{\alpha} = \frac{\rho^2}{\alpha} t.$$

Cette quantité admettant pour limite 0 quand t tend vers 0, f est dérivable en 0 suivant le vecteur u et la dérivée en 0 de f suivant u est nulle. Par conséquent f est différentiable au sens de Gâteaux en 0.

Supposons f différentiable en 0. Ce qui précède montre que la différentielle de f en 0 est nulle, donc on a $|f(z)| = |f(z) - f(0)| = o_{z \rightarrow 0}(|z|)$, d'où l'on déduit que :

$$(1) \quad \lim_{z \rightarrow 0} \frac{|f(z)|}{|z|} = 0.$$

Nous posons $v(\theta) = \theta e^{i\theta}$ pour tout $\theta \in]0, \pi]$. On a, pour tout $\theta \in]0, \pi]$:

$$\frac{|f(v(\theta))|}{|v(\theta)|} = \frac{\frac{\theta^2}{\theta}}{\theta} = 1,$$

donc $\lim_{\theta \rightarrow 0^+} \frac{|f(v(\theta))|}{|v(\theta)|} = 1$, ce qui, puisque $\lim_{\theta \rightarrow 0^+} |v(\theta)| = 0^+$, contredit (1).

Il en résulte que la fonction f n'est pas différentiable en 0. $\square$

DÉFINITION 14.6. — Soit E et F des espaces vectoriels normés, U un ouvert non vide de E et f une fonction de E dans F définie sur U . La fonction f est de classe $\mathcal{C}^1$ sur U si f est différentiable en tout point de U et si sa différentielle sur U , c'est-à-dire l'application $df : x \mapsto df(x)$, est une application continue de U dans l'espace vectoriel réel $\mathcal{L}_c(E, F)$ des applications linéaires continues de E dans F muni de la norme linéaire³.

3. La norme linéaire sur $\mathcal{L}_c(E, F)$ est l'application :

$$f \mapsto \|f\|_\ell = \sup_{x \in E \setminus \{0\}} \frac{\|f(x)\|}{\|x\|} = \sup_{\substack{x \in E \\ \|x\|=1}} \|f(x)\|$$

de $\mathcal{L}_c(E, F)$ dans $\mathbb{R}_+$; voir, dans le cas où $F = \mathbb{R}$, ce qui précède l'exemple 17.10, page 324.

THÉORÈME 14.2. — Une fonction f de $\mathbb{R}^2$ dans $\mathbb{R}$, définie sur un ouvert U de $\mathbb{R}^2$, est de classe $\mathcal{C}^1$ sur U si, et seulement si, les dérivées partielles de f par rapport aux deux variables existent en tout point de U et les deux fonctions dérivées partielles sont continues sur U .

14.10. Fonction différentiable sur $\mathbb{R}^2$ qui n'est pas de classe $\mathcal{C}^1$ sur $\mathbb{R}^2$.

Nous considérons les ouverts :

$$U_1 = \{(x, y) \mid y > x^2\},$$

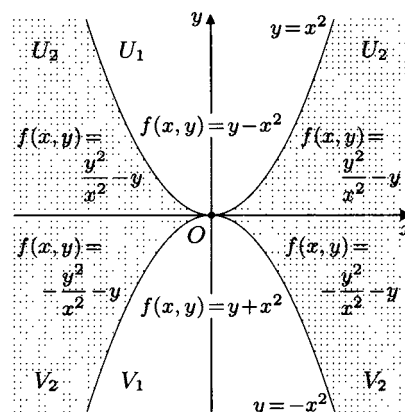
$$V_1 = \{(x, y) \mid y < -x^2\},$$

$$U_2 = \{(x, y) \mid 0 < y < x^2\},$$

$$V_2 = \{(x, y) \mid -x^2 < y < 0\}$$

de $\mathbb{R}^2$ et nous définissons l'application f de $\mathbb{R}^2$ dans $\mathbb{R}$ de la manière suivante :

$$f(x, y) = \begin{cases} y - x^2 & \text{si } y \geq x^2, \\ \frac{y^2}{x^2} - y & \text{si } 0 \leq y < x^2, \\ y + x^2 & \text{si } y < 0 \text{ et } y \leq -x^2, \\ -\frac{y^2}{x^2} - y & \text{si } -x^2 < y < 0 \end{cases}$$



Nous remarquons que $y - x^2$ et $(y^2/x^2) - y$ prennent la même valeur, à savoir 0, sur la courbe d'équation $y = x^2$, que $y + x^2$ et $-(y^2/x^2) - y$ ont la même valeur, également 0, sur la courbe d'équation $y = -x^2$ et enfin que $(y^2/x^2) - y$ et $-(y^2/x^2) - y$ prennent la même valeur, encore 0, sur la droite (Ox) . On en déduit facilement que la fonction f est continue sur $\mathbb{R}^2$.

La fonction f admet en tout point de l'ouvert U_1 (resp. V_1) des dérivées partielles par rapport aux deux variables et, pour tout point $z = (x, y)$ de U_1 (resp. de V_1) :

$$\frac{\partial f}{\partial x}(z) = -2x \text{ et } \frac{\partial f}{\partial y}(z) = 1 \quad \left(\text{resp. } \frac{\partial f}{\partial x}(z) = 2x \text{ et } \frac{\partial f}{\partial y}(z) = 1 \right).$$

De même f admet en tout point de l'ouvert U_2 (resp. V_2) des dérivées partielles par rapport aux deux variables et, pour tout point $z = (x, y)$ de U_2 (resp. de V_2) :

$$\frac{\partial f}{\partial x}(z) = -\frac{2y^2}{x^3} \text{ et } \frac{\partial f}{\partial y}(z) = \frac{2y}{x^2} - 1 \quad \left(\text{resp. } \frac{\partial f}{\partial x}(z) = \frac{2y^2}{x^3} \text{ et } \frac{\partial f}{\partial y}(z) = -\frac{2y}{x^2} - 1 \right).$$

En tout point $z = (x, y)$ tel que $y = x^2$ et $x \neq 0$ (resp. $y = -x^2$ et $x \neq 0$) en travaillant à droite, à gauche, au-dessus et au-dessous de z , on voit que f admet en z des dérivées partielles par rapport aux deux variables, données par les mêmes formules que pour $z \in U_1$ et $z \in V_1$ (resp. pour $z \in U_2$ et $z \in V_2$). Enfin, en tout point z de (Ox) différent de $(0, 0)$, f admet en z des dérivées partielles par rapport aux deux variables, données et, pour un tel point z :

$$\frac{\partial f}{\partial x}(z) = 0 \text{ et } \frac{\partial f}{\partial y}(z) = -1.$$

Finalement, f admet des dérivées partielles par rapport aux deux variables en tout point de l'ouvert $\Omega = \mathbb{R}^2 \setminus \{(0, 0)\}$ et on voit facilement que les deux fonctions dérivées partielles sont continues sur Ω . Il en résulte que f est de classe $\mathcal{C}^1$ sur Ω ; en particulier, f est différentiable en tout point de Ω .

Pour tout $x \neq 0$ (resp. $y \neq 0$), $f(x, 0) = 0$ (resp. $f(0, y) = y$), et $f(0, 0) = 0$, donc :

$$\frac{f(x, 0) - f(0, 0)}{x - 0} = 0 \quad \left(\text{resp.} \quad \frac{f(0, y) - f(0, 0)}{y - 0} = 1 \right).$$

Ainsi f admet en $(0, 0)$ des dérivées partielles par rapport aux deux variables et :

$$\frac{\partial f}{\partial x}(0, 0) = 0 \quad \text{et} \quad \frac{\partial f}{\partial y}(0, 0) = 1,$$

ce qui conduit à étudier $f(z) - f(0, 0) - y$ pour $z = (x, y)$ au voisinage de $(0, 0)$.

Soit $z = (x, y)$ un point de $\mathbb{R}^2$ différent de $(0, 0)$. On a :

$$|f(z) - f(0, 0) - y| = \begin{cases} 0 & \text{si } y = 0, \\ x^2 & \text{si } y > 0 \text{ et } y \geq x^2, \text{ ou } y < 0 \text{ et } y \leq -x^2, \\ \left| \frac{y^2}{x^2} - 2y \right| & \text{si } 0 < y < x^2, \\ \left| \frac{y^2}{x^2} + 2y \right| & \text{si } -x^2 < y < 0. \end{cases}$$

Si $0 < y < x^2$, alors $0 < 1/x^2 < 1/y$ donc :

$$\left| \frac{y^2}{x^2} - 2y \right| = \frac{|y^2 - 2x^2y|}{x^2} \leq \frac{|y^2 - 2x^2y|}{y} = |y - 2x^2| \leq |y| + 2x^2 = y + 2x^2 \leq 3x^2,$$

et on prouve de la même façon que, si $-x^2 < y < 0$, on a $\left| \frac{y^2}{x^2} + 2y \right| \leq 3x^2$.

Or $|x| \leq \|z\|$, donc, dans tous les cas, $|f(z) - f(0, 0) - y| \leq 3\|z\|^2 = (3\|z\|) \times \|z\|$.

Ainsi $f(x, y) - f(0, 0) - y = o_{(x,y) \rightarrow 0}(\|(x, y)\|)$, donc f est différentiable en $(0, 0)$.

Par suite f est différentiable sur $\mathbb{R}^2$. Pour tout réel $x \neq 0$, $\frac{\partial f}{\partial y}(x, 0) = -1$, donc :

$$\lim_{x \rightarrow 0} \frac{\partial f}{\partial y}(x, 0) = -1 \neq 1 = \frac{\partial f}{\partial y}(0, 0).$$

Or $(x, 0)$ admet pour limite $(0, 0)$ quand x tend vers 0, donc la fonction dérivée partielle en $(0, 0)$ par rapport à la deuxième variable est discontinue en $(0, 0)$, d'où l'on conclut que f n'est pas de classe $\mathcal{C}^1$ sur $\mathbb{R}^2$. $\square$

■ Théorèmes d'inversion

Pour les théorèmes d'inversion 14.3 et 14.4 qui suivent, E et F sont des espaces de Banach⁴, U est un ouvert de E et f est une application de U dans F .

Les exemples suivants 14.11 et 14.12 concernent le cas $E = F = \mathbb{R}^2$. Rappelons que si $f : (x, y) \mapsto f(x, y) = (u(x, y), v(x, y))$ est une fonction de $\mathbb{R}^2$ dans $\mathbb{R}^2$ définie et de classe $\mathcal{C}^1$ sur un ouvert U , la matrice jacobienne de f en un point $z = (x, y)$ de l'ouvert U est la matrice, carrée d'ordre 2 :

$$J_f(z) = \begin{pmatrix} \frac{\partial u}{\partial x}(z) & \frac{\partial u}{\partial y}(z) \\ \frac{\partial v}{\partial x}(z) & \frac{\partial v}{\partial y}(z) \end{pmatrix},$$

matrice représentative de la différentielle $df(z)$ dans la base canonique de $\mathbb{R}^2$.

4. Un espace de Banach est un espace vectoriel normé complet. Ces espaces sont introduits en 1920 par le mathématicien polonais Stefan Banach (1892-1945).

THÉORÈME 14.3. — Théorème d'inversion globale⁵.

Si la fonction f est de classe $\mathcal{C}^1$ sur U et injective, et si, pour tout point z de U , la différentielle $df(z)$ de f en z est inversible — ce qui signifie que c'est un isomorphisme d'espace vectoriel de E sur F —, alors $V = f(U)$ est un ouvert de F et f un difféomorphisme⁶ de classe $\mathcal{C}^1$ de U sur V .

Nous démontrons la nécessité de l'injectivité.

14.11. Fonction de classe $\mathcal{C}^1$ sur U de différentielle inversible en tout point de U , mais qui ne vérifie pas sur U le théorème d'inversion globale.

Nous posons $U = \mathbb{R}^2 \setminus \{(0, 0)\}$ et nous introduisons l'application :

$$\left| \begin{array}{l} f : U \longrightarrow \mathbb{R}^2 \\ z = (x, y) \longmapsto f(z) = (x^2 - y^2, 2xy). \end{array} \right.$$

Ses fonctions composantes étant de classe $\mathcal{C}^1$ sur U , f est de classe $\mathcal{C}^1$ sur U et, en tout point $z = (x, y)$ de U , la matrice jacobienne de f est :

$$J_f(z) = \begin{pmatrix} 2x & -2y \\ 2y & 2x \end{pmatrix}$$

dont le déterminant est $4(x^2 + y^2) > 0$, ce qui montre que $df(z)$ est inversible. Or, si x et y sont des réels différents de 0, (x, y) et $(-x, -y)$ appartiennent à l'ouvert U , $(x, y) \neq (-x, -y)$ et $f(x, y) = f(-x, -y)$, donc la fonction f n'est pas un difféomorphisme de U sur un ouvert de $\mathbb{R}^2$. $\square$

THÉORÈME 14.4. — Théorème d'inversion locale⁷.

Si f est de classe $\mathcal{C}^1$ sur U , si a est un point de U et si la différentielle $df(a)$ de f en a est inversible, il existe un voisinage V de a inclus dans U tel que $W = f(V)$ est un ouvert de F et la restriction $f|_V$ de f à V un difféomorphisme de classe $\mathcal{C}^1$ de V sur W .

L'exemple suivant montre que l'hypothèse «de classe $\mathcal{C}^1$ » est nécessaire : la différentiabilité ne suffit pas.

14.12. Fonction différentiable sur $\mathbb{R}^2$ qui ne vérifie pas en $a = (0, 0)$ le théorème d'inversion locale.

Nous utilisons l'application f de l'exemple 14.10 (pages 276 et 277), dont nous avons vu qu'elle est différentiable sur $\mathbb{R}^2$ mais qu'elle n'est pas de classe $\mathcal{C}^1$ sur $\mathbb{R}^2$, et nous considérons l'application :

$$\left| \begin{array}{l} g : \mathbb{R}^2 \longrightarrow \mathbb{R}^2 \\ z = (x, y) \longmapsto g(z) = (x, f(x, y)). \end{array} \right.$$

Ses fonctions composantes étant différentiables sur $\mathbb{R}^2$, g est différentiable sur $\mathbb{R}^2$.

5. Voir [GOUR], chapitre V, §3.1, théorème 1.

6. Un difféomorphisme de classe $\mathcal{C}^1$ d'un ouvert U de E sur un ouvert V de F est une bijection de U sur V de classe $\mathcal{C}^1$ sur U dont l'application réciproque est de classe $\mathcal{C}^1$ sur V .

7. Voir [GOUR], chapitre V, §3.1, corollaire 2.

On déduit des résultats de l'exemple 14.10 que la matrice jacobienne de g en $(0, 0)$ est :

$$J_g(0, 0) = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix},$$

ce qui montre que $dg(z) = \text{Id}_{\mathbb{R}^2}$, donc que $dg(z)$ est inversible. Or, pour tout nombre réel $x \neq 0$, $g(x, x^2) = g(x, -x^2)$, donc il n'existe aucun voisinage V de $(0, 0)$ tel que la restriction de g à V est injective. $\square$

■ Dérivées partielles secondes

Nous revenons à des fonctions de $\mathbb{R}^2$ dans $\mathbb{R}$ définies sur un ouvert U de $\mathbb{R}^2$.

Rappelons que si une telle fonction f admet en tout point de U une dérivée partielle par rapport à la deuxième (resp. la première) variable, la dérivée partielle seconde :

$$\frac{\partial^2 f}{\partial x \partial y}(z) \quad \left(\text{resp.} \quad \frac{\partial^2 f}{\partial y \partial x}(z) \right)$$

de f en un point $z = (x, y)$ de U est, si elle existe, la dérivée partielle en z par rapport à la première (resp. la deuxième) variable de la fonction :

$$\frac{\partial f}{\partial y} \quad \left(\text{resp.} \quad \frac{\partial f}{\partial x} \right).$$

THÉORÈME 14.5. — Théorème de Schwarz⁸.

Si f admet, en tout point $z = (x, y)$ de U , les deux dérivées partielles secondes :

$$\frac{\partial^2 f}{\partial x \partial y}(z) \quad \text{et} \quad \frac{\partial^2 f}{\partial y \partial x}(z),$$

si $c = (a, b)$ est un point de U et si les fonctions dérivées partielles secondes :

$$\frac{\partial^2 f}{\partial x \partial y} \quad \text{et} \quad \frac{\partial^2 f}{\partial y \partial x}$$

sont continues en c , alors :

$$\frac{\partial^2 f}{\partial x \partial y}(c) = \frac{\partial^2 f}{\partial y \partial x}(c).$$

14.13. Fonction f telle que $\frac{\partial^2 f}{\partial x \partial y}(0, 0)$ et $\frac{\partial^2 f}{\partial y \partial x}(0, 0)$ existent mais sont différentes⁹.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto f(z) = \begin{cases} 0 & \text{si } x = 0 \text{ ou } y = 0, \\ x^2 \arctan \frac{y}{x} - y^2 \arctan \frac{x}{y} & \text{si } x \neq 0 \text{ et } y \neq 0. \end{cases} \end{array} \right.$$

8. Du nom du mathématicien allemand Hermann Schwarz (1843-1921). Pour une démonstration, voir [HAIR], chapitre IV, §4, théorème 4.3 ou [GOUR], chapitre V, §1.2, théorème 2.

9. Après avoir étudié un exemple, Euler affirme en 1734 que l'on a en général l'égalité :

$$\frac{\partial^2 f}{\partial x \partial y}(c) = \frac{\partial^2 f}{\partial y \partial x}(c).$$

Il faut attendre 1873 pour que Hermann Schwarz donne ce contre-exemple.

En tout point de l'ouvert $\Omega = \{(x, y) \mid x \neq 0 \text{ et } y \neq 0\}$, f admet des dérivées partielles par rapport aux deux variables et, pour tout point $z = (x, y)$ de Ω :

$$\frac{\partial f}{\partial x}(z) = 2x \arctan \frac{y}{x} + x^2 \frac{\frac{-y}{x^2}}{1 + \frac{y^2}{x^2}} - y^2 \frac{\frac{1}{y}}{1 + \frac{x^2}{y^2}} = 2x \arctan \frac{y}{x} - y$$

et de même $\frac{\partial f}{\partial y}(z) = -2y \arctan \frac{x}{y} + x$.

Comme $f(t, 0) = 0$ pour tout réel t , f admet en tout point $z = (x, 0)$ de (Ox) une dérivée partielle par rapport à la première variable, égale à 0 ; de même, f admet en tout point $z = (0, y)$ de (Oy) une dérivée partielle par rapport à la deuxième variable, égale à 0. Si x est un nombre réel différent de 0, on a, pour tout réel $t \neq 0$:

$$\frac{f(x, t) - f(x, 0)}{t} = \frac{f(x, t)}{t} = x \frac{\arctan \frac{t}{x}}{t} - t \arctan \frac{x}{t},$$

et de plus $\lim_{(t \rightarrow 0)} (t/x) = 0$ et $\lim_{(t \rightarrow 0)} (x/t) = \pm\infty$; il en résulte que f admet en tout point $z = (x, 0)$ de (Ox) différent de $(0, 0)$ une dérivée partielle par rapport à la deuxième variable, égale à x . De même, f admet en tout point $z = (0, y)$ de (Oy) différent de $(0, 0)$ une dérivée partielle par rapport à la première variable, égale à $-y$. En conclusion, f admet sur $\mathbb{R}^2$ une dérivée partielle par rapport à chacune des deux variables et on a, pour tout point $z = (x, y)$ de $\mathbb{R}^2$:

$$\frac{\partial f}{\partial x}(z) = \begin{cases} -y & \text{si } x = 0, \\ 2x \arctan \frac{y}{x} - y & \text{si } x \neq 0 \end{cases} \quad \text{et} \quad \frac{\partial f}{\partial y}(z) = \begin{cases} x & \text{si } y = 0, \\ -2y \arctan \frac{x}{y} + x & \text{si } y \neq 0. \end{cases}$$

On voit facilement que les deux fonctions dérivées partielles sont continues sur $\mathbb{R}^2$, ce qui montre que la fonction f est de classe $\mathcal{C}^1$ sur $\mathbb{R}^2$.

Pour tout nombre réel t , $\frac{\partial f}{\partial y}(t, 0) = t$ et $\frac{\partial f}{\partial x}(0, t) = -t$, donc $\frac{\partial^2 f}{\partial x \partial y}(0, 0)$ et $\frac{\partial^2 f}{\partial y \partial x}(0, 0)$ existent et :

$$\frac{\partial^2 f}{\partial x \partial y}(0, 0) = 1 \neq -1 = \frac{\partial^2 f}{\partial y \partial x}(0, 0). \quad \square$$

14.14. Autre fonction f telle que $\frac{\partial^2 f}{\partial x \partial y}(0, 0)$ et $\frac{\partial^2 f}{\partial y \partial x}(0, 0)$ existent mais sont différentes.

Nous étudions l'application¹⁰ :

$$\left\{ \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto f(z) = \begin{cases} 0 & \text{si } x = y = 0, \\ \frac{xy(x^2 - y^2)}{x^2 + y^2} & \text{si } x \neq 0 \text{ ou } y \neq 0. \end{cases} \end{array} \right.$$

Un calcul simple et distrayant montre que f admet en tout point $\mathbb{R}^2$ différent de $(0, 0)$ une dérivée partielle par rapport à chacune des deux variables et que, pour

10. Peano montre en 1884 que toute fonction $f : (x, y) \mapsto f(x, y) = xyg(x, y)$ où g est une fonction bornée sur un voisinage de $(0, 0)$ et où :

$$\lim_{y \rightarrow 0} \lim_{x \rightarrow 0} g(x, y) = \alpha \neq \beta = \lim_{x \rightarrow 0} \lim_{y \rightarrow 0} g(x, y)$$

fournit un contre-exemple. La fonction de l'exemple 14.14 en est un cas particulier (voir [HAIR], chapitre IV, § 4.2).

tout point $z = (x, y) \neq (0, 0)$ de $\mathbb{R}^2$, on a :

$$\frac{\partial f}{\partial x}(z) = \frac{(x^2 + y^2)(3x^2y - y^3) + 2x^2y(y^2 - x^2)}{(x^2 + y^2)^2}$$

et :

$$\frac{\partial f}{\partial y}(z) = \frac{(x^2 + y^2)(x^3 - 3xy^2) + 2xy^2(y^2 - x^2)}{(x^2 + y^2)^2}.$$

On a donc, pour tout nombre réel $y \neq 0$: (1) $\frac{\partial f}{\partial x}(0, y) = -y$.

De même, pour tout nombre réel $x \neq 0$: (2) $\frac{\partial f}{\partial y}(x, 0) = x$.

De plus $f(t, 0) = f(0, t) = 0$ pour tout réel t , donc f admet en $(0, 0)$ une dérivée partielle égale à 0 par rapport à chacune des deux variables. Il en résulte que f admet en tout point de $\mathbb{R}^2$ une dérivée partielle par rapport à chacune des deux variables, que (1) est vraie pour tout réel y et que (2) est vraie pour tout réel x .

On en déduit comme dans l'exemple précédent 14.13 que $\frac{\partial^2 f}{\partial x \partial y}(0, 0)$ et $\frac{\partial^2 f}{\partial y \partial x}(0, 0)$ existent et que :

$$\frac{\partial^2 f}{\partial x \partial y}(0, 0) = 1 \neq -1 = \frac{\partial^2 f}{\partial y \partial x}(0, 0). \quad \square$$

■ Extremums

Nous étudions une fonction f de $\mathbb{R}^2$ dans $\mathbb{R}$ définie sur un ouvert U de $\mathbb{R}^2$. Comme pour les fonctions d'une variable réelle, f admet en un point a de U un maximum (resp. un minimum) relatif en a s'il existe un voisinage ouvert V de a inclus dans U tel que, pour tout point x de V , $f(x) \leq f(a)$ (resp. $f(x) \geq f(a)$), et f admet un extremum relatif en a si f admet un maximum ou un minimum relatif en a .

THÉORÈME 14.6. — Si a est un point de U , si f admet en a un extremum relatif et si f est différentiable en a , sa différentielle $df(a)$ en a est nulle.

Comme pour les fonctions d'une variable, la réciproque est fausse.

14.15. Fonction dont la différentielle en $(0, 0)$ est nulle mais qui n'admet pas d'extremum en $(0, 0)$.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto f(z) = x^2 - y^2. \end{array} \right.$$

La fonction f admet en tout point des dérivées partielles par rapport aux deux variables et on a, pour tout point $z = (x, y)$ de $\mathbb{R}^2$:

$$\frac{\partial f}{\partial x}(z) = 2x \text{ et } \frac{\partial f}{\partial y}(z) = -2y.$$

Les deux fonctions dérivées partielles étant clairement continues sur $\mathbb{R}^2$, la fonction f est de classe $\mathcal{C}^1$ sur $\mathbb{R}^2$. Les dérivées partielles sont nulles en $(0, 0)$, donc la différentielle de f est nulle en $(0, 0)$. On a $f(x, 0) = x^2 > 0$ pour tout réel $x \neq 0$ et

$f(0, y) = -y^2 < 0$ pour tout réel $y \neq 0$. Comme tout voisinage de $(0, 0)$ contient des points $(x, 0)$ tels que $x \neq 0$ et des points $(0, y)$ où $y \neq 0$, f n'admet pas d'extremum en $(0, 0)$. $\square$

Si f admet un minimum relatif en $a \in U$, sa restriction à chaque droite affine contenant a admet encore un minimum relatif ; cependant la réciproque est fausse.

14.16. Fonction qui n'admet pas de minimum relatif en $(0, 0)$ alors que sa restriction à chaque droite affine passant par $(0, 0)$ admet un minimum relatif strict ¹¹.

Nous introduisons l'application :

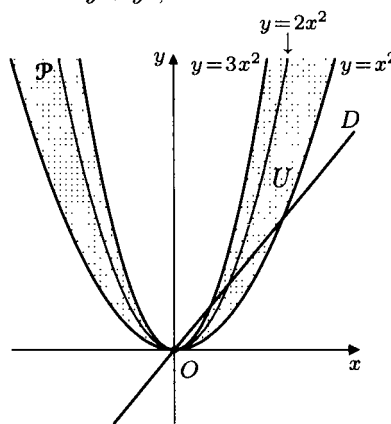
$$\begin{cases} f : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto f(z) = 3x^4 - 4x^2y + y^2, \end{cases}$$

l'ouvert $U = \{(x, y) \mid x^2 < y < 3x^2\}$ de $\mathbb{R}^2$ — en grisé sur le dessin ci-contre —, le fermé $G = \{(x, y) \mid x^2 \leq y \leq 3x^2\}$ de $\mathbb{R}^2$ et l'ouvert $V = \mathring{\mathbb{C}}_{\mathbb{R}^2} G$ de $\mathbb{R}^2$.

On constate que, pour tout point $z = (x, y)$ de $\mathbb{R}^2$, $f(x, y) = (y - x^2)(y - 3x^2)$, donc $f(z) < 0$ si, et seulement si, $z \in U$ et $f(z) > 0$ si, et seulement si, $z \in V$.

Soit D une droite affine de $\mathbb{R}^2$ contenant le point $(0, 0)$. Supposons d'abord que D soit différente de (Ox) et de (Oy) . Alors D est la droite d'équation $y = \lambda x$ où λ est un réel > 0 . Pour tout point $z = (x, \lambda x)$ de D , on a $f(z) = x^2(\lambda - x)(\lambda - 3x)$, donc $f(z) > 0$ pour tout $x \in]-\infty, \lambda/3[\setminus \{0\}$ et $f(z) = 0$ pour $x = 0$. Par conséquent la restriction de f à D admet un minimum relatif strict en 0. De plus, pour tout point $z = (x, 0)$ de la droite $D = (Ox)$, $f(z) = 3x^4$ et, pour tout point $z = (0, y)$ de la droite $D = (Oy)$, $f(z) = y^2$. En conclusion, la restriction de f à toute droite affine D contenant le point $(0, 0)$ admet un minimum relatif strict en 0.

Cependant, pour tout point $z = (x, y)$ différent de $(0, 0)$ appartenant à la parabole $\mathcal{P}$ d'équation $y = 2x^2$, on a $f(z) = -x^4 < 0$. Or $f(0, 0) = 0$ et tout voisinage de $(0, 0)$ contient des points de $\mathcal{P} \setminus \{(0, 0)\}$ et des points de l'ouvert V , donc la fonction f n'admet ni minimum relatif ni maximum relatif en $(0, 0)$. $\square$



■ Intégration

Pour l'intégrale double — c'est-à-dire l'intégrale sur une partie $\mathcal{P}$ de $\mathbb{R}^2$ d'une fonction de $\mathbb{R}^2$ dans $\mathbb{R}$ définie sur $\mathcal{P}$ —, on définit une notion de fonction intégrable sur un rectangle $\mathcal{A} = [a, b] \times [c, d]$. Le problème se pose de savoir si l'intégrale qui en découle est égale à la valeur obtenue en intégrant d'abord par rapport

¹¹. Cet exemple est dû, en 1884, au mathématicien italien Giuseppe Peano, âgé alors de vingt-cinq ans. Il contredit une affirmation du livre de Joseph Serret, qui faisait référence à l'époque.

à une variable et ensuite par rapport à l'autre. Pour l'intégrale de Riemann comme pour celle de Lebesgue, on dispose de théorèmes d'interversion de l'ordre des intégrations.

THÉORÈME 14.7. — Soit a, b, c et d des réels tels que $a < b$ et $c < d$ et f une fonction définie et intégrable au sens de Riemann sur le rectangle $\mathcal{A} = [a, b] \times [c, d]$. Si¹² la fonction $y \mapsto f(x, y)$ est, pour tout point x de $[a, b]$, intégrable sur le segment $[c, d]$, alors la fonction $x \mapsto \int_c^d f(x, y) dy$ est intégrable sur $[a, b]$ et :

$$\int_a^b dx \int_c^d f(x, y) dy = \iint_{\mathcal{A}} f(x, y) dx dy.$$

THÉORÈME 14.8. — **Théorème de Fubini**¹³.

Si a, b, c et d sont des réels tels que $a < b$ et $c < d$ et f une fonction définie et intégrable au sens de Lebesgue sur le rectangle $\mathcal{A} = [a, b] \times [c, d]$, alors la fonction $y \mapsto f(x, y)$ est, pour presque tout $x \in [a, b]$, intégrable sur le segment $[c, d]$, la fonction $x \mapsto f(x, y)$ est, pour presque tout $y \in [c, d]$, intégrable sur le segment $[a, b]$, la fonction $\varphi : x \mapsto \varphi(x) = \int_c^d f(x, y) dy$ est intégrable sur $[a, b]$ et la fonction $\psi : x \mapsto \psi(x) = \int_a^b f(x, y) dx$ est intégrable sur $[c, d]$ et :

$$(1) \quad \int_a^b dx \int_c^d f(x, y) dy = \int_c^d dy \int_a^b f(x, y) dx = \iint_{\mathcal{A}} f(x, y) dx dy,$$

et si de plus f est à valeurs réelles positives ou nulles et si l'une des deux premières intégrales existe, alors f est intégrable sur $\mathcal{A}$ et on a l'égalité (1).

14.17. Fonction f qui n'est pas intégrable au sens de Riemann sur un rectangle $[a, b] \times [c, d]$ alors que :

$$\int_a^b dx \int_c^d f(x, y) dy = \int_c^d dy \int_a^b f(x, y) dx.$$

Nous posons $\mathcal{A} = [-1, 1] \times [-1, 1]$ et nous considérons l'application :

$$\left| \begin{array}{l} f : \mathcal{A} \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto f(z) = \begin{cases} 0 & \text{si } z = (0, 0), \\ \frac{xy}{(x^2 + y^2)^2} & \text{si } z \neq (0, 0). \end{cases} \end{array} \right.$$

On a :

$$\int_{-1}^1 f(0, y) dy = \int_{-1}^1 0 dy = 0$$

et, pour $x \neq 0$:

$$\int_{-1}^1 f(x, y) dy = \int_{-1}^1 \frac{xy}{(x^2 + y^2)^2} dy = \frac{x}{2} \left[\frac{-1}{(x^2 + y^2)} \right]_{y=-1}^{y=1} = 0$$

donc :

$$\int_{-1}^1 dx \int_{-1}^1 f(x, y) dy = 0.$$

12. Ce théorème est établi en 1886 par le mathématicien autrichien Otto Stolz (1842-1905). Pour une démonstration, voir [HAIR], chapitre IV, §5, théorème 5.3.

13. Du nom du mathématicien italien Guido Fubini, qui démontre ce théorème en 1907.

Comme $f(x, y) = f(y, x)$ pour tout point (x, y) de $\mathcal{A}$, on a de même :

$$\int_{-1}^1 dy \int_{-1}^1 f(x, y) dx = 0.$$

Cependant, $f(x, x) = 1/(4x^2)$ pour tout point x de $]0, 1]$, donc f n'est pas bornée sur $\mathcal{A}$, ce qui montre que la fonction f n'est pas intégrable au sens de Riemann sur le rectangle $\mathcal{A} = [-1, 1] \times [-1, 1]$. $\square$

14.18. Fonction f bornée sur un rectangle $\mathcal{A} = [a, b] \times [c, d]$, qui n'est pas intégrable au sens de Riemann sur $\mathcal{A}$, alors que :

$$\int_a^b dx \int_c^d f(x, y) dy = \int_c^d dy \int_a^b f(x, y) dx.$$

Nous posons $D_0 = \{0, 1\}$ et nous notons, pour tout entier $k \geq 1$, D_k l'ensemble des nombres décimaux appartenant à $[0, 1]$ qui s'écrivent avec k décimales, la k -ième décimale n'étant pas nulle ; par exemple $0,47 \in D_2$ et $0,1789 \in D_4$. La famille $(D_k)_{k \in \mathbb{N}}$ est une partition de l'ensemble des nombres décimaux appartenant au segment $[0, 1]$. Nous posons :

$$\Delta = \bigcup_{k \in \mathbb{N}} (D_k \times D_k) \text{ et } \mathcal{A} = [0, 1] \times [0, 1]$$

et nous introduisons l'application :

$$\left| \begin{array}{l} f : \mathcal{A} \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto f(z) = \begin{cases} 1 & \text{si } z \in \Delta, \\ 0 & \text{si } z \in \mathcal{A} \setminus \Delta. \end{cases} \end{array} \right.$$

Soit x un point de $[0, 1]$. Si x n'est pas un nombre décimal, alors $f(x, y) = 0$ pour tout point y de $[0, 1]$, donc $\int_0^1 f(x, y) dy = 0$. Si x est un nombre décimal, il existe un entier naturel k tel que $x \in D_k$, donc $f(x, y) = 1$ seulement si $y \in D_k$, et comme D_k est fini, la fonction $y \mapsto f(x, y)$ est intégrable sur $[0, 1]$ et $\int_0^1 f(x, y) dy = 0$. On en déduit que :

$$\int_0^1 dx \int_0^1 f(x, y) dy = 0.$$

Comme $f(x, y) = f(y, x)$ pour tout point (x, y) de $\mathcal{A}$, on a de même :

$$\int_0^1 dy \int_0^1 f(x, y) dx = 0.$$

Soit a, b, c et d des réels tels que $0 \leq a < b \leq 1$ et $0 \leq c < d \leq 1$. Nous choisissons un entier naturel k tel que :

$$\frac{1}{10^k} < \min\left(\frac{b-a}{2}, \frac{d-c}{2}\right).$$

La propriété d'Archimède justifie l'existence de nombres décimaux $x \in]a, b[$ et $y \in]c, d[$ tels que $x \in D_{k+1}$ et $y \in D_{k+1}$, et alors $(x, y) \in \Delta$ donc $f(x, y) = 1$. Ainsi une fonction, constante sur $]a, b[\times]c, d[$ et majorant f sur ce rectangle, est au moins égale à 1. Clairement, une fonction constante sur $]a, b[\times]c, d[$ minorant f sur ce rectangle est au plus égale à 0. Comme ceci est vrai pour tout rectangle ouvert inclus dans $\mathcal{A} = [0, 1] \times [0, 1]$, la fonction f n'est pas intégrable au sens de Riemann sur $\mathcal{A}$. $\square$

14.19. Fonction f définie sur le rectangle $[0, 1] \times [0, 1]$ telle que :

$$\int_0^1 dx \int_0^1 f(x, y) dy \neq \int_0^1 dy \int_0^1 f(x, y) dx.$$

Nous posons $\mathcal{A} = [0, 1] \times [0, 1]$ et nous considérons l'application :

$$\left| \begin{array}{l} f : \mathcal{A} \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto f(z) = \begin{cases} 0 & \text{si } z = (0, 0), \\ \frac{x^2 - y^2}{(x^2 + y^2)^2} & \text{si } z \neq (0, 0). \end{cases} \end{array} \right.$$

On a, pour tout point x de $]0, 1]$:

$$\begin{aligned} \int_0^1 f(x, y) dy &= \int_0^1 \frac{x^2 - y^2}{(x^2 + y^2)^2} dy = \int_0^1 \frac{x^2 + y^2}{(x^2 + y^2)^2} dy - 2 \int_0^1 \frac{y^2}{(x^2 + y^2)^2} dy \\ &= \int_0^1 \frac{1}{x^2 + y^2} dy - \int_0^1 y \frac{2y}{(x^2 + y^2)^2} dy \\ &= \int_0^1 \frac{1}{x^2 + y^2} dy + \left[\frac{y}{x^2 + y^2} \right]_{y=0}^{y=1} - \int_0^1 \frac{1}{x^2 + y^2} dy = \frac{1}{x^2 + 1}. \end{aligned}$$

Remarquons que, pour $x = 0$, l'intégrale :

$$\int_0^1 f(0, y) dy = \int_0^1 \frac{-y^2}{y^4} dy = \int_0^1 \frac{-1}{y^2} dy$$

n'a pas de sens mais, comme $\{0\}$ est de mesure nulle, ceci n'est pas un problème pour le calcul qui suit — bien sûr f n'est pas intégrable sur $\mathcal{A}$. On a donc :

$$\int_0^1 dx \int_0^1 f(x, y) dy = \int_0^1 \frac{1}{x^2 + 1} dx = [\arctan x]_{x=0}^{x=1} = \frac{\pi}{4}.$$

Comme $f(x, y) = -f(y, x)$ pour tout point (x, y) de $\mathcal{A}$, on obtient :

$$\int_0^1 dy \int_0^1 f(x, y) dx = -\frac{\pi}{4} \neq \frac{\pi}{4} = \int_0^1 dx \int_0^1 f(x, y) dy. \quad \square$$

14.20. Autre fonction f définie sur $[0, 1] \times [0, 1]$ telle que :

$$\int_0^1 dx \int_0^1 f(x, y) dy \neq \int_0^1 dy \int_0^1 f(x, y) dx.$$

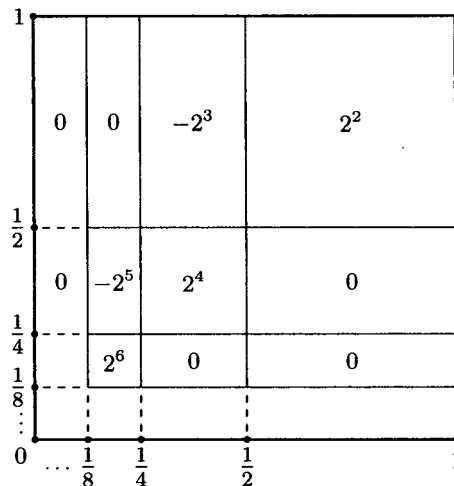
Nous posons $\mathcal{A} = [0, 1] \times [0, 1]$ et nous considérons l'application :

$$\left| \begin{array}{l} f : \mathcal{A} \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto f(z) = \begin{cases} 2^{2n} & \text{si } n \in \mathbb{N}^* \text{ et } z \in \left[\frac{1}{2^n}, \frac{1}{2^{n-1}} \right[\times \left[\frac{1}{2^n}, \frac{1}{2^{n-1}} \right[, \\ -2^{2n+1} & \text{si } n \in \mathbb{N}^* \text{ et } z \in \left[\frac{1}{2^{n+1}}, \frac{1}{2^n} \right[\times \left[\frac{1}{2^n}, \frac{1}{2^{n-1}} \right[, \\ 0 & \text{dans tous les autres cas.} \end{cases} \end{array} \right.$$

Une explication graphique donne une idée plus précise de cette fonction. Sur le schéma en haut de la page suivante, les nombres inscrits dans les rectangles indiquent la valeur constante qu'y prend la fonction f .

Soit x un point de l'intervalle $]0, 1/2[$. Il existe un entier $n \geq 2$ et un seul tel que $1/2^n \leq x < 1/2^{n-1}$. On a $f(x, y) = 2^{2n}$ pour tout point y de $[1/2^n, 1/2^{n-1}[$, $f(x, y) = -2^{2(n-1)+1}$ pour tout point y de $[1/2^{n-1}, 1/2^{n-2}[$ et $f(x, y) = 0$ pour les autres valeurs de y , donc :

$$\begin{aligned} \int_0^1 f(x, y) dy &= \int_{\frac{1}{2^n}}^{\frac{1}{2^{n-1}}} 2^{2n} dy - \int_{\frac{1}{2^{n-1}}}^{\frac{1}{2^{n-2}}} 2^{2(n-1)+1} dy \\ &= \frac{2^{2n}}{2^n} - \frac{2^{2n-1}}{2^{n-1}} = 0. \end{aligned}$$



Soit x un point de $]1/2, 1[$. On a $f(x, y) = 0$ si $y < 1/2$ et $f(x, y) = 4$ si $y \geq 1/2$, donc $\int_0^1 f(x, y) dy = \int_{1/2}^1 4 dy = 2$. Il en résulte que :

$$\int_0^1 dx \int_0^1 f(x, y) dy = \int_{\frac{1}{2}}^1 2 dx = 1.$$

Soit $y \in]0, 1[$. Il existe un entier $n \geq 1$ et un seul tel que $1/2^n \leq x < 1/2^{n-1}$. On a $f(x, y) = 2^{2n}$ pour tout point x de $[1/2^n, 1/2^{n-1}[$, $f(x, y) = -2^{2n+1}$ pour tout $x \in [1/2^{n+1}, 1/2^n[$ et $f(x, y) = 0$ pour les autres valeurs de x , donc :

$$\int_0^1 f(x, y) dx = - \int_{\frac{1}{2^{n+1}}}^{\frac{1}{2^n}} 2^{2n+1} dx + \int_{\frac{1}{2^n}}^{\frac{1}{2^{n-1}}} 2^{2n} dx = 0.$$

On en déduit que $\int_0^1 dy \int_0^1 f(x, y) dx = 0 \neq 1 = \int_0^1 dx \int_0^1 f(x, y) dy$. $\square$

14.21. Fonction f continue sur $[1, +\infty[\times [1, +\infty[$ telle que :

$$\int_1^{+\infty} dx \int_1^{+\infty} f(x, y) dy \neq \int_1^{+\infty} dy \int_1^{+\infty} f(x, y) dx.$$

Nous posons $\mathcal{A} = [1, +\infty[\times [1, +\infty[$ et nous considérons l'application :

$$\left| \begin{array}{l} f : \mathcal{A} \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto f(z) = \frac{x-y}{(x+y)^3}. \end{array} \right.$$

On a, pour tout réel $y \geq 1$, $\int_1^{+\infty} \frac{x-y}{(x+y)^3} dx = \left[\frac{-x}{(x+y)^2} \right]_{x=1}^{+\infty} = -\frac{1}{(1+y)^2}$, donc :

$$\int_1^{+\infty} dy \int_1^{+\infty} f(x, y) dx = - \int_1^{+\infty} \frac{1}{(1+y)^2} dy = \left[\frac{1}{(1+y)} \right]_{y=1}^{+\infty} = \frac{1}{2}.$$

En intervertissant les rôles de x et y , on obtient : $\int_1^{+\infty} dx \int_1^{+\infty} f(y, x) dy = \frac{1}{2}$.

Or $f(x, y) = -f(y, x)$ pour tout point $z = (x, y)$ de $\mathcal{A} = [1, +\infty[\times [1, +\infty[$, donc :

$$\int_1^{+\infty} dx \int_1^{+\infty} f(x, y) dy = -\frac{1}{2} \neq \frac{1}{2} = \int_1^{+\infty} dy \int_1^{+\infty} f(x, y) dx. \quad \square$$

14.22. Fonction f définie sur le rectangle $[0, 1] \times [0, 1]$ telle que, au sens de Riemann, l'intégrale :

$$\int_0^1 dy \int_0^1 f(x, y) dx$$

existe alors que l'intégrale $\int_0^1 dx \int_0^1 f(x, y) dy$ n'existe pas.

Nous introduisons le rectangle $\mathcal{A} = [0, 1] \times [0, 1]$ et l'application de l'exemple 11.6 (pages 211 et 212) :

$$\left| \begin{array}{l} g : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \begin{cases} 0 & \text{si } x \text{ est irrationnel ou égal à } 0, \\ \frac{1}{q} & \text{si } x \in \mathbb{Q} \cap]0, 1] \text{ et si } (p, q) \text{ est le représentant} \\ & \text{irréductible de } x \text{ de dénominateur positif,} \end{cases} \end{array} \right.$$

nous notons φ la fonction de Dirichlet, définie par $\varphi(x) = 1$ si $x \in \mathbb{Q}$ et $\varphi(x) = 0$ si $x \in \mathbb{R} \setminus \mathbb{Q}$ (exemple 8.1, page 134), et nous considérons l'application :

$$\left| \begin{array}{l} f : \mathcal{A} \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto f(z) = g(x)\varphi(y). \end{array} \right.$$

Soit $x \in [0, 1]$. Si x est nul ou irrationnel, $y \mapsto f(x, y)$ est l'application nulle de $[0, 1]$ dans $\mathbb{R}$. Si $x \in \mathbb{Q} \cap]0, 1]$ et si nous notons (p, q) le représentant irréductible de x tel que $q \in \mathbb{N}^*$, la fonction $y \mapsto f(x, y)$, produit par la constante $(1/q) \neq 0$ de la restriction de φ à $[0, 1]$, n'est pas intégrable au sens de Riemann sur $[0, 1]$ (voir l'exemple 11.1, page 208), ce qui montre que l'intégrale $\int_0^1 f(x, y) dy$ n'existe pas.

A fortiori, l'intégrale $\int_0^1 dx \int_0^1 f(x, y) dy$ n'existe pas.

Soit y un point de $[0, 1]$. Si y est irrationnel, $x \mapsto f(x, y)$ est l'application nulle de $[0, 1]$ dans $\mathbb{R}$ et, si y est rationnel, $x \mapsto f(x, y)$ est la restriction de g à $[0, 1]$, intégrable au sens de Riemann sur $[0, 1]$ et d'intégrale nulle (exemple 11.5).

Par conséquent l'intégrale $\int_0^1 dy \int_0^1 f(x, y) dx$ existe et vaut 0. $\square$

14.23. Fonction f intégrable sur $[0, 1] \times [0, 1]$ telle que $\int_0^1 f(x, y) dx$ n'existe pas pour certaines valeurs de y .

Nous posons $\mathcal{A} = [0, 1] \times [0, 1]$ et nous considérons l'application :

$$\left| \begin{array}{l} f : \mathcal{A} \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto f(z) = \begin{cases} 1 & \text{si } x \in \{0, 1\} \text{ et si } y \text{ est rationnel,} \\ 1 & \text{si } y \in \{0, 1\} \text{ et si } x \text{ est rationnel,} \\ 0 & \text{dans tous les autres cas.} \end{cases} \end{array} \right.$$

Pour $x = 0$ ou $x = 1$, la fonction $y \mapsto f(x, y)$ est la restriction à $[0, 1]$ de la fonction de Dirichlet (exemple 8.1, page 134) ; elle n'est donc pas intégrable au sens de Riemann sur $[0, 1]$ (exemple 11.1, page 208). De même, pour $y = 0$ ou $y = 1$, la fonction $x \mapsto f(x, y)$ n'est pas intégrable au sens de Riemann sur $[0, 1]$. Cependant f est constante égale à 0 sur l'intérieur $]0, 1[\times]0, 1[$ de $\mathcal{A}$, donc la fonction f est intégrable au sens de Riemann sur $\mathcal{A}$ et d'intégrale nulle. $\square$

Chapitre 15

Topologie générale

Le cadre des espaces métriques s'est avéré trop étroit pour traiter les problèmes de topologie ; certains espaces où des notions intuitives de continuité ou de limite semblaient utiles ne pouvait recevoir aucune distance pour les définir. En 1914, le mathématicien allemand Félix Hausdorff introduit les axiomes de la topologie générale en définissant la famille des ouverts, ou, ce qui est équivalent, la famille des voisinages de chaque point.

DÉFINITION 15.1. — Soit E un ensemble. Une topologie sur E est un ensemble $\mathcal{O}$ de parties de E vérifiant les trois axiomes suivants :

- (I) $\emptyset$ et E appartiennent à $\mathcal{O}$.
- (II) La réunion de toute famille d'éléments de $\mathcal{O}$ est un élément de $\mathcal{O}$.
- (III) Quels que soient les éléments O_1 et O_2 de $\mathcal{O}$, l'intersection $O_1 \cap O_2$ est un élément de $\mathcal{O}$.

DÉFINITION 15.2. — Un espace topologique est un couple $(E, \mathcal{O})$ où E est un ensemble et $\mathcal{O}$ une topologie sur E et, si $(E, \mathcal{O})$ est un espace topologique, les ouverts de $(E, \mathcal{O})$ sont les éléments de $\mathcal{O}$.

Si E est un espace topologique — comme de coutume, la donnée de la topologie $\mathcal{O}$ est sous-entendue —, les éléments de E sont appelés les points de E , une partie de E est ouverte si c'est un ouvert de E et on déduit de l'assertion (III) de la définition 15.1 que l'intersection de toute famille *finie* d'ouverts de E est ouverte.

DÉFINITION 15.3. — Si E est un espace topologique, une partie de E est un fermé de E si son complémentaire dans E est un ouvert de E .

Si E est un espace topologique, une partie de E est fermée si c'est un fermé de E . Les propriétés du passage au complémentaire donnent le théorème suivant.

THÉORÈME 15.1. — Soit E un espace topologique.

- (I') $\emptyset$ et E sont des fermés de E .
- (II') L'intersection de toute famille de fermés de E est fermée.
- (III') La réunion de toute famille finie de fermés de E est fermée.

DÉFINITION 15.4. — Si E est un espace topologique et a un point de E , un voisinage de a dans E est une partie V de E possédant la propriété suivante : il existe un ouvert inclus dans V et contenant a .

DÉFINITION 15.5. — La topologie grossière sur un ensemble E est celle dont les seuls ouverts sont $\emptyset$ et E — ce qui signifie que la topologie est $\mathcal{O} = \{\emptyset, E\}$.

DÉFINITION 15.6. — La topologie discrète sur un ensemble E est celle pour laquelle toute partie de E est un ouvert — ce qui signifie que la topologie est $\mathcal{O} = \mathcal{P}(E)$.

La topologie canonique d'un espace métrique E — voir le chapitre 16 — est la topologie sur E dont les ouverts sont les parties O de E possédant la propriété suivante : pour tout point x de O , il existe un réel $r > 0$ tel que la boule ouverte $B(x, r)$ de centre x et de rayon r est incluse dans O ; pour cette topologie, un voisinage d'un point a de E est une partie V de E contenant au moins une boule ouverte de centre a . La topologie canonique de $\mathbb{R}$ est la topologie de l'espace métrique $(\mathbb{R}, d)$, où d est la distance canonique $d : (x, y) \mapsto d(x, y) = |x - y|$ de $\mathbb{R}$.

■ Séparation

Pour l'étude de l'unicité de la limite d'une suite ou d'une fonction, on est amené à introduire des axiomes de séparation¹ appelés T_0 , T_1 et T_2 .

On associe à tout espace topologique E les assertions suivantes.

- T_0 : Quels que soient les points distincts x et y de E , il existe un voisinage V_1 de x tel que $y \notin V_1$ ou un voisinage V_2 de y tel que $x \notin V_2$.
- T_1 : Quels que soient les points distincts x et y de E , il existe un voisinage V_1 de x tel que $y \notin V_1$.
- T_2 : Quels que soient les points distincts x et y de E , il existe des voisinages V_1 de x et V_2 de y tels que $V_1 \cap V_2 = \emptyset$.

DÉFINITION 15.7. — Un espace topologique est T_0 (resp. T_1 , resp. T_2) si l'assertion T_0 (resp. T_1 , resp. T_2) est vraie pour cet espace topologique.

DÉFINITION 15.8. — Un espace topologique séparé est un espace topologique T_2 .

Si $\mathcal{D}$ est la topologie discrète sur un ensemble E , les singletons $\{x\}$ pour $x \in E$ sont des ouverts, donc l'espace topologique $(E, \mathcal{D})$ est séparé. Si l'on munit un espace métrique E de sa topologie canonique, l'espace topologique ainsi obtenu est séparé ; c'est le cas de $\mathbb{R}$ muni de sa topologie canonique.

Si E est un espace topologique, l'assertion T_2 entraîne la proposition T_1 qui elle-même implique l'assertion T_0 , mais toutes les réciproques sont en général fausses.

15.1. Espace topologique qui n'est pas T_0 .

Soit E un ensemble possédant au moins deux éléments distincts a et b , que l'on munit de la topologie grossière (voir la définition 15.5). L'ensemble E est le seul voisinage de a — en effet, si V est un voisinage de a , il existe un ouvert O tel que $a \in O$ et $O \subset V$, et comme O n'est pas vide, $O = E$, donc $V = E$ — et de même

1. Le mathématicien allemand Felix Hausdorff (1868-1942) définit en 1914 les voisinages d'un point, à l'aide desquels il donne des axiomes d'une topologie équivalents à ceux donnés ci-dessus par les ouverts. On lui doit l'axiome de séparation T_2 , appelé axiome de Hausdorff.

E est le seul voisinage de b . Comme de plus a et b sont des points distincts de l'ensemble E , l'assertion T_0 est fausse. $\square$

15.2. Espace topologique qui est T_0 mais n'est pas T_1 .

Si $(E, \leq)$ est un ensemble ordonné, une section finissante de E est une partie A de E telle que, pour tout élément x de A , tous les éléments de E supérieurs ou égaux à x appartiennent à A .

Dans l'ensemble ordonné $\mathbb{R}$, une section finissante est donc une partie A de $\mathbb{R}$ telle que, pour tout élément x de A , l'intervalle $[x, +\infty[$ est inclus dans A . On voit que les sections finissantes de $\mathbb{R}$ sont $\emptyset$ et les intervalles non majorés — c'est-à-dire $\emptyset, \mathbb{R}$, les $[a, +\infty[$ pour a parcourant $\mathbb{R}$ et les $]a, +\infty[$ pour a parcourant $\mathbb{R}$.

Nous notons $\mathcal{O}$ l'ensemble des sections finissantes de $\mathbb{R}$. Nous savons que l'ensemble vide et $\mathbb{R}$ appartiennent à $\mathcal{O}$.

Soit $(O_i)_{i \in I}$ une famille d'éléments de $\mathcal{O}$ et notons O la réunion de cette famille. Pour tout élément x de O , il existe un indice $i_0 \in I$ tel que $x \in O_{i_0}$ et, comme $[x, +\infty[\subset O_{i_0}$ et que $O_{i_0} \subset O$, $[x, +\infty[$ est inclus dans O . Par conséquent, O est une section finissante de $\mathbb{R}$.

Soit O_1 et O_2 des éléments de $\mathcal{O}$. Posons $O = O_1 \cap O_2$. Pour tout élément x de O , x appartient à O_1 et à O_2 , donc $[x, +\infty[$ est inclus dans O_1 et dans O_2 , ce qui montre que $[x, +\infty[$ est inclus dans O . Par suite, O est une section finissante de $\mathbb{R}$. Ceci achève de prouver que $\mathcal{O}$ est une topologie sur $\mathbb{R}$.

Nous munissons $\mathbb{R}$ de la topologie $\mathcal{O}$. Soit x et y des réels distincts. Quitte à échanger les lettres x et y , on suppose que $x < y$. Alors $V = [y, +\infty[$ est un voisinage de y — c'est un ouvert contenant y — et x n'appartient pas à V . Par contre, pour tout voisinage W de x , il existe une section finissante A contenant x et incluse dans W , donc $[x, +\infty[\subset A \subset W$, ce qui montre que y appartient à W . En conclusion, l'espace topologique $(\mathbb{R}, \mathcal{O})$ est T_0 mais il n'est pas T_1 . $\square$

15.3. Espace topologique qui est T_1 mais n'est pas T_2 .

Nous choisissons des nombres entiers relatifs distincts u et v n'appartenant pas à $\mathbb{N}$ — par exemple $u = -2$ et $v = -1$ —, nous posons $E = \{u, v\} \cup \mathbb{N}$ et nous notons $\mathcal{O}$ l'ensemble des parties O de E possédant l'une des trois propriétés suivantes :

- (1) O est une partie de $\mathbb{N}$.
- (2) u appartient à O et il existe un entier naturel p tel que tout entier naturel $n \geq p$ appartient à O .
- (3) v appartient à O et il existe un entier naturel q tel que tout entier naturel $n \geq q$ appartient à O .

L'ensemble vide vérifie (1) et E vérifie (2), donc $\emptyset$ et E appartiennent à $\mathcal{O}$.

Soit $(O_i)_{i \in I}$ une famille d'éléments de $\mathcal{O}$ et notons O la réunion de cette famille. Si l'un des O_i vérifie (2) ou (3), il en est clairement de même pour O ; sinon, $O_i \subset \mathbb{N}$ pour tout $i \in I$, donc O est inclus dans $\mathbb{N}$. Ainsi, dans tous les cas, la réunion O de la famille $(O_i)_{i \in I}$ appartient à $\mathcal{O}$.

Soit O_1 et O_2 des éléments de $\mathcal{O}$. Posons $O = O_1 \cap O_2$. Si O_1 ou O_2 vérifie (1), il est inclus dans $\mathbb{N}$, donc, puisque $O \subset O_1$ et $O \subset O_2$, O est inclus dans $\mathbb{N}$. Si O_1 vérifie (2) mais ne vérifie pas (3), et si O_2 vérifie (3) mais ne vérifie pas (2), alors $v \notin O_1$ et $u \notin O_2$, donc O est inclus dans $\mathbb{N}$. Si O_1 et O_2 vérifient tous

les deux la propriété (2) (resp. (3)), u (resp. v) appartient à O et il existe des entiers naturels p_1 et p_2 tels que tout entier $n \geq p_1$ appartient à O_1 et tout entier $n \geq p_2$ à O_2 donc, en posant $p = \text{Max}(p_1, p_2)$, tout entier naturel $n \geq p$ appartient à O , ce qui montre que O vérifie (2) (resp. (3)). Par suite, dans tous les cas, O appartient à $\mathcal{O}$. Ceci achève de prouver que $\mathcal{O}$ est une topologie sur E .

Nous munissons E de la topologie $\mathcal{O}$. Si a est un point de E , $E \setminus \{a\}$ vérifie (2) si $a \neq u$ et (3) si $a \neq v$, donc $E \setminus \{a\}$ est un ouvert. Si x et y sont des points distincts de E , $V = E \setminus \{x\}$ est un ouvert qui contient y , donc V est un voisinage de y ne contenant pas x . L'espace topologique $(E, \mathcal{O})$ est donc T_1 .

Soit U un voisinage de u et V un voisinage de v . Alors U contient un ouvert qui contient u , donc il existe un entier naturel p tel que tout entier $n \geq p$ appartient à U ; de même, il existe un entier naturel q tel que tout entier $n \geq q$ appartient à V . Par conséquent, l'entier naturel $m = \text{Max}(p, q)$ appartient à U et à V , donc $U \cap V$ n'est pas vide. En conclusion, l'espace topologique $(E, \mathcal{O})$ n'est pas T_2 . $\square$

THÉORÈME 15.2. — Si $(E, \mathcal{O})$ est un espace topologique, $\mathcal{R}$ une relation d'équivalence sur E , $E/\mathcal{R}$ l'ensemble quotient de E par $\mathcal{R}$ — c'est-à-dire l'ensemble des classes d'équivalence modulo $\mathcal{R}$ — et p la surjection canonique de E sur $E/\mathcal{R}$, alors l'ensemble des parties de $E/\mathcal{R}$ dont l'image réciproque par p est un ouvert de E , est une topologie sur $E/\mathcal{R}$, appelée la topologie quotient².

15.4. Espace topologique quotient non séparé d'un espace topologique séparé.

Muni de sa topologie canonique, $\mathbb{R}$ est un espace topologique séparé. Considérons la relation d'équivalence $\mathcal{R}$ définie sur $\mathbb{R}$ par : $x \mathcal{R} y$ si $x - y$ est rationnel. Nous notons p la surjection canonique de $\mathbb{R}$ sur $E = \mathbb{R}/\mathcal{R}$ et nous munissons E de la topologie quotient. La classe d'équivalence de 0 modulo $\mathcal{R}$ est $\mathbb{Q}$, donc $p^{-1}(\{p(0)\}) = \mathbb{Q}$.

Nous choisissons un réel irrationnel a (par exemple $a = e$). Comme $a - 0 = a \notin \mathbb{Q}$, les points $p(0)$ et $p(a)$ de E sont distincts. Soit V_1 et V_2 des voisinages respectifs de $p(0)$ et de $p(a)$ dans E . Il existe des ouverts O_1 et O_2 de $\mathbb{R}$ tels que $p(0) \in O_1$, $O_1 \subset V_1$, $p(a) \in O_2$ et $O_2 \subset V_2$. Alors $\Omega_1 = p^{-1}(O_1)$ est un ouvert de $\mathbb{R}$ qui contient $p^{-1}(\{p(0)\}) = \mathbb{Q}$ et $\Omega_2 = p^{-1}(O_2)$ un ouvert de $\mathbb{R}$ contenant a . Il existe un réel $\varepsilon > 0$ tel que la boule ouverte $]a - \varepsilon, a + \varepsilon[$ de centre a et de rayon ε est incluse dans Ω_2 . L'intervalle $]a - \varepsilon, a + \varepsilon[$ contient des rationnels, donc $\Omega_1 \cap \Omega_2$ n'est pas vide et, comme $\Omega_1 \cap \Omega_2 \subset p^{-1}(V_1) \cap p^{-1}(V_2) = p^{-1}(V_1 \cap V_2)$, $V_1 \cap V_2$ n'est pas vide. En conclusion, l'espace topologique quotient $E = \mathbb{R}/\mathcal{R}$ n'est pas séparé. $\square$

■ Suites convergentes, limites de suites, continuité

DÉFINITION 15.9. — Soit E un espace topologique et (x_n) une suite de points de E .

- a) Si $\ell \in E$, (x_n) converge vers ℓ dans E si, pour tout voisinage V de ℓ , il existe un entier naturel N tel que x_n appartient à V pour tout entier $n \geq N$.
- b) Un point ℓ de E est une limite de la suite (x_n) si (x_n) converge vers ℓ dans E .

2. Voir [SKAN], chapitre 3, § 2.

L'importance des suites provient du fait que dans $\mathbb{R}$, et plus généralement dans un espace métrique, les notions de convergence, de limite ou de continuité peuvent se ramener à des considérations sur des limites de suites. Ceci provient du fait que, dans un espace métrique, tout point possède une base dénombrable de voisinages³.

Soit E un espace topologique et A une partie de E . L'adhérence $\bar{A}$ de A dans E est le plus petit pour l'inclusion des fermés de E contenant A — c'est l'intersection de la famille des fermés de E contenant A — et un point a de E est adhérent à A s'il appartient à $\bar{A}$. Si ℓ est un point de E et si une suite (x_n) de points de A converge vers ℓ , ℓ est adhérent à A ; la réciproque est vraie dans un espace métrique : tout point adhérent à A est la limite dans E d'une suite convergente de points de A .

Dans un espace métrique, et plus généralement dans un espace topologique séparé, une suite convergente possède une seule limite. Ceci ne se généralise pas à un espace topologique quelconque.

15.5. Suite ayant plusieurs limites.

Dans le cas de la topologie grossière, toute suite est convergente et admet tout point pour limite. Aussi, nous utilisons une topologie contenant une infinité d'ouverts.

Nous notons $\mathcal{O}$ l'ensemble de parties de $\mathbb{N}$ constitué de l'ensemble vide et de tous les sous-ensembles de $\mathbb{N}$ dont le complémentaire dans $\mathbb{N}$ est une partie finie de $\mathbb{N}$. Comme $\mathbb{C}_{\mathbb{N}} \mathbb{N} = \emptyset$, $\mathbb{C}_{\mathbb{N}} \mathbb{N}$ est fini, donc $\mathbb{N}$ et l'ensemble vide appartiennent à $\mathcal{O}$. Soit $(O_i)_{i \in I}$ une famille d'éléments de $\mathcal{O}$ et O sa réunion. Si O_i est vide pour tout $i \in I$, O est vide donc O appartient à $\mathcal{O}$; sinon, en choisissant un indice $i_0 \in I$ tel que $O_{i_0} \neq \emptyset$, l'ensemble $\mathbb{C}_{\mathbb{N}} O_{i_0}$ est fini et $\mathbb{C}_{\mathbb{N}} O \subset \mathbb{C}_{\mathbb{N}} O_{i_0}$, donc $\mathbb{C}_{\mathbb{N}} O$ est une partie finie de $\mathbb{N}$, ce qui montre que O appartient à $\mathcal{O}$.

Soit O_1 et O_2 des éléments de $\mathcal{O}$ et $O = O_1 \cap O_2$. Si O_1 ou O_2 est vide, il en est de même de O ; sinon, $\mathbb{C}_{\mathbb{N}} O = \mathbb{C}_{\mathbb{N}} O_1 \cup \mathbb{C}_{\mathbb{N}} O_2$ est une partie finie de $\mathbb{N}$ comme réunion de deux parties finies de $\mathbb{N}$. Par suite, dans tous les cas, O appartient à $\mathcal{O}$.

Ceci achève de prouver que $\mathcal{O}$ est une topologie sur $\mathbb{N}$.

Nous munissons $\mathbb{N}$ de la topologie $\mathcal{O}$, nous choisissons une injection φ de $\mathbb{N}$ dans $\mathbb{N}$ — il en existe — et nous introduisons la suite $(u_n)_{n \geq 0}$ de terme général $u_n = \varphi(n)$. Soit a un entier naturel. Soit V un voisinage de a . Alors V contient un ouvert contenant a , donc il existe une partie A de V dont le complémentaire dans $\mathbb{N}$ est fini. L'application φ étant injective, $\varphi^{-1}(\mathbb{C}_{\mathbb{N}} A)$ est une partie finie de $\mathbb{N}$, donc majorée dans $\mathbb{N}$. Choisissons un majorant m de $\varphi^{-1}(\mathbb{C}_{\mathbb{N}} A)$ dans $\mathbb{N}$ et posons $N = m + 1$. Alors, pour tout entier $n \geq N$, $n \notin \varphi^{-1}(\mathbb{C}_{\mathbb{N}} A)$ donc u_n appartient à V . Il en résulte que la suite (u_n) converge vers a . Ainsi, la suite (u_n) converge dans $\mathbb{N}$ et admet tout élément de $\mathbb{N}$ pour limite. $\square$

La notion de suite convergente se définit à l'aide de la topologie, mais elle ne la caractérise pas comme le montre l'exemple suivant.

3. Une base de voisinages d'un point a est un ensemble (ou une famille) $\mathcal{B}$ de voisinages de a tel(telle) que tout voisinage de a contient un élément de l'ensemble (de la famille) $\mathcal{B}$; c'est une base dénombrable si l'ensemble $\mathcal{B}$ (ou l'ensemble des indices de la famille $\mathcal{B}$) est dénombrable.

15.6. Deux topologies différentes sur $\mathbb{R}$ ayant les mêmes suites convergentes.

Nous notons $\mathcal{O}$ l'ensemble de parties de $\mathbb{R}$ constitué de l'ensemble vide et de tous les sous-ensembles de $\mathbb{R}$ dont le complémentaire dans $\mathbb{R}$ est fini ou dénombrable. En raisonnant comme dans l'exemple précédent 15.5, on prouve que $\mathcal{O}$ est une topologie sur $\mathbb{R}$.

Soit $(u_n)_{n \geq 0}$ une suite de réels qui converge dans l'espace topologique $(\mathbb{R}, \mathcal{O})$. Nous notons ℓ une limite de (u_n) dans $(\mathbb{R}, \mathcal{O})$ et nous posons $U = \{u_n \mid n \in \mathbb{N}\}$ et $V = (\mathbb{R} \setminus U) \cup \{\ell\}$. Le complémentaire de V dans $\mathbb{R}$ est inclus dans U , ensemble fini ou dénombrable, donc V est un ouvert de $(\mathbb{R}, \mathcal{O})$ contenant ℓ . Par conséquent, V est un voisinage de ℓ , donc il existe un entier naturel N tel que u_n appartient à V pour tout entier $n \geq N$. Or V ne contient aucun terme de la suite autre que ℓ , donc $u_n = \ell$ pour tout entier $n \geq N$. La suite (u_n) est donc stationnaire⁴. La réciproque étant immédiate, les suites de réels qui convergent dans l'espace topologique $(\mathbb{R}, \mathcal{O})$ sont les suites stationnaires.

Nous notons $\mathcal{D}$ la topologie discrète sur $\mathbb{R}$ (définition 15.6, page 232). Si ℓ est un nombre réel, $\{\ell\}$ est un voisinage de ℓ dans l'espace topologique $(\mathbb{R}, \mathcal{D})$ donc, si une suite (u_n) de réels converge vers ℓ dans $(\mathbb{R}, \mathcal{D})$, il existe un entier naturel N tel que $u_n = \ell$ pour tout entier $n \geq N$: la suite (u_n) est stationnaire. Les suites de réels qui convergent dans $(\mathbb{R}, \mathcal{D})$ sont donc les suites stationnaires. Nous avons établi que les espaces topologiques $(\mathbb{R}, \mathcal{O})$ et $(\mathbb{R}, \mathcal{D})$ ont les mêmes suites convergentes, à savoir les suites stationnaires.

La réunion de deux parties de $\mathbb{R}$, finies ou dénombrables, étant elle-même finie ou dénombrable, n'est pas égale à $\mathbb{R}$, donc une partie A de $\mathbb{R}$, finie non vide ou dénombrable, n'est pas un ouvert de $(\mathbb{R}, \mathcal{O})$, alors que A est un ouvert de $(\mathbb{R}, \mathcal{D})$. Il en résulte que les topologies $\mathcal{O}$ et $\mathcal{D}$ diffèrent. $\square$

DÉFINITION 15.10. — Soit E et F des espaces topologiques et f une application de E dans F .

- a) Si a est un point de E , f est continue en a si l'image réciproque par f de tout voisinage de $f(a)$ dans F est un voisinage de a dans E .
- b) L'application f est continue sur E , ou f est une application continue de E dans F , si f est continue en tous les points de E .

Une application f d'un espace topologique E dans un espace topologique F est discontinue en un point a de E si elle n'est pas continue en a .

THÉORÈME 15.3. — Une application f d'un espace topologique E dans un espace topologique F est une application continue de E dans F si, et seulement si, l'image réciproque par f de tout ouvert de F est un ouvert de E .

Si E est un espace métrique, F un espace topologique, f une application de E dans F et a un point de E , f est continue en a si, et seulement si, pour toute suite (x_n) de points de E qui converge vers a , la suite $(f(x_n))$ converge vers $f(a)$. Ceci devient faux, en général, lorsque E est un espace topologique quelconque.

4. Une suite (u_n) à valeurs dans un ensemble E est stationnaire s'il existe un élément a de E et un entier naturel n_0 tels que $u_n = a$ pour tout entier $n \geq n_0$. Cette notion généralise celle de suite constante.

15.7. Application f discontinue en un point a telle que, pour toute suite (x_n) qui converge vers a , $(f(x_n))$ converge vers $f(a)$.

Nous considérons l'identité $f = \text{Id}_{\mathbb{R}}$ de $\mathbb{R}$ comme une application de $\mathbb{R}$, muni de la topologie $\mathcal{O}$ définie dans l'exemple précédent 15.6, dans $\mathbb{R}$ muni de la topologie discrète $\mathcal{D}$. Nous utilisons bien sûr les résultats établis dans l'exemple 15.6.

Soit a un réel. Soit (x_n) une suite de réels qui converge vers a dans $(\mathbb{R}, \mathcal{O})$. Elle est stationnaire, donc il existe $n_0 \in \mathbb{N}$ tel que $x_n = a$ pour tout entier $n \geq n_0$. Par conséquent, $f(x_n) = f(a)$ pour tout entier $n \geq n_0$, donc la suite $(f(x_n))$ converge vers $f(a)$ dans $(\mathbb{R}, \mathcal{D})$. Or $V = \{f(a)\}$ est un voisinage de $f(a)$ dans $(\mathbb{R}, \mathcal{D})$, alors que $f^{-1}(V) = \{a\}$ n'est pas un voisinage de a dans $(\mathbb{R}, \mathcal{O})$ — son complémentaire $\mathbb{R} \setminus \{a\}$ dans $\mathbb{R}$ n'est ni fini ni dénombrable —, donc f n'est pas continue en a . $\square$

THÉORÈME 15.4. — Si E est un espace topologique et F une partie de E , l'ensemble des intersections avec F de tous les ouverts de E est une topologie sur F .

DÉFINITION 15.11. — Si E est un espace topologique et F une partie de E , la topologie induite sur F par celle de E est la topologie sur F dont les ouverts sont les intersections avec F de tous les ouverts de E .

DÉFINITION 15.12. — Une application f d'un espace topologique E dans un espace topologique F est un homéomorphisme de E sur F si f est une bijection de E sur F et si les applications f et f^{-1} sont continues.

DÉFINITION 15.13. — L'espace topologique E est homéomorphe à l'espace topologique F s'il existe un homéomorphisme de E sur F .

La relation d'homéomorphisme sur la classe des espaces topologiques est réflexive, transitive et symétrique. Compte tenu de la symétrie, on dit que les espaces topologiques E et F sont homéomorphes pour exprimer que E (ou F) est homéomorphe à F (ou E). Les homéomorphismes sont les isomorphismes de la structure d'espace topologique. Intéressons nous à un problème qui apparaît pour chaque structure : Si E et F sont des espaces topologiques et s'il existe une bijection continue de E sur F et une bijection continue de F sur E , E et F sont-ils homéomorphes ? La réponse est en général négative.

15.8. Espaces topologiques E et F qui ne sont pas homéomorphes alors qu'il existe une bijection continue de E sur F et une bijection continue de F sur E .

Nous posons, pour tout entier naturel n , $A_n =]3n, 3n + 1[\cup \{3n + 2\}$, puis :

$$E = \bigcup_{n \in \mathbb{N}} A_n =]0, 1[\cup \{2\} \cup]3, 4[\cup \{5\} \cup]6, 7[\cup \{8\} \cup]9, 10[\cup \{11\} \cup \dots$$

et :

$$F = (E \cup \{1\}) \setminus \{2\} =]0, 1[\cup]3, 4[\cup \{5\} \cup]6, 7[\cup \{8\} \cup]9, 10[\cup \{11\} \cup \dots$$

et nous munissons les parties E et F de $\mathbb{R}$ des topologies induites par la topologie canonique de $\mathbb{R}$. Nous introduisons l'application :

$$\left| \begin{array}{l} f : E \longrightarrow F \\ x \longmapsto f(x) = \begin{cases} x & \text{si } x \neq 2, \\ 1 & \text{si } x = 2 \end{cases} \end{array} \right.$$

et aussi l'application :

$$\left| \begin{array}{l} g : F \longrightarrow E \\ x \longmapsto g(x) = \begin{cases} \frac{x}{2} & \text{si } x \leq 1, \\ \frac{x}{2} - 1 & \text{si } 3 < x < 4, \\ x - 3 & \text{sinon.} \end{cases} \end{array} \right.$$

On prouve facilement que f est une bijection de E sur F et g une bijection de F sur E . Comme $\{2\} = E \cap]2 - (1/2), 2 + (1/2)[$, $\{2\}$ est un voisinage de 2 dans E ; on en déduit que f est continue sur E . Clairement, g est continue sur F .

Supposons qu'il existe un homéomorphisme h de E sur F . Comme $h^{-1}(1)$ est un élément de E , il existe un entier naturel n tel que $h^{-1}(1)$ appartient à A_n .

Si $h^{-1}(1) = 3n + 2$, $V =]3n + 1, 3n + 3[\cap E$ est un voisinage de $h^{-1}(1)$ dans E , et comme h^{-1} est continue en 1, $(h^{-1})^{-1}(V)$ est un voisinage de 1 dans F , ce qui montre qu'il existe un réel $\varepsilon > 0$, que l'on peut choisir dans $]0, 1[$, tel que :

$$]1 - \varepsilon, 1 + \varepsilon[\cap F \subset (h^{-1})^{-1}(V) = h(V) = h(\{3n + 2\}) = \{h(3n + 2)\},$$

inclusion impossible puisque $]1 - \varepsilon, 1 + \varepsilon[\cap F =]1 - \varepsilon, 1[$ est infini.

Il en résulte que $h^{-1}(1)$ appartient à $]3n, 3n + 1[$. Notons φ la restriction de h à $]3n, 3n + 1[$. L'application φ est continue de $]3n, 3n + 1[$ dans F (topologies induites par celle de $\mathbb{R}$), donc φ est continue de l'intervalle $]3n, 3n + 1[$ dans $\mathbb{R}$, et φ est injective, donc φ est strictement croissante ou strictement décroissante⁵. Par conséquent, $h([3n, 3n + 1]) = \varphi([3n, 3n + 1])$ est un intervalle ouvert $]\alpha, \beta[$ où α et β sont des points de $\overline{\mathbb{R}}$ tels que $\alpha < \beta$. Le réel $h^{-1}(1)$ appartient à $]3n, 3n + 1[$, donc 1 appartient à $]\alpha, \beta[$, ce qui montre que $\alpha < 1 < \beta$; de plus, $]\alpha, \beta[$ est inclus dans $\text{Im}(h) = F$. En choisissant un réel c tel que $1 < c < \min(3, \beta)$, c appartient à F et $1 < c < 3$, en contradiction avec la définition de F . L'existence de h conduisant à une contradiction, les espaces topologiques E et F ne sont pas homéomorphes. $\square$

■ Connexité

La connexité d'un espace topologique veut exprimer qu'il est « d'un seul tenant ».

DÉFINITION 15.14. — Un espace topologique E est connexe s'il n'existe aucune partition de E en deux ouverts non vides⁶.

DÉFINITION 15.15. — Un connexe — ou une partie connexe — d'un espace topologique E est une partie de E qui, munie de la topologie induite par celle de E , est un espace topologique connexe.

Si E est un espace topologique, les singletons $\{x\}$ pour x parcourant E , ainsi que l'ensemble vide, sont des connexes de E .

THÉORÈME 15.5. — Les connexes de $\mathbb{R}$, muni de sa topologie canonique, sont les intervalles⁷. En particulier l'espace topologique $\mathbb{R}$ est connexe.

5. Voir [ARN2], théorème IV.2.3.

6. Voir [SKAN], chapitre 4, §3 ou [CHOQ], chapitre V.II, §13.

7. Le théorème des valeurs intermédiaires est un corollaire des théorèmes 15.5 et 15.7 (page suivante); voir par exemple [BOUA], chapitre 18, §18.2.

THÉORÈME 15.6. — Si C est un connexe d'un espace topologique E et $\bar{C}$ son adhérence, toute partie A de E telle que $C \subset A \subset \bar{C}$ est connexe.

THÉORÈME 15.7. — L'image directe par une application continue d'un espace topologique dans un autre d'un connexe de l'espace de départ est un connexe de l'espace d'arrivée.

Ce dernier résultat est faux, en général, pour l'image réciproque.

15.9. Image réciproque non connexe d'un connexe par une application continue.

Nous munissons $\mathbb{R}$ de sa topologie canonique et nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = x^2. \end{array} \right.$$

Elle est évidemment continue sur $\mathbb{R}$. L'intervalle $I = [1, +\infty[$ est connexe. L'image réciproque du connexe I par f est $A = \{x \in \mathbb{R} \mid x^2 \geq 1\} =]-\infty, -1] \cup [1, +\infty[$. Or -1 et 1 appartiennent à A alors que le segment $[-1, 1]$ n'est pas inclus dans A , donc A n'est pas un intervalle, ce qui montre que A n'est pas un connexe de $\mathbb{R}$. $\square$

THÉORÈME 15.8. — Si A et B sont des connexes d'un espace topologique dont l'intersection n'est pas vide, leur réunion $A \cup B$ est connexe. Ceci reste vrai si l'on suppose que A rencontre l'adhérence de B (ou B l'adhérence de A).

THÉORÈME 15.9. — La réunion d'une famille de connexes d'un espace topologique dont l'intersection n'est pas vide est connexe.

Le résultat du théorème 15.8 est en général faux pour des connexes A et B quelconques, même si l'adhérence de A rencontre celle de B .

15.10. Réunion non connexe de deux parties connexes dont les adhérences se rencontrent.

Nous munissons $\mathbb{R}$ de sa topologie canonique et nous posons $A =]-\infty, 0[$ et $B =]0, +\infty[$. Comme ce sont des intervalles, A et B sont des connexes de $\mathbb{R}$. L'adhérence de A est $]-\infty, 0]$ et celle de B est $[0, +\infty[$, donc l'adhérence de A rencontre celle de B . Or -1 et 1 appartiennent à $A \cup B$ alors que $[-1, 1]$ n'est pas inclus dans $A \cup B$, donc $A \cup B$ n'est pas un intervalle, ce qui montre que la réunion des connexes A et B de $\mathbb{R}$ n'est pas connexe. $\square$

Une intersection de connexes n'a aucune raison d'être connexe.

15.11. Intersection non connexe de deux parties connexes⁸.

Nous munissons les espaces $\mathbb{R}$ et $\mathbb{C}$ de leurs topologies canoniques et nous posons $U = \{z \in \mathbb{C} \mid |z| = 1\}$. L'application $f : t \mapsto f(t) = e^{it}$ de $\mathbb{R}$ dans $\mathbb{C}$ est continue, $\mathbb{R}$ est connexe et $U = f(\mathbb{R})$, donc U est un connexe de $\mathbb{C}$ (théorème 15.7). Enfin $\mathbb{R}$ est un connexe de $\mathbb{C}$, alors que $U \cap \mathbb{R} = \{-1, 1\}$ n'est pas connexe. $\square$

8. Un exemple géographique : le Tyrol autrichien n'est pas connexe bien que le Tyrol et l'Autriche soient tous les deux connexes.

Si a est un point d'un espace topologique E , la réunion de la famille des connexes de E contenant a est un connexe de E (théorème 15.9) et c'est évidemment le plus grand pour l'inclusion des connexes de E contenant a — le plus petit est $\{a\}$.

DÉFINITION 15.16. — La composante connexe d'un point a d'un espace topologique E est le plus grand pour l'inclusion des connexes de E contenant a .

Soit E un espace topologique. Nous définissons sur E la relation binaire $\mathcal{R}$ par : $x \mathcal{R} y$ s'il existe un connexe de E contenant x et y . Alors $\mathcal{R}$ est une relation d'équivalence sur E et les composantes connexes de E sont les classes d'équivalence modulo $\mathcal{R}$, donc la famille des composantes connexes de E est une partition de E . L'adhérence d'un connexe de E étant connexe (théorème 15.6, page 296), les composantes connexes de E sont fermées. En général, elles ne sont pas ouvertes ; c'est cependant vrai dans un espace topologique n'ayant qu'un nombre fini de composantes connexes et dans un espace topologique localement connexe⁹.

DÉFINITION 15.17. — Un espace topologique E est totalement discontinu si ses seuls connexes sont l'ensemble vide et les singletons $\{x\}$ pour x parcourant E , ce qui équivaut à l'assertion suivante : les composantes connexes de E sont les singletons $\{x\}$ pour x parcourant E .

Si $\mathcal{D}$ est la topologie discrète sur un ensemble E , l'espace topologique $(E, \mathcal{D})$ est totalement discontinu. La réciproque est en général fausse, et on peut même imposer à l'espace de n'avoir aucun point isolé¹⁰, alors que si l'on munit un ensemble E de la topologie discrète, tous les points de E sont isolés — en effet, pour tout point a de E , $\{a\}$ est un ouvert, donc un voisinage de a .

15.12. Espace topologique totalement discontinu sans point isolé.

Nous munissons l'ensemble $\mathbb{Q}$ des nombres rationnels de la topologie induite sur $\mathbb{Q}$ par la topologie canonique de $\mathbb{R}$. Soit A une partie de $\mathbb{Q}$ contenant au moins deux rationnels distincts r_1 et r_2 . Quitte à échanger les indices, nous supposons que $r_1 < r_2$. Nous munissons A de la topologie induite par celle de $\mathbb{Q}$, donc par celle de $\mathbb{R}$, nous choisissons un réel irrationnel α tel que $r_1 < \alpha < r_2$ et nous posons $O_1 = A \cap]-\infty, \alpha[$ et $O_2 = A \cap]\alpha, +\infty[$. Alors O_1 et O_2 sont des ouverts non vides de A , $O_1 \cap O_2$ est vide et $O_1 \cup O_2 = A$ — en effet, $\alpha \notin A$. Il en résulte que A n'est pas un connexe de $\mathbb{Q}$. Par conséquent les seuls connexes de $\mathbb{Q}$ sont l'ensemble vide et les singletons, donc $\mathbb{Q}$ est totalement discontinu.

Si $r \in \mathbb{Q}$, alors, pour tout réel $\varepsilon > 0$, l'intervalle ouvert $]r - \varepsilon, r + \varepsilon[$ contient des rationnels, donc $\{r\}$ n'est pas un voisinage de r dans $\mathbb{Q}$. Ainsi l'espace topologique $\mathbb{Q}$ n'a aucun point isolé et ses composantes connexes ne sont pas ouvertes. $\square$

Une topologie $\mathcal{O}_2$ sur un ensemble E est plus fine (resp. moins fine) que la topologie $\mathcal{O}_1$ sur E si tout ouvert de $(E, \mathcal{O}_1)$ (resp. de $(E, \mathcal{O}_2)$) est un ouvert de $(E, \mathcal{O}_2)$ (resp. de $(E, \mathcal{O}_1)$). Ainsi, plus une topologie sur un ensemble E est fine, plus elle a d'ouverts, donc il paraît intuitivement clair qu'elle a moins de connexes : il est plus facile pour une partie de E d'en trouver une partition en deux

9. Voir la définition 15.18 dans la suite de ce chapitre, page 298.

10. Un point a d'un espace topologique E est isolé si $\{a\}$ est un voisinage de a .

ouverts non vides. En utilisant le langage des ensembles ordonnés, nous voyons que l'application de l'ensemble des topologies sur E , muni de la relation de finesse (l'inclusion), dans l'ensemble $\mathcal{P}(\mathcal{P}(E))$, lui aussi muni de l'inclusion, qui à une topologie fait correspondre l'ensemble de ses parties connexes, est décroissante. Nous allons voir que cette application n'est pas strictement décroissante; en particulier, la donnée des parties connexes ne caractérise pas la topologie.

15.13. Deux topologies différentes sur $\mathbb{Q}$ ayant les mêmes parties connexes.

Nous munissons $\mathbb{Q}$ de la topologie discrète $\mathcal{D}$ et de la topologie $\mathcal{O}$ induite sur $\mathbb{Q}$ par la topologie canonique de $\mathbb{R}$. Dans l'espace topologique $(\mathbb{Q}, \mathcal{D})$, toutes les parties de $\mathbb{Q}$ sont ouvertes, donc toute partie de $\mathbb{Q}$ ayant au moins deux éléments distincts est la réunion disjointe de deux ouverts non vides; par conséquent les connexes de $(\mathbb{Q}, \mathcal{D})$ sont l'ensemble vide et les singletons. Nous avons vu dans l'exemple précédent 15.12 que les connexes de $(\mathbb{Q}, \mathcal{O})$ sont également l'ensemble vide et les singletons, donc les espaces topologiques $(\mathbb{Q}, \mathcal{D})$ et $(\mathbb{Q}, \mathcal{O})$ ont les mêmes connexes. Or, pour tout rationnel r , $\{r\}$ est un ouvert de $(\mathbb{Q}, \mathcal{D})$, mais ce n'est pas un voisinage de r dans $(\mathbb{Q}, \mathcal{O})$ (exemple 15.12) donc $\{r\}$ n'est pas un ouvert de $(\mathbb{Q}, \mathcal{O})$, ce qui montre que les topologies $\mathcal{D}$ et $\mathcal{O}$ sur $\mathbb{Q}$ sont différentes. $\square$

Si une topologie est moins fine qu'une autre, elle a plus de parties connexes. Cependant on peut trouver des espaces non séparés totalement discontinus.

15.14. Espace topologique non séparé et totalement discontinu.

Nous choisissons des relatifs distincts u et v n'appartenant pas à $\mathbb{N}$ — par exemple $u = -2$ et $v = -1$ —, nous posons $E = \{u, v\} \cup \mathbb{N}$ et nous munissons E de la topologie $\mathcal{O}$ définie dans l'exemple 15.3 (pages 290 et 291). Alors, pour tout $a \in E$, $E \setminus \{a\}$ est un ouvert de E , et l'espace topologique E n'est pas séparé.

Soit A une partie de E contenant au moins deux points distincts, que nous munissons de la topologie induite par celle de E . Il y a deux cas.

1^{er} cas : A contient au moins un entier naturel a . Alors $\{a\}$ est une partie de $\mathbb{N}$ donc un ouvert de E , ce qui montre que $O_1 = \{a\}$ est un ouvert non vide de A . De plus $O_2 = A \setminus \{a\} = A \cap (E \setminus \{a\})$ est un ouvert de A , non vide puisque $A \neq \{a\}$, $O_1 \cap O_2 = \emptyset$ et $O_1 \cup O_2 = A$.

2^e cas : A ne contient aucun entier naturel. Alors $A \subset \{u, v\}$, donc $A = \{u, v\}$. Les ensembles $U = \{u\} \cup \mathbb{N}$ et $V = \{v\} \cup \mathbb{N}$ sont des ouverts de E , donc $O_1 = A \cap U$ et $O_2 = A \cap V$ sont des ouverts non vides de A et, comme $O_1 = \{u\}$ et que $O_2 = \{v\}$, $O_1 \cap O_2 = \emptyset$ et $O_1 \cup O_2 = A$.

Il en résulte que, dans tous les cas, A n'est pas un connexe de E . En conclusion, les connexes de E sont l'ensemble vide et les singletons $\{x\}$ pour x parcourant E , donc l'espace topologique E est totalement discontinu. $\square$

DÉFINITION 15.18. — Un espace topologique E est localement connexe si, pour tout point a de E et tout voisinage V de a dans E , il existe un voisinage connexe de a inclus dans V .

Un espace topologique est donc localement connexe si tous ses points possèdent une base de voisinages connexes (voir la note n° 3, page 292).

Comme son nom l'indique, la connexité locale est une notion locale, alors que la connexité est une notion globale. Montrons qu'aucune des deux n'entraîne l'autre.

15.15. Espace topologique localement connexe qui n'est pas connexe.

Nous munissons $\mathbb{R}^* = \mathbb{R} \setminus \{0\}$ de la topologie induite par la topologie canonique de $\mathbb{R}$. On a $\mathbb{R}^* =]-\infty, 0[\cup]0, +\infty[$, $]-\infty, 0[$ et $]0, +\infty[$ sont des ouverts non vides de $\mathbb{R}^*$ et $]-\infty, 0[\cap]0, +\infty[= \emptyset$, donc l'espace topologique $\mathbb{R}^*$ n'est pas connexe. Or, si a est un point de $\mathbb{R}^*$, la suite $(V_n)_{n \geq 0}$ de parties de $\mathbb{R}^*$, de terme général :

$$V_n = \left] a - \frac{|a|}{n+2}, a + \frac{|a|}{n+2} \right[,$$

est clairement une base de voisinages connexes de a dans $\mathbb{R}^*$. $\square$

15.16. Espace topologique connexe qui n'est pas localement connexe.

Nous munissons $\mathbb{R}^2$ de la distance euclidienne, qui définit sa topologie canonique, et nous notons Δ « l'axe des ordonnées », donc $\Delta = \{0\} \times \mathbb{R}$, et, pour tout rationnel r , D_r la droite d'équation $y = r$, c'est-à-dire $D_r = \mathbb{R} \times \{r\}$. L'ensemble Δ est l'image du connexe $\mathbb{R}$ par l'application continue $y \mapsto (0, y)$ de $\mathbb{R}$ dans $\mathbb{R}^2$, donc Δ est un connexe de $\mathbb{R}^2$. De même, pour tout $r \in \mathbb{Q}$, D_r est l'image de $\mathbb{R}$ par l'application continue $x \mapsto (x, r)$ de $\mathbb{R}$ dans $\mathbb{R}^2$, donc D_r est un connexe de $\mathbb{R}^2$. Nous posons :

$$E = \left(\bigcup_{r \in \mathbb{Q}} D_r \right) \cup \Delta = \bigcup_{r \in \mathbb{Q}} (D_r \cup \Delta).$$

Pour tout $r \in \mathbb{Q}$, $(0, r) \in D_r \cap \Delta$, donc $D_r \cup \Delta$, réunion de deux connexes dont l'intersection n'est pas vide, est un connexe de $\mathbb{R}^2$. Ainsi E est la réunion d'une famille de connexes de $\mathbb{R}^2$, d'intersection non vide, donc E est un connexe de $\mathbb{R}^2$.

Nous considérons l'espace topologique connexe obtenu en munissant E de la topologie induite par la topologie canonique de $\mathbb{R}^2$.

Soit a et b des réels tels que $a \neq 0$. Nous posons $c = (a, b)$ et nous notons $\mathcal{B}_0$ la boule ouverte de centre c et de rayon $|a|$. Alors $V_0 = \mathcal{B}_0 \cap E$ est un voisinage de c dans E et, comme $\mathcal{B}_0 \cap \Delta = \emptyset$, tous les points de V_0 ont une ordonnée rationnelle.

Supposons qu'il existe un voisinage connexe V de c dans E inclus dans V_0 . Tous les points de V ont alors une ordonnée rationnelle. Il existe un nombre réel $\varepsilon > 0$ tel que, $\mathcal{B}$ étant la boule ouverte de centre c et de rayon ε , $\mathcal{B} \cap E$ est inclus dans V .

La projection :

$$\begin{array}{l} q : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto q(z) = y \end{array}$$

est une application continue de $\mathbb{R}^2$ dans $\mathbb{R}$ (topologies canoniques), donc sa restriction à E est une application continue de E dans $\mathbb{R}$. Comme V est un connexe de E , son image $W = q(V)$ par $q|_E$ est un connexe de $\mathbb{R}$; de plus, $W \subset \mathbb{Q}$, donc W est un connexe de $\mathbb{Q}$ muni de la topologie induite par la topologie canonique de $\mathbb{R}$. En choisissant des nombres rationnels r_1 et r_2 tels que $b - \varepsilon < r_1 < r_2 < b + \varepsilon$ — il en existe —, les points $c_1 = (a, r_1)$ et $c_2 = (a, r_2)$ appartiennent à V , donc $r_1 = q(c_1)$ et $r_2 = q(c_2)$ appartiennent à W , en contradiction avec les résultats de

l'exemple 15.12 (page 297) : une partie de $\mathbb{Q}$ contenant au moins deux rationnels distincts n'est pas connexe. Ainsi, aucun voisinage connexe de c n'est inclus dans le voisinage V_0 de c , donc l'espace topologique E n'est pas localement connexe. $\square$

DÉFINITION 15.19. — Un espace topologique E est connexe par arcs si, quels que soient les points a et b de E , il existe une application continue f de $[0, 1]$ dans E telle que $f(0) = a$ et $f(1) = b$.

La connexité par arcs de E exprime intuitivement que l'on peut relier deux points quelconques de E par un chemin continu inclus dans E . La connexité par arcs entraîne la connexité, mais la réciproque est en général fausse.

15.17. Espace topologique connexe qui n'est pas connexe par arcs.

Nous munissons $\mathbb{R}^2$ de la distance euclidienne, qui définit sa topologie canonique, nous considérons les suites $(D_n)_{n \geq 1}$ et $(E_n)_{n \geq 1}$ de parties de $\mathbb{R}^2$ de termes généraux :

$$D_n = \left\{ \frac{1}{n} \right\} \times \mathbb{R} \quad \left(\text{la droite d'équation } x = \frac{1}{n} \right) \quad \text{et} \quad E_n = \left[\frac{1}{n+1}, \frac{1}{n} \right] \times \{n\}$$

et nous introduisons la suite $(F_n)_{n \geq 1}$ de parties de $\mathbb{R}^2$ et les parties F et Δ de $\mathbb{R}^2$ définies par :

$$F_n = \bigcup_{1 \leq k \leq n} (E_k \cup D_k), \quad F = \bigcup_{n \in \mathbb{N}^*} F_n \quad \text{et} \quad \Delta = \{0\} \times \mathbb{R} \quad (\text{l'axe des ordonnées}).$$

En raisonnant comme dans l'exemple précédent 15.16, on prouve que, pour tout entier $n \geq 1$, E_n et D_n sont des connexes de $\mathbb{R}^2$, ce qui, puisque $(1/n, n)$ appartient à $E_n \cap D_n$, montre que $E_n \cup D_n$ est connexe.

Nous démontrons, par récurrence sur $n \in \mathbb{N}^*$, que, pour tout entier $n \geq 1$, F_n est un connexe de $\mathbb{R}^2$. Comme $F_1 = E_1 \cup D_1$, F_1 est connexe. Soit n un entier ≥ 1 et supposons que F_n soit connexe. On a $F_{n+1} = F_n \cup (E_{n+1} \cup D_{n+1})$, F_n et $E_{n+1} \cup D_{n+1}$ sont connexes et leur intersection n'est pas vide — elle contient $(1/(n+1), n)$ — donc F_{n+1} est un connexe de $\mathbb{R}^2$. Ceci achève le raisonnement par récurrence.

La suite $(F_n)_{n \geq 1}$ est une suite de connexes de E , croissante pour l'inclusion. Son intersection est F_1 , qui n'est pas vide, donc F est un connexe de $\mathbb{R}^2$.

Nous posons enfin $G = F \cup \Delta$ et nous munissons G de la topologie induite sur G par celle de $\mathbb{R}^2$. Nous démontrons que l'espace topologique G est connexe mais qu'il n'est pas connexe par arcs.

Pour tout nombre réel y , la suite $((1/n, y))$ de points de F admet pour limite $(0, y)$, donc le point $(0, y)$ appartient à l'adhérence $\bar{F}$ de F dans $\mathbb{R}^2$. Par conséquent, $\Delta \subset \bar{F}$, donc $F \subset G \subset \bar{F}$, et on déduit du théorème 15.6 (page 297) que G est un connexe de $\mathbb{R}^2$; l'espace topologique G est donc connexe. Les projections :

$$\left| \begin{array}{l} p : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto p(z) = x \end{array} \right. \quad \text{et} \quad \left| \begin{array}{l} q : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ z = (x, y) \longmapsto q(z) = y \end{array} \right.$$

sont des applications continues de $\mathbb{R}^2$ dans $\mathbb{R}$.

Les points $(0, 0)$ et $(1, 1)$ appartiennent au connexe G . Supposons l'existence d'une application continue f de $[0, 1]$ dans G telle que $f(0) = (0, 0)$ et $f(1) = (1, 1)$. Les composées $p \circ f$ et $q \circ f$ sont des applications continues de $[0, 1]$ dans $\mathbb{R}$.

La compacité¹¹ de $q \circ f([0, 1])$ nous assure l'existence d'un majorant a de l'ensemble des ordonnées des points $f(t)$ pour t parcourant le segment $[0, 1]$; l'ordonnée de $f(1)$ est 1, donc $a \geq 1$. De plus, $p \circ f([0, 1])$ est un intervalle contenant les points $p \circ f(0) = p((0, 0)) = 0$ et $p \circ f(1) = p((1, 1)) = 1$, donc $[0, 1]$ est inclus dans $p \circ f([0, 1])$. Nous choisissons un réel non entier strictement plus grand que $a + 1$ (il en existe) et nous notons b son inverse; alors :

$$0 < b < \frac{1}{a+1} \text{ et } \frac{1}{b} \notin \mathbb{N}^*.$$

En particulier, $b \in [0, 1]$; or $[0, 1] \subset p \circ f([0, 1])$, donc il existe $u \in [0, 1]$ tel que $b = p \circ f(u)$. Le point $f(u)$ appartient à G et $p(f(u)) = b > 0$, donc $f(u) \notin \Delta$, ce qui montre que $f(u) \in F$. Par suite, il existe un entier $p \geq 1$ tel que $f(u) \in F_p$, donc un entier $n \in [1, p]$ tel que $f(u) \in E_n \cup D_n$. Comme $p \circ f(u) = b \neq (1/n)$, $f(u)$ n'appartient pas à D_n , donc $f(u)$ appartient à E_n , d'où l'on déduit que :

$$\frac{1}{n+1} \leq b \leq \frac{1}{n} \text{ et } q \circ f(u) = n.$$

Il en résulte que $1/(n+1) < 1/(a+1)$, donc que $n > a$, en contradiction avec le fait que a majore l'ensemble des ordonnées des points $f(t)$ pour $t \in [0, 1]$. Finalement, il n'existe aucune application continue f de $[0, 1]$ dans G telle que $f(0) = (0, 0)$ et $f(1) = (1, 1)$: l'espace topologique G n'est pas connexe par arcs. $\square$

■ Compacité

Un recouvrement d'un ensemble E est une famille d'ensembles dont la réunion contient E .

DÉFINITION 15.20. — Si E est un espace topologique, un recouvrement ouvert d'une partie A de E est un recouvrement de A par une famille d'ouverts de E .

DÉFINITION 15.21. — Soit E un espace topologique. Une partie A de E possède la propriété de Borel-Lebesgue si, de tout recouvrement ouvert de A , on peut extraire un sous-recouvrement fini.

Toute partie finie d'un espace topologique possède la propriété de Borel-Lebesgue.

DÉFINITION 15.22. — Une partie K d'un espace topologique E est compacte, ou K est un compact de E , si la topologie induite sur K par celle de E est séparée et si K possède la propriété de Borel-Lebesgue. Un espace topologique est compact si c'est un compact de lui-même.

Il en résulte qu'un espace topologique E est compact s'il est séparé et s'il possède la propriété de Borel-Lebesgue.

11. Voir le paragraphe « Compacité » ci-dessous. La propriété utilisée ici est le fait que tout compact de $\mathbb{R}$ est borné.

Si A est une partie finie d'un espace topologique E et si la topologie induite sur A par celle de E est séparée, A est un compact de E ; en particulier, toute partie finie d'un espace topologique séparé est compacte. Si l'on munit $\mathbb{R}$ de sa distance canonique, et plus généralement $\mathbb{R}^p$ de la distance euclidienne, les compacts sont les parties bornées et fermées.

L'intérêt des compacts¹² est de pouvoir étendre des propriétés trivialement vérifiées par des applications définies sur un ensemble fini à des applications définies sur des espaces topologiques infinis, à condition bien sûr qu'elles soient continues.

15.18. Espace topologique non séparé dans lequel la propriété de Borel-Lebesgue est vérifiée.

Nous imposons évidemment à l'ensemble des ouverts d'être infini, sans quoi la propriété de Borel-Lebesgue est automatiquement vraie.

Nous notons $\mathcal{O}$ l'ensemble des parties de $\mathbb{N}$ constitué de l'ensemble vide et de tous les sous-ensembles de $\mathbb{N}$ dont le complémentaire dans $\mathbb{N}$ est une partie finie de $\mathbb{N}$. Alors $\mathcal{O}$ est une topologie sur $\mathbb{N}$ — voir l'exemple 15.5 (page 321).

Nous munissons $\mathbb{N}$ de la topologie $\mathcal{O}$. S'il existe des ouverts non vides O_1 et O_2 d'intersection vide, O_1 est inclus dans $\mathbb{N} \setminus O_2$, ce qui est absurde puisque O_1 est infini et $\mathbb{N} \setminus O_2$ fini. Par conséquent, $(\mathbb{N}, \mathcal{O})$ n'est pas un espace topologique séparé.

Soit $(O_i)_{i \in I}$ un recouvrement ouvert de $\mathbb{N}$. Les O_i ne sont pas tous vides. Nous choisissons un indice i_0 appartenant à I tel que $O_{i_0} \neq \emptyset$. Si $O_{i_0} = \mathbb{N}$, $(O_i)_{i \in I}$ est un sous-recouvrement fini du recouvrement $(O_i)_{i \in I}$ de $\mathbb{N}$. Il reste à examiner le cas où $O_{i_0} \neq \mathbb{N}$. Alors $\mathbb{N} \setminus O_{i_0}$ est une partie finie non vide de $\mathbb{N}$, donc on a $\mathbb{N} \setminus O_{i_0} = \{n_1, \dots, n_q\}$ où $q \in \mathbb{N}^*$. A chaque entier k appartenant à $[1, q]$, nous associons un indice $i_k \in I$ tel que $n_k \in O_{i_k}$. Comme $O_{i_0} \cup O_{i_1} \cup \dots \cup O_{i_q} = \mathbb{N}$, la famille $(O_i)_{i \in \{i_0, i_1, \dots, i_q\}}$ est un sous-recouvrement fini du recouvrement $(O_i)_{i \in I}$ de $\mathbb{N}$. L'espace topologique $(\mathbb{N}, \mathcal{O})$ vérifie donc la propriété de Borel-Lebesgue. $\square$

Par un raisonnement analogue, il est possible de montrer que dans l'espace topologique $(\mathbb{N}, \mathcal{O})$ de l'exemple précédent 15.18, toutes les parties de $\mathbb{N}$ vérifient la propriété de Borel-Lebesgue, alors que pourtant certaines d'entre-elles — les parties strictes infinies — ne sont pas fermées. Dans un espace topologique séparé, les parties compactes sont fermées. C'est sans doute cette propriété qui a amené à imposer la séparation dans la définition des espaces compacts.

15.19. Partie compacte qui n'est pas fermée.

Nous munissons $\mathbb{N}$ de la topologie grossière (définition 15.5, page 289) et nous posons $K = \{0\}$. La topologie induite sur K par celle de $\mathbb{N}$ est évidemment séparée et K est fini, donc K est compact. Or les fermés de $\mathbb{N}$ sont $\emptyset$ et $\mathbb{N}$, donc K n'est pas fermé. $\square$

12. Heine en 1872, puis Borel et Lebesgue peu après, utilisent cette propriété des parties fermées et bornées de $\mathbb{R}$ et des espaces $\mathbb{R}^p$. Cependant il faut attendre 1904, et le mathématicien français Maurice Fréchet, pour éprouver le besoin de définir la notion de compact ; il le fait dans les espaces métriques à l'aide des suites — voir le théorème 17.2, page 325.

On ne peut pas caractériser une topologie par la donnée de ses parties compactes, comme le montre l'exemple suivant.

15.20. Deux topologies différentes sur $\mathbb{R}$ ayant les mêmes parties compactes.

Nous munissons d'abord $\mathbb{R}$ de la topologie discrète $\mathcal{D}$. Elle est séparée, donc toute partie finie de $\mathbb{R}$ est compacte. Les singletons $\{x\}$ pour $x \in \mathbb{R}$ sont des ouverts donc, pour toute partie A de $\mathbb{R}$, la famille $(\{x\})_{x \in A}$ est un recouvrement ouvert de A , ce qui montre que si A est compacte, c'est une partie finie de $\mathbb{R}$. Les compacts de $(\mathbb{R}, \mathcal{D})$ sont donc les parties finies de $\mathbb{R}$.

Rappelons que tout ouvert Ω de la topologie canonique de $\mathbb{R}$ vérifie la propriété de Lindelöf¹³ : de tout recouvrement ouvert de Ω , on peut extraire un sous recouvrement fini ou dénombrable.

Nous notons $\mathcal{O}$ l'ensemble des $O = \Omega \setminus D$ où Ω est un ouvert de la topologie canonique et D une partie finie ou dénombrable de $\mathbb{R}$. Tout ouvert de la topologie canonique de $\mathbb{R}$ appartient à $\mathcal{O}$ (prendre $D = \emptyset$), donc $\emptyset$ et $\mathbb{R}$ appartiennent à $\mathcal{O}$.

Soit O_1 et O_2 des éléments de $\mathcal{O}$. Posons $O = O_1 \cap O_2$. Il existe des ouverts Ω_1 et Ω_2 de la topologie canonique de $\mathbb{R}$ et des parties finies ou dénombrables D_1 et D_2 de $\mathbb{R}$ tels que $O_1 = \Omega_1 \setminus D_1$ et $O_2 = \Omega_2 \setminus D_2$. Alors $O = (\Omega_1 \cap \Omega_2) \setminus (D_1 \cup D_2)$ et $D_1 \cup D_2$ est fini ou dénombrable, donc O appartient à $\mathcal{O}$.

Soit $(O_i)_{i \in I}$ une famille d'éléments de $\mathcal{O}$ et notons O sa réunion. Pour tout $i \in I$, il existe un ouvert Ω_i de la topologie canonique de $\mathbb{R}$ et une partie finie ou dénombrable D_i de $\mathbb{R}$ tels que $O_i = \Omega_i \setminus D_i$; notons Ω la réunion de la famille $(\Omega_i)_{i \in I}$. Alors $O \subset \Omega$ et, par la propriété de Lindelöf, il existe une partie finie ou dénombrable J de I telle que Ω est la réunion de la sous-famille $(\Omega_i)_{i \in J}$. On a :

$$\Omega \setminus \left(\bigcup_{i \in J} D_i \right) \subset \bigcup_{i \in J} (\Omega_i \setminus D_i) \subset \bigcup_{i \in I} (\Omega_i \setminus D_i) = \bigcup_{i \in I} O_i = O.$$

La réunion $\bigcup_{i \in J} D_i$ est une partie finie ou dénombrable de $\mathbb{R}$ et :

$$D = \Omega \setminus O \subset \bigcup_{i \in J} D_i,$$

donc la partie D de $\mathbb{R}$ est elle aussi finie ou dénombrable. Or $O = \Omega \setminus D$, donc O appartient à $\mathcal{O}$. Ceci achève de prouver que $\mathcal{O}$ est une topologie sur $\mathbb{R}$.

Nous munissons $\mathbb{R}$ de la topologie $\mathcal{O}$. La topologie $\mathcal{O}$ est plus fine que la topologie canonique de $\mathbb{R}$, donc $(\mathbb{R}, \mathcal{O})$ est un espace topologique séparé, d'où l'on déduit que toute partie finie de $\mathbb{R}$ est un compact de $(\mathbb{R}, \mathcal{O})$.

Soit A une partie infinie de $\mathbb{R}$. Pour toute partie finie B de A , $A \setminus B$ n'est pas vide, donc, en choisissant un point d_0 de A , un point d_1 de $A \setminus \{d_0\}$, un point d_2 de $A \setminus \{d_0, d_1\}$, un point d_3 de $A \setminus \{d_0, d_1, d_2\}$, etc., on construit une suite $(d_n)_{n \geq 0}$ de points de A deux à deux distincts. nous posons $D = \{d_n \mid n \in \mathbb{N}\}$ et, pour tout entier naturel p , $D_p = D \setminus \{d_p\}$ et $O_p = \mathbb{R} \setminus D_p$; D_p est dénombrable, donc O_p est un ouvert de $(\mathbb{R}, \mathcal{O})$. Comme $\mathbb{R}$ est la réunion de la famille $(O_n)_{n \in \mathbb{N}}$,

13. Du nom du mathématicien finlandais Lorentz Lindelöf (1827-1908).

cette famille est un recouvrement ouvert de A dans $(\mathbb{R}, \mathcal{O})$. Si $n \in \mathbb{N}$, $d_n \in A$ et, pour tout entier naturel $p \neq n$, $d_n \neq d_p$ donc $d_n \in D_p$, ce qui montre que :

$$d_n \notin \bigcup_{\substack{p \in \mathbb{N} \\ p \neq n}} O_p.$$

Par suite il n'existe aucune partie finie F de $\mathbb{N}$ telle que $A \subset \bigcup_{p \in F} O_p$, donc A n'est pas un compact de $(\mathbb{R}, \mathcal{O})$.

On conclut de cette étude que les compacts de $(\mathbb{R}, \mathcal{O})$ sont les parties finies de $\mathbb{R}$: les espaces topologiques $(\mathbb{R}, \mathcal{O})$ et $(\mathbb{R}, \mathcal{D})$ ont les mêmes parties compactes.

Il reste à montrer que $\mathcal{O}$ n'est pas la topologie discrète de $\mathbb{R}$. Remarquons qu'un ouvert non vide Ω de la topologie canonique de $\mathbb{R}$ n'est ni fini, ni dénombrable ; en effet, en choisissant un point a de Ω , il existe un réel $\varepsilon > 0$ tel que l'intervalle $I =]a - \varepsilon, a + \varepsilon[$ est inclus dans Ω , et I est équipotent à $\mathbb{R}$. Si O est un ouvert non vide de l'espace topologique $(\mathbb{R}, \mathcal{O})$, il existe un ouvert non vide Ω de la topologie canonique de $\mathbb{R}$ et une partie finie ou dénombrable D de $\mathbb{R}$ tels que $O = \Omega \setminus D$, donc si O était fini ou dénombrable, il en serait de même de Ω , ce qui est impossible. Ainsi, une partie finie non vide ou dénombrable de $\mathbb{R}$ (il y en a) n'est pas un ouvert de l'espace topologique $(\mathbb{R}, \mathcal{O})$, alors que c'est un ouvert de l'espace topologique $(\mathbb{R}, \mathcal{D})$: les topologies $\mathcal{O}$ et $\mathcal{D}$ sont différentes. $\square$

THÉORÈME 15.10. — L'image directe d'une partie compacte d'un espace topologique, par une application continue de cet espace dans un espace topologique séparé, est un compact de l'espace d'arrivée.

Ceci est faux, en général, si l'espace topologique d'arrivée n'est pas séparé.

15.21. Image directe non compacte d'une partie compacte par une application continue.

Nous notons $\mathcal{D}$ la topologie discrète et $\mathcal{G}$ la topologie grossière de $\mathbb{N}$, et f l'identité de $\mathbb{N}$. Toute partie de $\mathbb{N}$ étant un ouvert de $(\mathbb{N}, \mathcal{D})$, f est une application continue de $(\mathbb{N}, \mathcal{D})$ dans $(\mathbb{N}, \mathcal{G})$. Nous posons $A = \{0, 1\}$; c'est une partie finie de $\mathbb{N}$ et l'espace topologique $(\mathbb{N}, \mathcal{D})$ est séparé, donc A est un compact de $(\mathbb{N}, \mathcal{D})$. Les ouverts de $f(A) = A$, muni de la topologie induite par $\mathcal{G}$, sont $\emptyset$ et A , donc cette topologie induite n'est pas séparée : $f(A)$ n'est pas un compact de $(\mathbb{N}, \mathcal{G})$. $\square$

Le théorème 15.10 ne s'applique pas à l'image réciproque.

15.22. Image réciproque non compacte d'une partie compacte par une application continue.

Nous munissons $\mathbb{R}$ de sa topologie canonique. La fonction sinus est une application continue de $\mathbb{R}$ dans $\mathbb{R}$, $[-1, 1]$ est un compact de $\mathbb{R}$, l'image réciproque de $[-1, 1]$ par sinus est $\sin^{-1}([-1, 1]) = \{x \in \mathbb{R} \mid \sin x \in [-1, 1]\} = \mathbb{R}$ et l'espace topologique $\mathbb{R}$ n'est pas compact. $\square$

Chapitre 16

Espaces métriques

La notion de distance est introduite par Fréchet en 1905. Elle généralise à différents espaces la notion de proximité que l'on exprime dans $\mathbb{R}$ ou $\mathbb{C}$ par la valeur absolue ou le module et elle rapproche notre intuition de celle que nous avons dans l'espace qui nous environne. Elle permet surtout de traiter l'étude de la convergence uniforme indispensable en analyse fonctionnelle. On doit au mathématicien allemand Felix Hausdorff, en 1914, la définition actuelle d'espace métrique. Ceci l'inspirera pour définir la notion plus générale d'espace topologique.

DÉFINITION 16.1. — Une distance sur un ensemble E est une application :

$$\left| \begin{array}{l} d : E \times E \longrightarrow \mathbb{R}_+ \\ (x, y) \longmapsto d(x, y) \end{array} \right.$$

vérifiant les trois axiomes suivants, valables quels que soient les éléments x, y et z de E :

- (I) $d(x, y) = 0$ si, et seulement si, $x = y$.
- (II) $d(y, x) = d(x, y)$.
- (III) $d(x, z) \leq d(x, y) + d(y, z)$ [relation triangulaire].

DÉFINITION 16.2. — Un espace métrique¹ est un couple (E, d) où E est un ensemble non vide et d une distance sur E .

Le plus souvent, un espace métrique (E, d) est simplement noté E . Les éléments d'un espace métrique E sont aussi appelés les points de E .

DÉFINITION 16.3. — Si (E, d) est un espace métrique, a un point de E et r un nombre réel > 0 , la boule ouverte (resp. la boule fermée) de centre a et de rayon r est l'ensemble :

$$B(a, r) = \{x \in E \mid d(a, x) < r\} \quad (\text{resp. } B'(a, r) = \{x \in E \mid d(a, x) \leq r\}).$$

Si (E, d) est un espace métrique et a un point de E , alors, pour des réels r et r' tels que $0 < r < r'$, $B(a, r) \subset B'(a, r) \subset B(a, r')$.

1. Pour les généralités sur les espaces métriques, voir [CHOQ], chapitre V.III, qui reste la meilleure référence, ou [GOUR], chapitre I, §1.1.

DÉFINITION 16.4. — Si E est un espace métrique, une partie O de E est ouverte, ou O est un ouvert de E , si, pour tout point x de O , il existe au moins un réel $r > 0$ tel que la boule ouverte de centre x et de rayon r est incluse dans O .

On définit ainsi, sur un espace métrique E , une topologie séparée pour laquelle $\emptyset$, E et les boules ouvertes sont des ouverts, ce qui permet d'utiliser dans E toutes les notions introduites dans un espace topologique (chapitre 15) : fermé, intérieur et adhérence d'une partie, connexe, compact, suite convergente et limite d'une suite convergente (unique puisque la topologie est séparée) et continuité d'une application de E dans un espace topologique F ou d'une application d'un espace topologique F dans E . Dans un espace métrique, l'ensemble vide, l'espace tout entier et les boules fermées sont des fermés.

La distance canonique de $\mathbb{R}$ est l'application :

$$\left| \begin{array}{l} d : \mathbb{R} \times \mathbb{R} \longrightarrow \mathbb{R}_+ \\ (x, y) \longmapsto d(x, y) = |x - y| \end{array} \right.$$

et si a est un nombre réel et r un réel > 0 , la boule ouverte (resp. la boule fermée) de centre a et de rayon r de l'espace métrique $\mathbb{R}$ — c'est-à-dire de $(\mathbb{R}, d)$ — est l'intervalle ouvert (resp. fermé) $]a - r, a + r[$ (resp. $[a - r, a + r]$). Cette distance définit la topologie canonique de $\mathbb{R}$.

DÉFINITION 16.5. — Si d est une distance sur E et F une partie non vide de E , la distance induite sur F par la distance de E est la restriction de d à $F \times F$; l'ensemble F devient ainsi un espace métrique dont la topologie est la topologie induite sur F par celle de E .

■ Boules

Les espaces vectoriels normés (chapitre 17) sont des espaces métriques particuliers. Cependant, certaines propriétés valables dans un espace vectoriel normé sont mises en défaut dans un espace métrique quelconque.

16.1. Boule de rayon r strictement incluse dans une boule de rayon $r' < r$.

Considérons l'espace métrique obtenu en munissant l'ensemble $E = [0, +\infty[\cup \{-2\}$ de la distance induite par la distance canonique de $\mathbb{R}$.

La boule ouverte de centre -2 et de rayon 4 de E est $B(-2, 4) = \{-2\} \cup [0, 2[$, alors que la boule ouverte de centre 0 et de rayon 3 est $B(0, 3) = \{-2\} \cup [0, 3[$, d'où, dans l'espace métrique E , l'inclusion stricte $B(-2, 4) \subsetneq B(0, 3)$. $\square$

16.2. Espace métrique où les boules ouvertes sont fermées et où les boules fermées sont ouvertes.

Nous munissons $\mathbb{Z}$ de la distance d induite par la distance canonique de $\mathbb{R}$; $\mathbb{Z}$ devient ainsi un espace métrique et, quels que soient les entiers relatifs a et b , $d(a, b)$ est un entier naturel.

Si a est un entier relatif et r un réel > 0 n'appartenant pas à $\mathbb{N}^*$, il n'existe aucun

entier relatif x tel que $d(a, x) = r$, donc la boule ouverte $B(a, r)$ est égale à la boule fermée $B'(a, r)$. De plus, si $r \in \mathbb{N}^*$, alors, pour tout entier relatif x , $|x - a| < r$ entraîne $|x - a| \leq r - (1/2)$, donc $B(a, r) \subset B'(a, r - (1/2))$, d'où l'on déduit que $B(a, r) = B'(a, r - (1/2))$. En conclusion, toute boule ouverte est une boule fermée, donc un fermé de E . Par un raisonnement analogue, on montre que toute boule fermée est une boule ouverte, donc un ouvert de E . $\square$

Dans un espace vectoriel normé, l'adhérence d'une boule ouverte est la boule fermée correspondante. Dans un espace métrique quelconque, l'adhérence d'une boule ouverte est incluse dans la boule fermée correspondante, mais il n'y a pas toujours égalité.

16.3. Une boule ouverte dont l'adhérence n'est pas la boule fermée correspondante.

Nous choisissons un nombre premier p , nous notons, pour tout entier relatif n différent de zéro, $\mu(n)$ l'exposant de p dans la décomposition de n en produit de facteurs premiers, nous définissons l'application ν de $\mathbb{Q}^*$ dans $\mathbb{Z}$ de la manière suivante : si r est un nombre rationnel différent de zéro et si (n, m) est l'unique représentant irréductible de dénominateur positif de r , alors :

$$\nu(r) = \begin{cases} 0 & \text{si } p \text{ ne divise ni } n \text{ ni } m, \\ \mu(n) & \text{si } p \text{ divise } n, \\ -\mu(m) & \text{si } p \text{ divise } m, \end{cases}$$

et nous introduisons l'application :

$$\begin{cases} d : \mathbb{Q} \times \mathbb{Q} \longrightarrow \mathbb{R}_+ \\ (x, y) \longmapsto d(x, y) = \begin{cases} e^{-\nu(x-y)} & \text{si } x \neq y, \\ 0 & \text{si } x = y. \end{cases} \end{cases}$$

Nous démontrons que d est une distance sur $\mathbb{Q}$.

Soit x et y des nombres rationnels. Une exponentielle ne s'annule jamais, donc $d(x, y) = 0$ si, et seulement si, $x = y$. Par ailleurs, $\nu(-r) = \nu(r)$ pour tout rationnel $r \neq 0$, donc $d(x, y) = d(y, x)$.

Pour démontrer la relation triangulaire, nous prouvons d'abord que, si x et y sont des rationnels non nuls et si $x + y \neq 0$, $\nu(x + y) \geq \inf(\nu(x), \nu(y))$.

Soit $x, y \in \mathbb{Q}^*$ tels que $x + y \neq 0$; quitte à échanger les lettres x et y , nous supposons que $\nu(x) \leq \nu(y)$. Alors :

$$x = p^{\nu(x)} \times \frac{a}{b} \text{ et } y = p^{\nu(y)} \times \frac{c}{d}$$

où a et c sont des entiers relatifs différents de zéro et b et d des entiers naturels non nuls, a et b étant premiers entre eux et premiers avec p , et c et d premiers entre eux et premiers avec p . On a :

$$x + y = p^{\nu(x)} \left(\frac{a}{b} + p^{\nu(y)-\nu(x)} \times \frac{c}{d} \right) = p^{\nu(x)} \times \frac{n}{bd} \text{ où } n = ad + bcp^{\nu(y)-\nu(x)},$$

bd est premier avec p et n est un entier naturel, différent de zéro puisque $x + y \neq 0$, donc $\nu(x + y) = \nu(x) + \mu(n) \geq \nu(x) = \inf(\nu(x), \nu(y))$.

Soit maintenant des rationnels x, y et z . S'ils ne sont pas deux à deux distincts, on a par exemple $x = y$, donc $d(x, z) = 0 + d(y, z) = d(x, y) + d(y, z) \leq d(x, y) + d(y, z)$, et la preuve est aussi simple dans les cas $x = z$ et $y = z$. Nous supposons donc

que x , y et z sont deux à deux distincts. Comme $x - z = (x - y) + (y - z)$ et que les nombres rationnels $x - y$ et $y - z$ et leur somme $x - z$ sont différents de zéro, $\nu(x - z) \geq \inf(\nu(x - y), \nu(y - z))$ donc $-\nu(x - z) \leq \sup(-\nu(x - y), -\nu(y - z))$. La fonction exponentielle étant croissante, on obtient :

$$\begin{aligned} d(x, z) &= e^{-\nu(x-z)} \leq e^{\sup(-\nu(x-y), -\nu(y-z))} \\ &\leq \sup(e^{-\nu(x-y)}, e^{-\nu(y-z)}) = \sup(d(x, y), d(y, z)) \leq d(x, y) + d(y, z) \end{aligned}$$

puisque les réels $d(x, y)$ et $d(y, z)$ sont positifs.

En conclusion, d est une distance² sur $\mathbb{Q}$.

Remarquons maintenant que $d(x, y)$ ne prend aucune valeur entre $1/e$ et 1 ; par conséquent la boule ouverte $\mathcal{B}$ de centre 0 et de rayon 1 n'est autre que la boule fermée de même centre et de rayon $1/2$, donc $\mathcal{B}$ est à la fois ouverte et fermée. De même, $d(x, y)$ ne prend aucune valeur entre 1 et e , donc la boule fermée $\mathcal{B}'$ de centre 0 et de rayon 1 est aussi la boule ouverte de même centre et de rayon 2 , ce qui montre que $\mathcal{B}'$ est à la fois ouverte et fermée. On a $\mathcal{B} = \{r \in \mathbb{Q} \mid \nu(r) > 0\}$ et $\mathcal{B}' = \{r \in \mathbb{Q} \mid \nu(r) \geq 0\}$. En choisissant des nombres premiers q_1 et q_2 distincts et différents de p et en posant $r = q_1/q_2$, le rationnel r appartient à la boule $\mathcal{B}'$ mais pas à $\mathcal{B}$, donc la boule $\mathcal{B}$ est strictement incluse dans $\mathcal{B}'$. De plus, $\mathcal{B}$ est un fermé, donc l'adhérence de $\mathcal{B}$ est égale à $\mathcal{B}$, et comme $\mathcal{B}'$ est la boule fermée de centre 0 et de rayon 1 , l'adhérence de $B(0, 1)$ n'est pas égale à $\mathcal{B}'(0, 1)$. $\square$

16.4. Deux boules ouvertes disjointes de même rayon r incluses dans une boule fermée de rayon r .

Nous reprenons la distance d sur $\mathbb{Q}$ définie dans l'exemple précédent 16.3. On a alors $d(x, z) \leq \sup(d(x, y), d(y, z))$ quels que soient les rationnels x , y et z .

Comme $\nu(1) = 0$, $d(0, 1) = e^{\nu(1)} = e^0 = 1$. Posons $\mathcal{B}_0 = B(0, 1)$ et $\mathcal{B}_1 = B(1, 1)$, boules ouvertes de centres respectifs 0 et 1 et de rayon 1 , et $\mathcal{C} = \mathcal{B}'(0, 1)$, boule fermée de centre 0 et de rayon 1 . S'il existait au moins un point a dans $\mathcal{B}_0 \cap \mathcal{B}_1$, on aurait $1 = d(0, 1) \leq \sup(d(0, a), d(a, 1)) < 1$, ce qui est absurde ; par conséquent, $\mathcal{B}_0 \cap \mathcal{B}_1 = \emptyset$. Par ailleurs, $\mathcal{B}_0 \subset \mathcal{C}$. Pour tout point y de $\mathcal{B}_1$, $d(1, y) < 1$ donc $d(0, y) \leq \sup(d(0, 1), d(1, y)) = d(0, 1) = 1$; par suite, $\mathcal{B}_0 \subset \mathcal{C}$.

Ainsi, les boules ouvertes $B(0, 1) = \mathcal{B}_0$ et $B(1, 1) = \mathcal{B}_1$ sont incluses dans la boule fermée $\mathcal{C} = \mathcal{B}'(0, 1)$. $\square$

DÉFINITION 16.6. — Une partie bornée d'un espace métrique E est une partie de E que l'on peut inclure dans une boule fermée.

Si (x_n) est une suite de points d'un espace métrique (E, d) et ℓ un point de E , la suite (x_n) converge vers ℓ dans l'espace métrique E si elle converge vers ℓ dans l'espace topologique E , ce qui se traduit ici de la manière suivante : pour tout réel $\varepsilon > 0$, il existe un entier naturel N tel que $d(\ell, x_n) < \varepsilon$ pour tout entier $n \geq N$.

2. On a en fait démontré pour la distance d une inégalité plus forte que la relation triangulaire. Une distance d sur un ensemble E telle que $d(x, z) \leq \sup(d(x, y), d(y, z))$ quels que soient les éléments x , y et z de E , s'appelle une distance *ultramétrique*. La distance sur $\mathbb{Q}$ associée au nombre premier p dans l'exemple 16.3 est donc une distance ultramétrique, que l'on appelle la distance p -adique sur $\mathbb{Q}$.

En 1821, Augustin-Louis Cauchy considère comme intuitivement évident qu'une suite de nombres réels converge si, et seulement si, elle vérifie ce que nous appelons maintenant *le critère de Cauchy*.

DÉFINITION 16.7. — Une suite (x_n) de points d'un espace métrique (E, d) est une suite de Cauchy³ de E si, pour tout réel $\varepsilon > 0$, il existe un entier naturel N tel que, quels que soient les entiers $p \geq N$ et $q \geq N$, $d(x_p, x_q) < \varepsilon$.

Une suite qui converge dans un espace métrique E est une suite de Cauchy de E . Pour $\mathbb{R}$ ou $\mathbb{C}$ muni de sa distance canonique, la réciproque est vraie, ce qui, dans un espace métrique quelconque, peut tomber en défaut, comme nous le verrons à propos de la complétude (page 312 et suivantes).

Si (E, d) et (F, d') sont des espaces métriques et f une application de E dans F , la continuité de f en un point a de E se traduit ainsi : pour tout réel $\varepsilon > 0$, il existe un réel $\delta > 0$ tel que, pour tout point x de E , $d(a, x) < \delta$ implique $d'(f(a), f(x)) < \varepsilon$.

DÉFINITION 16.8. — Soit (E, d) et (F, d') des espaces métriques et f une application de E dans F .

- a) L'application f est uniformément continue sur E si, pour tout réel $\varepsilon > 0$, il existe un nombre réel $\delta > 0$ tel que, quels que soient les points x et y de E , $d(x, y) < \delta$ implique $d'(f(x), f(y)) < \varepsilon$.
- b) Si k est un réel positif ou nul, f est lipschitzienne de rapport k sur E si, quels que soient les points x et y de E , $d'(f(x), f(y)) < kd(x, y)$.
- c) L'application f est lipschitzienne sur E s'il existe un nombre réel $k \geq 0$ tel que f est lipschitzienne de rapport k sur E .

Soit (E, d) et (F, d') des espaces métriques et f une application de E dans F . Si la fonction f est uniformément continue sur E , elle est continue sur E , et si f est lipschitzienne sur E , elle est uniformément continue sur E . Les réciproques sont en général fausses — voir les exemples 8.14 et 8.15 (page 145).

Les notions de continuité, de suite convergente et de limite d'une suite sont de nature topologique, alors que celles de partie bornée, de suite de Cauchy, de continuité uniforme et d'application lipschitzienne sont des notions purement métriques, c'est-à-dire inhérentes à la distance. Aussi est-il naturel de s'intéresser aux distances pour lesquelles ces notions coïncident.

3. Cauchy écrit ([CAU1], chapitre VI, § 1^{er}, pages 125 et 126) : « [...] pour que la série

$$(1) \quad u_0, u_1, u_2, \dots, u_n, u_{n+1}, \&c. \dots$$

soit convergente, il est nécessaire et il suffit que des valeurs croissantes de n fassent converger indéfiniment la somme

$$s_n = u_0 + u_1 + u_2 + \&c. \dots + u_{n-1}$$

vers une limite fixe s ; en d'autres termes, il est nécessaire et il suffit que, pour des valeurs infiniment grandes du nombre n , les sommes

$$s_n, s_{n+1}, s_{n+2}, \&c. \dots$$

diffèrent de la limite s , et par conséquent entre elles, de quantités infiniment petites [...] c'est-à-dire [que] les sommes des quantités

$$u_n, u_{n+1}, u_{n+2}, \&c. \dots$$

prises, à partir de la première, en tel nombre que l'on voudra, finissent par obtenir constamment des valeurs [absolues] inférieures à toute limite assignable. »

■ Equivalences de distances

DÉFINITION 16.9. — Soit d_1 et d_2 des distances sur un même ensemble E .

- a) Les distances d_1 et d_2 sont topologiquement équivalentes si elles définissent la même topologie sur E — c'est-à-dire les mêmes ouverts —, ce qui signifie que l'identité de E est une application continue de (E, d_1) dans (E, d_2) et de (E, d_2) dans (E, d_1) — donc que Id_E est un homéomorphisme de (E, d_1) sur (E, d_2) .
- b) Les distances d_1 et d_2 sont uniformément équivalentes si Id_E est une application uniformément continue de (E, d_1) dans (E, d_2) et de (E, d_2) dans (E, d_1) .
- c) Les distances d_1 et d_2 sont équivalentes s'il existe des nombres réels k et k' strictement positifs tels que $k d_1 \leq d_2 \leq k' d_1$, ce qui signifie que Id_E est une application lipschitzienne de (E, d_1) dans (E, d_2) et de (E, d_2) dans (E, d_1) .

Pour des distances topologiquement équivalentes, les notions purement topologiques coïncident : les ouverts, fermés, compacts et connexes sont les mêmes, ainsi que les notions de limite et de continuité. Si des distances sont uniformément équivalentes, elles sont topologiquement équivalentes et les notions de suites de Cauchy, de continuité uniforme et de complétude sont les mêmes pour les deux distances. Enfin, si des distances sont équivalentes, elles sont uniformément équivalentes et la notion de partie bornée est la même pour les deux distances. Nous allons montrer sur deux exemples que si d_1 et d_2 sont des distances sur un ensemble E et si l'on considère les assertions :

- (1) d_1 et d_2 sont topologiquement équivalentes,
- (2) d_1 et d_2 sont uniformément équivalentes,
- (3) d_1 et d_2 sont équivalentes,

les implications (1) implique (2) et (2) implique (3) ne sont pas, en général, des équivalences.

16.5. Deux distances topologiquement équivalentes qui ne sont pas uniformément équivalentes.

Nous posons $E =]0, 1]$, nous notons d_1 la distance induite sur E par la distance canonique de $\mathbb{R}$ et nous définissons l'application :

$$\left| \begin{array}{l} d_2 : E \times E \longrightarrow \mathbb{R}_+ \\ (x, y) \longmapsto d_2(x, y) = \left| \frac{1}{x} - \frac{1}{y} \right|, \end{array} \right.$$

qui est clairement une distance sur E . L'inégalité :

$$d_1(x, y) = |x - y| \leq \frac{|x - y|}{xy} = d_2(x, y),$$

valable quels que soient les éléments x et y de E , montre que la topologie définie par d_2 est plus fine que celle qui est définie par d_1 — c'est-à-dire que la topologie définie par d_2 a plus d'ouverts que celle définie par d_1 — et que Id_E est une bijection lipschitzienne, donc continue, de (E, d_2) sur (E, d_1) .

Nous posons $F = [1, +\infty[$, nous notons d la distance induite sur F par la distance

canonique de $\mathbb{R}$ et nous considérons les bijections :

$$\left| \begin{array}{l} f : E \longrightarrow F \\ t \longmapsto f(t) = \frac{1}{t} \end{array} \right. \quad \text{et} \quad \left| \begin{array}{l} g : F \longrightarrow E \\ x \longmapsto g(x) = \frac{1}{x} \end{array} \right.$$

respectivement de E sur F et de F sur E . L'application f est évidemment une bijection continue de (E, d_1) sur (F, d) (topologie canonique des parties de $\mathbb{R}$). Quels que soient les éléments x et y de F :

$$d_2(g(x), g(y)) = \left| \frac{1}{g(x)} - \frac{1}{g(y)} \right| = |x - y| = d(x, y),$$

donc g est une bijection continue, car lipschitzienne, de (F, d) sur (E, d_2) ; or $\text{Id}_E = g \circ f$, donc Id_E est une bijection continue de (E, d_1) sur (E, d_2) .

Par conséquent, les distances d_1 et d_2 sont topologiquement équivalentes.

Considérons maintenant la suite $(u_n)_{n \geq 1}$ de points de E de terme général $u_n = \frac{1}{n}$. On a, pour tout entier $n \geq 1$:

$$d_1(u_n, u_{n+1}) = \left| \frac{1}{n} - \frac{1}{n+1} \right| = \frac{1}{n(n+1)} \quad \text{et} \quad d_2(u_n, u_{n+1}) = |n - (n+1)| = 1.$$

Soit α un réel > 0 . La suite $(1/n)$ converge vers 0 dans $\mathbb{R}$ muni de sa distance canonique, ce qui justifie le choix d'un entier $N \geq 1$ tel que, pour tout $n \geq N$:

$$\frac{1}{n(n+1)} < \alpha.$$

On a, pour tout entier naturel $n \geq N$, $d_1(u_n, u_{n+1}) < \alpha$ et $d_2(u_n, u_{n+1}) = 1$. Si Id_E était une application uniformément continue de (E, d_1) dans (E, d_2) , il existerait un réel $\alpha > 0$ tel que, quels que soient les points x et y de E , $d_1(x, y) < \alpha$ implique $d_2(x, y) < 1$, en contradiction avec ce qui précède. Il en résulte que Id_E n'est pas une application uniformément continue de (E, d_1) dans (E, d_2) , donc que les distances d_1 et d_2 ne sont pas uniformément équivalentes. $\square$

Remarquons que la suite (u_n) de l'exemple précédent 16.5 est une suite de Cauchy de (E, d_1) mais que ce n'est une suite de Cauchy de (E, d_2) . On aurait pu alors directement conclure à partir de là, sachant que des distances uniformément équivalentes ont les mêmes suites de Cauchy.

16.6. Deux distances uniformément équivalentes qui ne sont pas équivalentes.

On a, pour tout réel t , $-1 < \frac{t}{1+|t|} < 1$, ce qui justifie l'existence des applications :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow]-1, 1[\\ t \longmapsto f(t) = \frac{t}{1+|t|} \end{array} \right. \quad \text{et} \quad \left| \begin{array}{l} g :]-1, 1[\longrightarrow \mathbb{R} \\ x \longmapsto g(x) = \frac{x}{1-|x|} \end{array} \right.$$

Alors f est une bijection de $\mathbb{R}$ sur $]-1, 1[$ et $f^{-1} = g$. Or f et g sont continues, donc f est un homéomorphisme de $\mathbb{R}$ sur $]-1, 1[$ (topologie canonique des parties de $\mathbb{R}$). On a $f(0) = 0$, f est dérivable sur $\mathbb{R}$ et, pour tout nombre réel t :

$$f'(t) = \frac{1}{(1+|t|)^2} > 0,$$

d'où l'on déduit en particulier que f est strictement croissante.

On a de plus, quels que soient les réels $a \geq 0$ et $b \geq 0$:

$$(1) \quad f(a+b) = \frac{a+b}{1+(a+b)} = \frac{a}{(1+a)+b} + \frac{b}{(1+b)+a} \leq f(a) + f(b).$$

Notons d_1 la distance canonique de $\mathbb{R}$ et définissons l'application :

$$\left| \begin{array}{l} d_2 : \mathbb{R} \times \mathbb{R} \longrightarrow \mathbb{R}_+ \\ (x, y) \longmapsto d_2(x, y) = f(|x - y|). \end{array} \right.$$

Les propriétés de f montrent que d_2 est une distance sur $\mathbb{R}$ — pour la relation triangulaire, on part de la relation triangulaire pour d_1 et on utilise la croissance de f et l'inégalité (1) ci-dessus.

Pour tout réel $t \geq 0$, $1+t \geq 1$ donc $f(t) \leq t$. Par conséquent, quels que soient les nombres réels x et y , $d_2(x, y) \leq d_1(x, y)$. En particulier, $\text{Id}_{\mathbb{R}}$ est une application lipschitzienne, donc uniformément continue, de $(\mathbb{R}, d_1)$ dans $(\mathbb{R}, d_2)$.

Soit ε un réel > 0 . La continuité de $f^{-1} = g$ en 0 nous donne un réel $\alpha > 0$ tel que, pour tout réel u , $f(u) < \alpha$ entraîne $|u| < \varepsilon$. On appliquant ceci à $u = |x - y|$, nous obtenons que, quels que soient les réels x et y , $d_2(x, y) < \alpha$ implique $|x - y| < \varepsilon$, c'est-à-dire $d_1(x, y) < \varepsilon$. Par conséquent, $\text{Id}_{\mathbb{R}}$ est une application uniformément continue de $(\mathbb{R}, d_2)$ dans $(\mathbb{R}, d_1)$, donc les distances d_1 et d_2 sont uniformément équivalentes.

On a, pour tout réel $y \geq 0$, $d_1(0, y) = y$ et $d_2(0, y) = f(y) < 1$, donc on ne peut pas trouver de réel $k > 0$ tel que, pour tout réel x , $d_1(0, x) \leq k d_2(0, x)$ — prendre $x = k + 1$. Il en résulte que les distances d_1 et d_2 ne sont pas équivalentes. $\square$

■ Complétude

DÉFINITION 16.10. — Un espace métrique complet est un espace métrique dans lequel toute suite de Cauchy converge.

Par exemple $\mathbb{R}$, muni de sa distance canonique, est un espace métrique complet.

THÉORÈME 16.1. — Si E est un espace métrique complet, tout fermé non vide de E , muni de la distance induite, est un espace métrique complet.

16.7. Espace métrique qui n'est pas complet.

Il est facile de construire un espace métrique non complet. Il suffit d'enlever à un espace métrique complet la limite ℓ d'une suite convergente qui ne prend jamais la valeur ℓ . Il est par contre intéressant de montrer que $\mathbb{Q}$, muni de sa métrique habituelle, n'est pas un espace métrique complet.

Nous munissons donc $\mathbb{Q}$ de la distance d induite sur $\mathbb{Q}$ par celle de $\mathbb{R}$.

Nous considérons la suite $(u_n)_{n \geq 0}$ de terme général :

$$u_n = \sum_{k=0}^n \frac{1}{k!} = 1 + \frac{1}{1!} + \frac{1}{2!} + \cdots + \frac{1}{n!}.$$

Nous démontrons que (u_n) est une suite de Cauchy de $(\mathbb{Q}, d)$ mais qu'elle ne converge pas⁴ dans $\mathbb{Q}$. Si n est un entier naturel, on a, pour tout entier $p \geq 1$:

$$\begin{aligned} u_{n+p} - u_n &= \frac{1}{(n+1)!} + \frac{1}{(n+2)!} + \cdots + \frac{1}{(n+p)!} \\ &= \frac{1}{(n+1)!} \left(1 + \frac{1}{(n+2)} + \frac{1}{(n+2)(n+3)} + \cdots + \frac{1}{(n+2) \cdots (n+p)} \right) \\ &\leq \frac{1}{(n+1)!} \left(1 + \frac{1}{(n+1)} + \frac{1}{(n+1)^2} + \cdots + \frac{1}{(n+1)^{p-1}} \right) \\ &\leq \frac{1}{(n+1)!} \times \frac{1}{1 - \frac{1}{n+1}} = \frac{1}{n!n}. \end{aligned}$$

Or la suite $\left(\frac{1}{n!n}\right)$ converge vers 0, donc la suite (u_n) est de Cauchy dans $(\mathbb{Q}, d)$.

Supposons que (u_n) converge dans $(\mathbb{Q}, d)$ vers un rationnel r et notons (a, b) le représentant irréductible de r de dénominateur positif ; alors $a, b \in \mathbb{N}^*$. La suite (u_n) étant strictement croissante, $u_n < r$ pour tout entier naturel n .

On choisit un entier $n > b$. Pour tout entier $p \geq 1$, $u_{n+p} - u_n < 1/(n!n)$; le passage à la limite quand p tend vers l'infini et le fait que $u_n < r$ donnent les inégalités :

$$(1) \quad 0 < r - u_n \leq \frac{1}{n!n}.$$

Comme $k!$ divise $n!$ pour tout $k \in \llbracket 0, n \rrbracket$, on obtient, en réduisant la somme qui définit u_n au même dénominateur, un entier $c \geq 1$ tel que $u_n = c/n!$. D'autre part, b divise $n!$, donc $r = d/n!$ où d est un entier ≥ 1 . Alors $r - u_n = (d - c)/n!$, donc on déduit de (1) que $d - c \in \mathbb{N}^*$ et que :

$$\frac{d - c}{n!} \leq \frac{1}{n!n} \text{ donc } d - c \leq \frac{1}{n}.$$

Or $n > b \geq 1$ donc $n \geq 2$, ce qui montre que $d - c \leq \frac{1}{2}$, en contradiction avec l'appartenance de $d - c$ à $\mathbb{N}^*$.

Par conséquent, la suite (u_n) ne converge pas dans $\mathbb{Q}$. En conclusion, $(\mathbb{Q}, d)$ n'est pas un espace métrique complet. $\square$

Nous avons expliqué précédemment que la complétude est une propriété métrique. Ceci apparaîtra plus clairement dans l'exemple qui suit, où deux distances sur un même espace définissent la même topologie, bien que l'espace soit complet pour l'une des distances alors qu'il ne l'est pas pour l'autre.

16.8. Deux distances topologiquement équivalentes telles que l'une seulement définit un espace métrique complet.

Nous reprenons les distances d_1 et d_2 définies sur $E =]0, 1]$ dans l'exemple 16.5 par :

$$d_1(x, y) = |x - y| \text{ et } d_2(x, y) = \left| \frac{1}{x} - \frac{1}{y} \right|.$$

Nous avons vu que ces deux distances sont topologiquement équivalentes.

La suite $(u_n)_{n \geq 1}$, de terme général $u_n = 1/n$, est à valeurs dans E , et comme elle converge dans $\mathbb{R}$ vers 0, c'est une suite de Cauchy de (E, d_1) . Or $0 \notin E$, donc

4. On sait que cette suite converge vers e dans $\mathbb{R}$; ceci redémontre donc que e est irrationnel.

la suite (u_n) ne converge pas dans (E, d_1) . Par conséquent, (E, d_1) n'est pas un espace métrique complet.

Nous avons également établi dans l'exemple 16.5 qu'en notant d la distance induite sur $F = [1, +\infty[$ par celle de $\mathbb{R}$, la bijection $g : x \mapsto g(x) = 1/x$ de $[1, +\infty[$ sur $]0, 1]$ est une isométrie de (F, d) sur (E, d_2) , donc g conserve les distances et par conséquent les suites de Cauchy. Comme F est un fermé de $\mathbb{R}$, on déduit du théorème 16.1 que (F, d) est un espace métrique complet. Il en résulte que (E, d_2) est un espace métrique complet. $\square$

Remarquons que dans l'exemple précédent 16.8, la suite (u_n) de terme général $u_n = 1/n$ est une suite de Cauchy de (E, d_1) mais n'est pas une suite de Cauchy de (E, d_2) ; en effet $d_2(u_n, u_{n+1}) = |n - (n+1)| = 1$ pour tout entier $n \geq 1$. La convergence des suites est une notion de nature topologique et les distances d_1 et d_2 définissent sur $E =]0, 1]$ la même topologie, donc les deux espaces métriques (E, d_1) et (E, d_2) ont les mêmes suites convergentes. Si nous notons $\mathcal{C}$ l'ensemble des suites convergentes de points de E pour la distance d_1 , donc aussi pour d_2 , l'ensemble des suites de Cauchy de (E, d_1) contient strictement $\mathcal{C}$ alors que l'ensemble des suites de Cauchy de (E, d_2) est égal à $\mathcal{C}$.

16.9. Espace métrique non complet totalement discontinu⁵.

Posons $E = \left\{ \frac{1}{p} \mid p \in \mathbb{N}^* \right\}$ et munissons E de la distance d induite par celle de $\mathbb{R}$.

La suite $(u_n)_{n \geq 1}$, de terme général $u_n = 1/n$, est à valeurs dans E , et comme elle converge dans $\mathbb{R}$, c'est une suite de Cauchy de (E, d) . Or $0 \notin E$, donc (u_n) ne converge pas dans E . Par suite, (E, d) n'est pas un espace métrique complet.

Soit A une partie de E admettant au moins deux éléments distincts. La partie A contient les inverses $1/p$ et $1/q$ d'entiers naturels tels que $1 \leq p < q$. Choisissons un point a de $]0, 1[$ tel que $1/q < a < 1/(q-1)$ et posons $O_1 =]-\infty, a[\cap A$ et $O_2 =]a, +\infty[\cap A$. Alors $1/q \in O_1$, $1/p \in O_2$, O_1 et O_2 sont des ouverts de A pour la topologie induite sur A par celle de $\mathbb{R}$, donc par celle de E , $O_1 \cup O_2 = A$ et $O_1 \cap O_2 = \emptyset$. Il en résulte que A n'est pas une partie connexe de E .

En conclusion, les connexes de E sont l'ensemble vide et les singletons, donc l'espace topologique E est totalement discontinu. $\square$

Dans $\mathbb{R}$ — et même dans $\mathbb{R}^p$ — muni de sa distance canonique, les boules fermées sont compactes, donc l'intersection de toute suite décroissante modulo l'inclusion de boules fermées non vides en est une partie non vide. Ce résultat ne se généralise pas à un espace métrique quelconque, même si celui-ci est complet.

16.10. Suite décroissante de boules fermées dont l'intersection est vide.

Nous introduisons l'application :

$$\left| \begin{array}{l} d : \mathbb{N} \times \mathbb{N} \longrightarrow \mathbb{R}_+ \\ (m, n) \longmapsto d(m, n) = \begin{cases} 0 & \text{si } m = n, \\ 1 + \frac{1}{m+n} & \text{si } m \neq n. \end{cases} \end{array} \right.$$

5. Voir le chapitre 15, pages 297 et 298, définitions 15.16 et 15.17 et ce qui suit.

Soit m et n des entiers naturels. Il est clair que $d(m, n) = d(n, m)$ et que $d(m, n) = 0$ si, et seulement si, $m = n$. Si m, n et p sont des entiers naturels deux à deux distincts, on a $d(m, p) \leq 2 \leq d(m, n) + d(n, p)$, et on voit que l'inégalité $d(m, p) \leq d(m, n) + d(n, p)$ est également vraie lorsque m, n et p ne sont pas deux à deux distincts — il suffit d'examiner les deux cas : $m = p$, $m \neq p$. Par conséquent d est une distance sur $\mathbb{N}$.

Soit (u_n) une suite de Cauchy de l'espace métrique $(\mathbb{N}, d)$. Il existe un entier naturel n_0 tel que, quels que soient les entiers $p \geq n_0$ et $q \geq n_0$, $d(u_p, u_q) < 1/2$ donc $u_p = u_q$. Nous posons $\ell = u_{n_0}$. Alors $u_n = \ell$ pour tout entier $n \geq n_0$, donc la suite (u_n) converge dans $\mathbb{N}$ vers l'entier naturel ℓ . On conclut de cette étude que $(\mathbb{N}, d)$ est un espace métrique complet.

Nous notons, pour tout $n \in \mathbb{N}^*$, $\mathcal{B}_n$ la boule fermée de centre n et de rayon $1 + \frac{1}{2n}$. Soit $n \in \mathbb{N}^*$. Alors $d(n, n) = 0$ et, pour tout entier naturel $m \neq n$:

$$d(n, m) = 1 + \frac{1}{m+n},$$

donc $m \in \mathcal{B}_n$ si, et seulement si, $m > n$. Par conséquent, $\mathcal{B}_n = \{n, n+1, n+2, \dots\}$. Il en résulte que $(\mathcal{B}_n)_{n \in \mathbb{N}^*}$ est une suite décroissante modulo l'inclusion de boules fermées non vides de $\mathbb{N}$ dont l'intersection est vide. $\square$

■ Distance d'un point à une partie et distance entre deux parties

DÉFINITION 16.11. — Si (E, d) est un espace métrique, A une partie non vide de E et a un point de E , la distance de a à A est le réel positif ou nul :

$$\delta(a, A) = \inf \{ d(a, x) \mid x \in A \}.$$

Si (E, d) est un espace métrique, A une partie non vide de E et a un point de E , la distance de a à A est nulle si, et seulement si, a appartient à l'adhérence de A .

THÉORÈME 16.2. — Si (E, d) est un espace métrique, K un compact⁶ non vide de E et a un point de E , il existe au moins un point b_0 de K tel que :

$$\delta(a, K) = d(a, b_0).$$

On exprime l'existence du point b_0 de K affirmée dans le texte du théorème 16.2 en disant que la distance de a à K est atteinte (en b_0).

Dans $\mathbb{R}$ (ou $\mathbb{R}^p$ pour $p \in \mathbb{N}^*$) muni de sa distance canonique, le théorème 16.2 est valable en remplaçant le compact non vide K par un fermé non vide. Cependant, ceci est faux pour un espace métrique quelconque.

16.11. Distance d'un point à un fermé qui n'est pas atteinte.

Nous utilisons l'espace métrique $(\mathbb{N}, d)$ où d est la distance définie dans l'exemple précédent 16.10.

6. Voir le chapitre 15, pages 301 et suivantes.

On a donc :

$$\left| \begin{array}{l} d : \mathbb{N} \times \mathbb{N} \longrightarrow \mathbb{R}_+ \\ (m, n) \longmapsto d(m, n) = \begin{cases} 0 & \text{si } m = n, \\ 1 + \frac{1}{m+n} & \text{si } m \neq n. \end{cases} \end{array} \right.$$

Pour tout entier $n \geq 1$, $d(0, n) = 1 + (1/n) > 1$, donc le singleton $\{0\}$ est dans $(\mathbb{N}, d)$ la boule ouverte de centre 0 et de rayon 1, ce qui prouve que son complémentaire $\mathbb{N}^*$ dans $\mathbb{N}$ est un fermé de $\mathbb{N}$.

La suite $(\alpha_n)_{n \geq 1}$, de terme général $\alpha_n = d(0, n)$, converge vers 1 et $d(0, n) > 1$ pour tout « point » n de $\mathbb{N}^*$, donc la distance de 0 à $\mathbb{N}^*$ est égale à 1, alors qu'il n'existe aucun « point » n_0 de $\mathbb{N}^*$ en lequel cette distance est atteinte. $\square$

DÉFINITION 16.12. — Si (E, d) est un espace métrique et A_1 et A_2 des parties non vides de E , la distance entre A_1 et A_2 est le réel positif ou nul :

$$\delta(A_1, A_2) = \inf \{ d(x, y) \mid x \in A_1 \text{ et } y \in A_2 \}.$$

Il faut prendre garde que si (E, d) est un espace métrique, l'application :

$$\delta : (A_1, A_2) \mapsto \delta(A_1, A_2)$$

n'est pas une distance sur l'ensemble $\mathcal{E} = \mathcal{P}(E) \setminus \{\emptyset\}$ des parties non vides de E .

Si (E, d) est un espace métrique, F un fermé non vide et K un compact non vide de E , la distance entre F et K est nulle si, et seulement si, F rencontre K — ce qui signifie que $F \cap K \neq \emptyset$. Ceci devient faux si K est seulement fermé.

16.12. Des fermés non vides F_1 et F_2 qui sont disjoints alors que la distance entre F_1 et F_2 est nulle⁷.

Nous munissons $\mathbb{R}^2$ de sa distance canonique d et nous posons :

$$F_1 = \{ (x, y) \mid x, y \in \mathbb{R} \text{ et } xy = 1 \} \text{ et } F_2 = \mathbb{R} \times \{0\}.$$

Les applications $f : z = (x, y) \mapsto f(z) = xy - 1$ et $q : z = (x, y) \mapsto q(z) = y$ de $\mathbb{R}^2$ dans $\mathbb{R}$ sont continues, F_1 est l'image réciproque du fermé $\{0\}$ de $\mathbb{R}$ par f et F_2 l'image réciproque du même fermé par q , donc F_1 et F_2 sont des fermés de $\mathbb{R}^2$.

Si $c = (a, b)$ appartient à F_1 , $ab = 1$ donc $b \neq 0$, ce qui montre que c n'appartient pas à F_2 . Par conséquent, $F_1 \cap F_2$ est vide.

Considérons les suites $(u_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$ de points de $\mathbb{R}^2$, de termes généraux :

$$u_n = \left(n, \frac{1}{n} \right) \text{ et } v_n = (n, 0),$$

Pour tout entier $n \geq 1$, u_n appartient à F_1 et v_n à F_2 , donc :

$$0 \leq \delta(F_1, F_2) \leq d(u_n, v_n) = \frac{1}{n}.$$

Le passage à la limite quand n tend vers l'infini dans ces inégalités montre que :

$$\delta(F_1, F_2) = 0. \quad \square$$

7. Voir aussi les exemples 5.27 et 5.28 (page 94) du chapitre 5 consacré aux nombres réels.

Chapitre 17

Espaces vectoriels normés

Pour généraliser les notions de continuité et de différentiabilité étudiées sur $\mathbb{R}$, et plus généralement sur les espaces $\mathbb{R}^n$, plusieurs mathématiciens, en particulier Stefan Banach (1892-1945), introduisent en 1920 la notion d'espace vectoriel normé, déduite de celle de norme ; les axiomes de la norme reprennent les propriétés essentielles de la valeur absolue sur $\mathbb{R}$ et du module sur $\mathbb{C}$.

David Hilbert (1862-1943) avait généralisé des propriétés de $\mathbb{R}^n$ à des espaces de suites en y introduisant un produit scalaire. John Von Neumann (1903-1957) formalise ces espaces en 1927 et les appellent fort justement des espaces hilbertiens.

DÉFINITION 17.1. — Une norme sur un espace vectoriel réel ou complexe E est une application :

$$\left| \begin{array}{l} N : E \longrightarrow \mathbb{R}_+ \\ x \longmapsto N(x) \end{array} \right.$$

qui vérifie les trois axiomes suivants, valables pour tout scalaire λ et quels que soient les vecteurs x et y de E :

- (I) $N(x) = 0$ entraîne $x = 0$.
- (II) $N(\lambda x) = |\lambda| N(x)$.
- (III) $N(x + y) \leq N(x) + N(y)$ [relation triangulaire].

Si N est une norme, on déduit de (II) que $N(0) = 0$, ce qui montre que $N(x) = 0$ si, et seulement si, $x = 0$, et que $x \neq 0$ si, et seulement si, $N(x) > 0$.

DÉFINITION. — Un espace vectoriel normé est un couple (E, N) où E est un espace vectoriel réel ou complexe différent de $\{0\}$ et N une norme sur E .

THÉORÈME 17.1. — Si (E, N) est un espace vectoriel normé, l'application ;

$$\left| \begin{array}{l} d : E \times E \longrightarrow \mathbb{R}_+ \\ (x, y) \longmapsto d(x, y) = N(x - y) \end{array} \right.$$

est une distance sur E , appelée la distance sur E associée à N .

L'intérêt des distances provenant d'une norme est qu'elles rendent continues les opérations de l'espace vectoriel et qu'en particulier les translations sont des homéomorphismes.

La valeur absolue est une norme sur l'espace vectoriel réel $\mathbb{R}$ et le module une norme sur l'espace vectoriel réel ou complexe $\mathbb{C}$, la distance associée étant la distance canonique. Pour le rappeler, une norme sur un espace vectoriel réel ou complexe E est souvent notée $\|\cdot\|$, ce qui signifie que c'est l'application $\|\cdot\| : x \mapsto \|x\|$ de E dans $\mathbb{R}_+$, de distance associée $d : (x, y) \mapsto d(x, y) = \|x - y\|$.

■ Nécessité des axiomes

Nous allons montrer grâce à des exemples que chacun des trois axiomes de la définition 17.1 est nécessaire pour définir la notion de norme telle que nous la connaissons, c'est-à-dire qu'aucun des trois n'est une conséquence des deux autres.

17.1. Application N ne vérifiant que les axiomes (I) et (III).

Nous considérons l'application :

$$\left| \begin{array}{l} \varphi : \mathbb{R}_+ \longrightarrow \mathbb{R}_+ \\ t \longmapsto \varphi(t) = \ln(1 + t). \end{array} \right.$$

La fonction logarithme étant croissante, on a, quels que soient $t_1, t_2 \in \mathbb{R}_+$:

$$\begin{aligned} \varphi(t_1 + t_2) &= \ln(1 + t_1 + t_2) \leq \ln(1 + t_1 + t_2 + t_1 t_2) = \ln((1 + t_1)(1 + t_2)) \\ &\leq \ln(1 + t_1) + \ln(1 + t_2) = \varphi(t_1) + \varphi(t_2). \end{aligned}$$

Nous introduisons un espace vectoriel normé $(E, \|\cdot\|)$ (il y en a !) et l'application :

$$\left| \begin{array}{l} N : E \longrightarrow \mathbb{R}_+ \\ x \longmapsto N(x) = \varphi(\|x\|) = \ln(1 + \|x\|). \end{array} \right.$$

L'application φ est strictement croissante et $\varphi(0) = 0$, donc, pour tout réel $t \geq 0$, $\varphi(t) = 0$ implique $t = 0$. Il en résulte que si x est un vecteur de E et si $N(x) = 0$, alors $\|x\| = 0$ donc $x = 0$. On a, quels que soient les vecteurs x et y de E :

$$N(x + y) = \varphi(\|x + y\|) \leq \varphi(\|x\| + \|y\|) \leq \varphi(\|x\|) + \varphi(\|y\|) = N(x) + N(y).$$

Les axiomes (I) et (III) sont donc vérifiés par N . Si l'on choisit un vecteur non nul a de E et si l'on pose $\alpha = \|a\|$ — c'est un réel > 0 — et $u = (1/\alpha)a$, alors $\|u\| = 1$ et :

$$N(2u) = \ln(1 + \|2u\|) = \ln(1 + 2\|u\|) = \ln 3 \neq \ln 4 = 2 \ln 2 = 2N(u).$$

L'axiome (II) n'est donc pas vérifié par N . $\square$

17.2. Application N ne vérifiant que les axiomes (I) et (II).

Nous considérons l'espace vectoriel réel $\mathbb{R}^2$ et l'application :

$$\left| \begin{array}{l} N : \mathbb{R}^2 \longrightarrow \mathbb{R}_+ \\ x = (x_1, x_2) \longmapsto N(x) = \begin{cases} \text{Max}(|x_1|, |x_2|) & \text{si } x_2 \neq 0, \\ 2|x_1| & \text{si } x_2 = 0. \end{cases} \end{array} \right.$$

Remarquons que, si $x = (x_1, x_2)$ est un vecteur de $\mathbb{R}^2$, alors $0 \leq |x_1| \leq N(x)$ et $0 \leq |x_2| \leq N(x)$. Par conséquent, si $x = (x_1, x_2)$ est un vecteur de $\mathbb{R}^2$ et si $N(x) = 0$, on a $|x_1| = 0$ et $|x_2| = 0$, donc $x = (0, 0)$.

Soit $x = (x_1, x_2)$ un vecteur de $\mathbb{R}^2$ et λ un nombre réel. Alors $\lambda x = (\lambda x_1, \lambda x_2)$. Si $\lambda = 0$, $\lambda x = (0, 0)$, donc $N(\lambda x) = 2|0| = 0 = 0 \times N(x)$. Nous supposons maintenant que $\lambda \neq 0$. Si $x_2 \neq 0$, $\lambda x_2 \neq 0$ et $\lambda > 0$, donc :

$$\begin{aligned} N(\lambda x) &= \text{Max}(|\lambda x_1|, |\lambda x_2|) = \text{Max}(|\lambda| |x_1|, |\lambda| |x_2|) \\ &= |\lambda| \text{Max}(|x_1|, |x_2|) = |\lambda| N(x) \end{aligned}$$

et, si $x_2 = 0$, $\lambda x_2 = 0$, donc $N(\lambda x) = 2|\lambda x_1| = |\lambda| |x_1| = |\lambda| N(x)$.

Les axiomes (i) et (ii) sont donc vérifiés par N . Or :

$$N((1, -1) + (1, 1)) = N((2, 0)) = 4 > 2 = 1 + 1 = N((1, -1)) + N((1, 1))$$

donc l'axiome (iii) n'est pas vérifié par N . $\square$

17.3. Application N ne vérifiant que les axiomes (ii) et (iii), en dimension finie.

Nous considérons l'espace vectoriel réel $\mathbb{R}^2$ et l'application :

$$\left| \begin{array}{l} N : \mathbb{R}^2 \longrightarrow \mathbb{R}_+ \\ x = (x_1, x_2) \longmapsto N(x) = |x_1|. \end{array} \right.$$

Soit $x = (x_1, x_2)$ et $y = (y_1, y_2)$ des vecteurs de l'espace $\mathbb{R}^2$ et λ un nombre réel. Alors $\lambda x = (\lambda x_1, \lambda x_2)$, donc $N(\lambda x) = |\lambda x_1| = |\lambda| |x_1| = |\lambda| N(x)$. On a $x + y = (x_1 + y_1, x_2 + y_2)$, donc $N(x + y) = |x_1 + y_1| \leq |x_1| + |y_1| = N(x) + N(y)$. Les axiomes (ii) et (iii) sont donc vérifiés par N , et comme $(0, 1) \neq (0, 0)$ alors que $N((0, 1)) = 0$, l'axiome (i) ne l'est pas. $\square$

17.4. Application N ne vérifiant que les axiomes (ii) et (iii), en dimension infinie.

Nous considérons des réels a et b tels que $a < b$, nous notons $E = L([a, b], \mathbb{R})$ l'espace vectoriel réel des applications de $[a, b]$ dans $\mathbb{R}$ intégrables au sens de Lebesgue sur $[a, b]$ et nous introduisons l'application :

$$\left| \begin{array}{l} N : E \longrightarrow \mathbb{R}_+ \\ f \longmapsto N(f) = \int_a^b |f(t)| dt. \end{array} \right.$$

Les propriétés de l'intégrale montrent que N vérifie les axiomes (ii) et (iii). Cependant, si l'on choisit une partie finie non vide ou dénombrable D de $[a, b]$ et si f est une application de $[a, b]$ dans $\mathbb{R}$, nulle sur $[a, b] \setminus D$ et telle que $f(t) \neq 0$ pour tout point t de D — par exemple $f(t) = 1$ pour tout $t \in D$ —, alors f est un vecteur de E , $N(f) = 0$ et $f \neq 0$: l'axiome (i) n'est pas vérifié. $\square$

Dans la situation des exemples 17.3 et 17.4 — un couple (E, N) vérifiant les axiomes (ii) et (iii) mais pas l'axiome (i) —, l'application N est appelée une semi-norme sur E , l'ensemble $F = \{x \in E \mid N(x) = 0\}$ est un sous-espace vectoriel de E et on obtient une norme N_0 sur l'espace vectoriel quotient E/F en posant, pour tout vecteur z de E/F , $N_0(z) = N(x)$ où x est un vecteur (quelconque) de E dont z est la classe d'équivalence.

■ Somme de deux parties

Soit A et B des parties d'un espace vectoriel normé. On note $A + B$ l'ensemble des vecteurs $x + y$ pour x décrivant A et y décrivant B . Si A est ouverte, la somme $A + B$ est ouverte, et ce quel que soit B . En revanche, si A et B sont fermées, la somme $A + B$ peut ne pas l'être.

17.5. Somme de deux parties fermées qui n'est pas fermée¹.

Munissons l'espace vectoriel réel $\mathbb{R}^2$ de l'une des ses normes usuelles, qui en définit la topologie canonique. Nous posons $A = \{x = (x_1, x_2) \mid x_1 \geq 0 \text{ et } x_1 x_2 = 1\}$ et $B = \{x = (x_1, x_2) \mid x_1 \geq 0 \text{ et } x_1 x_2 = -1\}$. Les applications :

$$\left| \begin{array}{l} p : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ x = (x_1, x_2) \longmapsto p(x) = x_1 \end{array} \right. \quad \text{et} \quad \left| \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R} \\ x = (x_1, x_2) \longmapsto f(x) = x_1 x_2 \end{array} \right.$$

sont continues sur $\mathbb{R}^2$. On a :

$$A = p^{-1}(\mathbb{R}_+) \cap f^{-1}(\{1\}) \quad \text{et} \quad B = p^{-1}(\mathbb{R}_+) \cap f^{-1}(\{-1\})$$

donc A et B sont des fermés de $\mathbb{R}^2$ — comme intersections d'images réciproques de fermés de $\mathbb{R}$ par p et par f .

Pour tout vecteur $x = (x_1, x_2)$ appartenant à A ou à B , $x_1 \geq 0$ et $x_1 \neq 0$, donc $x_1 > 0$. Par suite, pour tout vecteur $z = (z_1, z_2)$ de $A + B$, $z_1 > 0$, donc $A + B$ est inclus dans $]0, +\infty[\times \mathbb{R}$.

Soit $z = (z_1, z_2)$ un vecteur appartenant à $]0, +\infty[\times \mathbb{R}$. On a $z_1 > 0$. Nous supposons d'abord que $z_2 > 0$. Soit $x = (x_1, x_2)$ un vecteur de A et $y = (y_1, y_2)$ un vecteur de B ; alors $x_1 > 0$, $y_1 > 0$, $x_1 x_2 = 1$ et $y_1 y_2 = -1$, donc :

$$x_2 + y_2 = \frac{1}{x_1} - \frac{1}{y_1} = \frac{y_1 - x_1}{x_1 y_1}$$

d'où l'on déduit qu'il y a équivalence entre les assertions :

- (I) $x + y = z$.
- (II) $x_1 + y_1 = z_1$ et $x_2 + y_2 = z_2$.
- (III) $x_1 + y_1 = z_1$ et $\frac{y_1 - x_1}{x_1 y_1} = z_2$.
- (IV) $x_1 + y_1 = z_1$ et $\frac{z_1 - 2x_1}{x_1(z_1 - x_1)} = z_2$.
- (V) $x_1 + y_1 = z_1$ et $z_2(x_1)^2 - (2 + z_1 z_2)x_1 + z_1 = 0$.

Le discriminant de l'équation $z_2 t^2 - (2 + z_1 z_2)t + z_1 = 0$, du second degré et d'inconnue t , est $\Delta = 4 + (z_1)^2 (z_2)^2 > 0$, donc l'une de ses solutions dans $\mathbb{R}$ est :

$$t = \frac{2 + z_1 z_2 - \sqrt{\Delta}}{z_2} = z_1 + \frac{2}{z_2} - \sqrt{\frac{4}{(z_2)^2} + (z_1)^2}.$$

De plus :

$$\left(z_1 + \frac{2}{z_2}\right)^2 = (z_1)^2 + \frac{4}{(z_2)^2} + \frac{4z_1}{z_2} > \frac{4}{(z_2)^2} + (z_1)^2 \quad \text{et} \quad \sqrt{\frac{4}{(z_2)^2} + (z_1)^2} > \frac{2}{z_2}$$

donc $0 < t < z_1$. Nous posons :

$$x_1 = t, \quad x_2 = \frac{1}{t}, \quad y_1 = z_1 - t \quad \text{et} \quad y_2 = -\frac{1}{z_1 - t}.$$

1. Voir aussi l'exemple 5.25, page 93.

Alors $x = (x_1, x_2)$ est un vecteur de A et $y = (y_1, y_2)$ un vecteur de B , et l'équivalence entre les assertions (i) et (v) montre que $z = x + y$. Si $z_2 < 0$, l'étude précédente fournit des vecteurs $x = (x_1, x_2)$ de A et $y = (y_1, y_2)$ de B tels que $(z_1, -z_2) = x + y$, et alors $z = (y_1, -y_2) + (x_1, -x_2)$, $(y_1, -y_2)$ appartient à A et $(x_1, -x_2)$ à B . Enfin, si $z_2 = 0$:

$$z = (z_1, 0) = \left(\frac{z_1}{2}, \frac{z_1}{2}\right) + \left(\frac{z_1}{2}, -\frac{z_1}{2}\right), \quad \left(\frac{z_1}{2}, \frac{z_1}{2}\right) \in A \quad \text{et} \quad \left(\frac{z_1}{2}, -\frac{z_1}{2}\right) \in B.$$

Ainsi, dans tous les cas, z appartient à $A + B$. En conclusion, $A + B =]0, +\infty[\times \mathbb{R}$ donc la somme $A + B$ des fermés A et B n'est pas un fermé de $\mathbb{R}^2$. $\square$

■ Comparaison des normes

DÉFINITION 17.2. — Soit E un espace vectoriel réel ou complexe et N_1 et N_2 des normes sur E .

- a) N_1 est plus fine que N_2 s'il existe un réel $k > 0$ tel que $N_2 \leq k N_1$.
- b) N_1 est équivalente à N_2 s'il existe des réels $k > 0$ et $k' > 0$ tels que :

$$k' N_1 \leq N_2 \leq k N_1.$$

- c) N_1 est strictement plus fine que N_2 si N_1 est plus fine que N_2 et si N_1 n'est pas équivalente à N_2 .

Soit E un espace vectoriel réel ou complexe. La relation binaire «est équivalente à» est — heureusement — une relation d'équivalence sur l'ensemble des normes sur E . En raison de la symétrie, on dit que les normes N_1 et N_2 sont équivalentes pour exprimer que N_1 (ou N_2) est équivalente à N_2 (ou N_1). Les normes N_1 et N_2 sont équivalentes si, et seulement si, N_1 est plus fine que N_2 et N_2 plus fine que N_1 . La relation de finesse, à savoir la relation binaire «est plus fine que», est une relation de préordre sur l'ensemble des normes sur E , ce qui signifie qu'elle est réflexive et transitive. L'équivalence de deux normes équivaut à l'équivalence de leurs distances associées, et on démontre que pour des distances issues de normes, les notions de distances équivalentes, uniformément équivalentes et topologiquement équivalentes coïncident. Sur un espace vectoriel réel ou complexe de dimension finie, les normes sont deux à deux équivalentes.

17.6. Norme strictement plus fine qu'une autre.

Considérons l'espace vectoriel réel $E = \mathcal{C}([0, 1], \mathbb{R})$ des applications continues de $[0, 1]$ dans $\mathbb{R}$. Les applications :

$$\left| \begin{array}{l} N_1 : E \longrightarrow \mathbb{R}_+ \\ f \longmapsto N_1(f) = \int_0^1 |f(t)| dt \end{array} \right. \quad \text{et} \quad \left| \begin{array}{l} N_\infty : E \longrightarrow \mathbb{R}_+ \\ f \longmapsto N_\infty(f) = \sup_{t \in [0, 1]} |f(t)|. \end{array} \right.$$

sont des normes sur E — c'est un résultat classique.

Pour tout vecteur f de E , on a, en posant $M = N_\infty(f) = \sup_{t \in [0, 1]} |f(t)|$:

$$N_1(f) = \int_0^1 \underbrace{|f(t)|}_{\leq M} dt \leq (1 - 0)M = N_\infty(f),$$

donc $N_1 \leq N_\infty$: la norme N_∞ est plus fine que N_1 . Supposons que N_∞ et N_1 soient équivalentes. Alors N_1 est plus fine que N_∞ , donc il existe un réel $k > 0$ tel que $N_\infty \leq kN_1$. Nous associons à tout entier $n \geq 1$ l'application :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ t \longmapsto f_n(t) = \begin{cases} n - n^2 t & \text{si } t \in \left[0, \frac{1}{n}\right], \\ 0 & \text{si } t \in \left]\frac{1}{n}, 1\right]. \end{cases} \end{array} \right.$$

Soit n un entier ≥ 1 . Comme $f_n(1/n) = 0$, f_n est continue sur $[0, 1]$ donc f_n est un vecteur de E , on a $N_\infty(f_n) = n$ et :

$$N_1(f_n) = \int_0^1 |f_n(t)| dt = \int_0^{1/n} (n - n^2 t) dt = \left[nt - \frac{n^2 t^2}{2} \right]_{t=0}^{t=1/n} = \frac{1}{2}.$$

Par suite, pour tout $n \geq 1$, $n \leq k/2$, ce qui est bien sûr impossible. En conclusion, N_∞ n'est pas équivalente à N_1 , donc N_∞ est strictement plus fine que N_1 . $\square$

17.7. Normes qui ne sont pas comparables par la relation de finesse.

Considérons l'espace vectoriel réel $E = \mathbb{R}[X]$. Si $P = \sum_{k=0}^{+\infty} a_k X^k$ est un vecteur de E , alors, en choisissant un entier naturel $p \geq \deg P$, $a_k = 0$ pour tout entier $k > p$, ce qui justifie l'existence des réels positifs ou nuls :

$$\|P\|_\infty = \max_{k \in \mathbb{N}} |a_k| \left(= \max_{0 \leq k \leq p} |a_k| \right) \text{ et } \|P\|_h = \sum_{k=0}^{+\infty} \frac{|a_k|}{k+1} \left(= \sum_{k=0}^p \frac{|a_k|}{k+1} \right).$$

Il est clair que $\|\cdot\|_\infty$ et $\|\cdot\|_h$ sont des normes sur E .

Nous introduisons les suites $(P_n)_{n \geq 0}$ et $(Q_n)_{n \geq 0}$ de vecteurs de E de termes généraux :

$$P_n = X^n \text{ et } Q_n = \sum_{k=0}^n X^k.$$

On a, pour tout entier naturel n :

$$\|P_n\|_\infty = \|Q_n\|_\infty = 1, \quad \|P_n\|_h = \frac{1}{n+1} \text{ et } \|Q_n\|_h = \sum_{k=0}^n \frac{1}{k+1} = \sum_{p=1}^{n+1} \frac{1}{p}.$$

La suite $(\|P_n\|_h)_{n \geq 0}$ converge vers 0, donc $\|\cdot\|_h$ n'est pas plus fine que $\|\cdot\|_\infty$, et la suite $(\|Q_n\|_h)_{n \geq 0}$ diverge vers $+\infty$ (série harmonique), donc $\|\cdot\|_h$ n'est pas plus fine que $\|\cdot\|_\infty$. En conclusion, les normes $\|\cdot\|_\infty$ et $\|\cdot\|_h$ ne sont pas comparables par la relation de finesse. $\square$

■ Continuité des applications linéaires

La continuité d'une application linéaire f d'un espace vectoriel normé $(E, \|\cdot\|_E)$ dans un espace vectoriel normé $(F, \|\cdot\|_F)$ équivaut à sa continuité en 0, ou encore à l'existence d'un réel $k \geq 0$ tel que, pour tout vecteur x de E , $\|f(x)\|_F \leq k \|x\|_E$. De plus, toute application linéaire d'un espace vectoriel normé de dimension finie dans un espace vectoriel normé est continue.

17.8. Forme linéaire discontinue.

Nous reprenons l'espace vectoriel réel $E = \mathcal{C}([0, 1], \mathbb{R})$ et la norme, notée $\|\cdot\|_1$ au lieu de N_1 :

$$\|\cdot\|_1 : f \mapsto \|f\|_1 = \int_0^1 |f(t)| dt,$$

de l'exemple 17.6 et nous munissons $\mathbb{R}$ de la valeur absolue. L'application :

$$\left| \begin{array}{l} u : E \longrightarrow \mathbb{R} \\ f \longmapsto u(f) = f(0) \end{array} \right.$$

est une forme linéaire sur E . Nous utilisons la suite $(f_n)_{n \geq 1}$ de vecteurs de E définie dans l'exemple 17.6 et nous introduisons la suite $(g_n)_{n \geq 1}$ de vecteurs de E de terme général :

$$g_n = \frac{1}{n} f_n.$$

Pour tout $n \in \mathbb{N}^*$, $u(f_n) = f_n(0) = n$ donc $u(g_n) = 1$, et $\|g_n\|_1 = \frac{1}{n} \|f_n\|_1 = \frac{1}{2n}$.

Il en résulte que la suite (g_n) converge dans E vers le vecteur zéro de E , c'est-à-dire l'application nulle 0 de $[0, 1]$ dans $\mathbb{R}$, alors que la suite $(u(g_n))$ ne converge pas dans $\mathbb{R}$ vers $u(0) = 0$. Par suite, u n'est pas une forme linéaire continue sur E . $\square$

Remarquons que le noyau H de la forme linéaire u de l'exemple précédent 17.8 n'est pas un fermé de E ; en effet, en considérant la suite $(h_n)_{n \geq 1}$ de vecteurs de E de terme général $h_n = g_n - 1$ — c'est-à-dire $h_n(t) = g_n(t) - 1$ pour tout point t de $[0, 1]$ —, alors, pour tout entier $n \geq 1$, $u(h_n) = 0$ donc $h_n \in H$, la suite (h_n) converge dans E vers l'application constante $h : t \mapsto h(t) = -1$ de $[0, 1]$ dans $\mathbb{R}$ et $u(h) = -1$, ce qui montre h n'appartient pas à H . Ce résultat est particulier aux espaces vectoriels normés de dimension infinie, puisque dans un espace vectoriel normé de dimension finie, tous les sous-espaces vectoriels sont fermés.

17.9. Automorphisme discontinu.

Nous reprenons l'espace vectoriel normé $(E, \|\cdot\|_1)$ et la forme linéaire u de l'exemple précédent 17.8, nous choisissons un vecteur non nul g de E tel que $u(g) = 0$ — par exemple l'application $g : t \mapsto g(t) = t$ de $[0, 1]$ dans $\mathbb{R}$ — et nous considérons l'application :

$$\left| \begin{array}{l} \psi : E \longrightarrow E \\ f \longmapsto \psi(f) = f - u(f)g \end{array} \right.$$

qui, comme u , est linéaire. Nous démontrons que ψ est un automorphisme de E .

Si f est un vecteur de E et si $\psi(f) = 0$, alors $f = u(f)g$ donc $u(f) = u(f)u(g) = 0$, ce qui montre que $f = 0$. Par conséquent, l'endomorphisme ψ de E est injectif.

Comme $u(g) = 0$, on a, pour tout vecteur f de E :

$$\psi(f + u(f)g) = f + u(f)g - u(f + u(f)g)g = f + u(f)g - (u(f) + u(f)u(g))g = f,$$

donc ψ est une surjection de E sur E . Finalement, ψ est un automorphisme de l'espace vectoriel E .

Supposons que l'automorphisme ψ soit continu. On a, pour tout vecteur f de E , $u(f)g = f - \psi(f)$, donc l'endomorphisme $\varphi : f \mapsto \varphi(f) = u(f)g$ de E est continu.

Il existe donc un réel $k \geq 0$ tel que, pour tout $f \in E$, $\|\varphi(f)\|_1 \leq k \|f\|_1$. Pour tout vecteur f de E , $\|\varphi(f)\|_1 = |u(f)| \|g\|_1$ donc, puisque $\|g\|_1 > 0$:

$$|u(f)| = \frac{1}{\|g\|_1} \|\varphi(f)\|_1 \leq \frac{k}{\|g\|_1} \|f\|_1.$$

Par conséquent, la forme linéaire u est continue, en contradiction avec le résultat de l'exemple 17.8. En conclusion, l'automorphisme ψ de E n'est pas continu². $\square$

Soit $(E, \|\cdot\|)$ un espace vectoriel normé de « corps de base » $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$. Si φ est une forme linéaire continue sur E , il existe un réel $k \geq 0$ tel que $|\varphi(x)| \leq k \|x\|$ pour tout vecteur x de E , ce qui justifie l'existence du nombre réel positif ou nul :

$$\|\varphi\|_\ell = \sup_{x \in E \setminus \{0\}} \frac{|\varphi(x)|}{\|x\|} = \sup_{\substack{x \in E \\ \|x\|=1}} |\varphi(x)|.$$

Alors $\|\cdot\|_\ell$ est une norme sur l'espace vectoriel $\mathcal{L}_c(E, \mathbb{K})$ des formes linéaires continues sur E , appelée la norme linéaire.

La borne supérieure qui définit $\|\varphi\|_\ell$ n'est pas toujours atteinte.

17.10. Norme linéaire d'une forme linéaire continue qui n'est pas atteinte.

Nous reprenons l'espace vectoriel réel $E = \mathcal{C}([0, 1], \mathbb{R})$ des applications continues de $[0, 1]$ dans $\mathbb{R}$ et nous le munissons de la norme :

$$\left| \begin{array}{l} \|\cdot\|_\infty : E \longrightarrow \mathbb{R}_+ \\ f \longmapsto \|f\|_\infty = \sup_{t \in [0, 1]} |f(t)|. \end{array} \right.$$

La linéarité de l'intégrale sur un segment montre que l'application :

$$\left| \begin{array}{l} u : E \longrightarrow \mathbb{R} \\ f \longmapsto u(f) = \int_0^{1/2} f(t) dt - \int_{1/2}^1 f(t) dt \end{array} \right.$$

est une forme linéaire sur E . On a, pour tout vecteur f de E :

$$\begin{aligned} |u(f)| &= \left| \int_0^{1/2} f(t) dt - \int_{1/2}^1 f(t) dt \right| \leq \left| \int_0^{1/2} f(t) dt \right| + \left| \int_{1/2}^1 f(t) dt \right| \\ &\leq \int_0^{1/2} |f(t)| dt + \int_{1/2}^1 |f(t)| dt = \int_0^1 \underbrace{|f(t)|}_{\leq \|f\|_\infty} dt \leq \|f\|_\infty, \end{aligned}$$

donc la forme linéaire u est continue et sa norme linéaire est inférieure ou égale à 1. Nous démontrons que $\|u\|_\ell = 1$.

Nous introduisons pour ceci la suite $(f_n)_{n \geq 2}$ d'applications de $[0, 1]$ dans $\mathbb{R}$ de

2. On démontre à l'aide de l'axiome du choix (chapitre 1, page 14) que, sur tout espace vectoriel normé de dimension infinie, il existe des formes linéaires discontinues, ce qui permet, par la construction que nous venons de faire, d'obtenir des automorphismes discontinus de cet espace.

terme général :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ t \longmapsto f_n(t) = \begin{cases} 1 & \text{si } 0 \leq t \leq \frac{1}{2} - \frac{1}{n}, \\ \frac{n}{2} - nt & \text{si } \frac{1}{2} - \frac{1}{n} < t < \frac{1}{2} + \frac{1}{n}, \\ -1 & \text{si } \frac{1}{2} + \frac{1}{n} \leq t \leq 1. \end{cases} \end{array} \right.$$

Pour tout $n \geq 2$, on a $f_n(t) = (n/2) - nt$ pour tout $t \in [(1/2) - (1/n), (1/2) + (1/n)]$, donc f_n est continue sur le segment $[0, 1]$. Par conséquent, $(f_n)_{n \geq 2}$ est une suite de vecteurs de E . Pour tout entier $n \geq 2$, f_n est décroissante sur $[0, 1]$, donc $-1 = f_n(1) \leq f_n(t) \leq f_n(0) = 1$ pour tout point t de $[0, 1]$, ce qui montre que $\|f_n\|_\infty = 1$, et on a :

$$\begin{aligned} u(f_n) &= \int_0^{1/2} f_n(t) dt - \int_{1/2}^1 f_n(t) dt \\ &= \int_0^{\frac{1}{2} - \frac{1}{n}} dt + \int_{\frac{1}{2} - \frac{1}{n}}^{\frac{1}{2}} \left(\frac{n}{2} - nt \right) dt - \int_{\frac{1}{2}}^{\frac{1}{2} + \frac{1}{n}} \left(\frac{n}{2} - nt \right) dt + \int_{\frac{1}{2} + \frac{1}{n}}^1 dt \\ &= \left(\frac{1}{2} - \frac{1}{n} \right) + \frac{n}{2} \left[t - t^2 \right]_{t=\frac{1}{2} - \frac{1}{n}}^{\frac{1}{2}} - \frac{n}{2} \left[t - t^2 \right]_{t=\frac{1}{2}}^{\frac{1}{2} + \frac{1}{n}} + \left(1 - \left(\frac{1}{2} + \frac{1}{n} \right) \right) \\ &= \left(\frac{1}{2} - \frac{1}{n} \right) + \frac{n}{2} \frac{1}{n^2} - \frac{n}{2} \frac{1}{n^2} + \left(\frac{1}{2} - \frac{1}{n} \right) = 1 - \frac{1}{n} \geq 0 \end{aligned}$$

donc $|u(f_n)| = 1 - (1/n)$, d'où l'on déduit que $1 - (1/n) \leq \|u\|_\ell \leq 1$. En passant à la limite quand n tend vers l'infini, on obtient l'égalité $\|u\|_\ell = 1$.

Nous prouvons enfin qu'il n'existe aucun vecteur f de E tel que $\|f\|_\infty = 1$ et $|u(f)| = 1$. Soit f un vecteur de E tel que $\|f\|_\infty = 1$, c'est-à-dire un vecteur de la sphère unité de $(E, \|\cdot\|_\infty)$. Alors, pour tout point t de $[0, 1]$, $-1 \leq f(t) \leq 1$. Les fonctions $t \mapsto 1 - f(t)$ et $t \mapsto 1 + f(t)$ sont continues et positives sur les segments $[0, 1/2]$ et $[1/2, 1]$, donc :

$$1 - u(f) = \underbrace{\int_0^{1/2} (1 - f(t)) dt}_{= I \geq 0} + \underbrace{\int_{1/2}^1 (1 + f(t)) dt}_{= J \geq 0}.$$

S'il existe un point c du segment $[0, 1/2]$ tel que $f(c) \neq 1$, alors $I > 0$; sinon, $f(t) = 1$ pour tout $t \in [0, 1/2]$, en particulier $f(1/2) = 1 > -1$, et la continuité de f en $1/2$ fournit un point c de $]1/2, 1[$ tel que $f(c) > -1$, ce qui montre que $J > 0$. Il en résulte que $1 - u(f) > 0$. On prouve de même que $1 + u(f) > 0$. Finalement, $-1 < u(f) < 1$ donc $|u(f)| < 1$. En conclusion, la norme linéaire 1 de u n'est pas atteinte sur la sphère unité de l'espace vectoriel normé $(E, \|\cdot\|_\infty)$. $\square$

■ Parties compactes et parties fermées et bornées

THÉORÈME 17.2. — Une partie K d'un espace vectoriel normé est compacte si, et seulement si, de toute suite de vecteurs de K , on peut extraire une sous-suite qui converge dans K .

THÉORÈME 17.3. — Dans un espace vectoriel normé, toute partie compacte est fermée et bornée et, dans un espace vectoriel normé de dimension finie, toute partie fermée et bornée est compacte.

17.11. Partie fermée et bornée qui n'est pas compacte.

Nous considérons l'espace vectoriel réel $E = \mathbb{R}[X]$ des polynômes à coefficients réels et nous posons, pour tout vecteur $P = \sum_{k=0}^{+\infty} a_k X^k$ de E :

$$\|P\| = \max_{k \in \mathbb{N}} |a_k| \quad \left(= \max_{0 \leq k \leq p} |a_k| \text{ si } p \in \mathbb{N} \text{ et } \deg(P) \leq p \right).$$

Clairement, $\|\cdot\|$ est une norme sur E . La boule unité fermée $\mathcal{B} = \{P \in E \mid \|P\| \leq 1\}$ de l'espace vectoriel normé $(E, \|\cdot\|)$ est bien sûr fermée et bornée. Nous introduisons la suite $(P_n)_{n \geq 0}$ de vecteurs de E , de terme général $P_n = X^n$. Pour tout $n \in \mathbb{N}$, $\|P_n\| = 1$, donc (P_n) est une suite de vecteurs de $\mathcal{B}$. Quels que soient les entiers naturels distincts m et n , $P_m - P_n = X^m - X^n$ donc $\|P_m - P_n\| = 1$. Par conséquent (P_n) ne possède aucune sous-suite convergente, donc on déduit du théorème 17.2 que $\mathcal{B}$ n'est pas un compact³ de E . $\square$

L'image d'une partie compacte d'un espace vectoriel normé par une application continue de cet espace vectoriel normé dans un autre est compacte. Ce résultat devient faux si l'on remplace *compacte* par *fermée et bornée*.

17.12. Image non fermée d'une partie fermée et bornée par une application continue.

Nous reprenons l'espace vectoriel normé $(E, \|\cdot\|_\infty)$ et la forme linéaire u sur E de l'exemple 17.10. Nous savons que u est continue et que sa norme linéaire est égale à 1. La boule unité fermée $\mathcal{B}$ de $(E, \|\cdot\|_\infty)$ est connexe — car convexe ; voir ce qui concerne les convexes dans la suite du chapitre —, donc l'image de $\mathcal{B}$ par u est un connexe de $\mathbb{R}$; c'est donc un intervalle. Si f est un vecteur non nul de $\mathcal{B}$, de norme k , $g = (1/k)f$ est de norme 1, donc $|u(g)| < 1$ — voir l'exemple 17.10 —, ce qui montre que $|u(f)| < k \leq 1$; de plus $u(0) = 0$, donc $u(\mathcal{B}) \subset]-1, 1[$.

Soit c un point de $] -1, 1[$. Comme $1 = \sup \{|u(f)| \mid f \in E \text{ et } \|f\| = 1\}$ et que $c < 1$, il existe au moins un vecteur f_0 de E tel que $\|f_0\|_\infty = 1$ et $|u(f_0)| > c$ donc, quitte à remplacer f_0 par $-f_0$, $u(f_0) > c$; de même, on déduit de l'inégalité $-c < 1$ l'existence d'un vecteur g_0 de E tel que $\|g_0\|_\infty = 1$ et $u(g_0) > -c$, donc $u(-g_0) < c$. Par conséquent, c est un point du segment $[u(-g_0), u(f_0)]$; or $u(\mathcal{B})$ est un intervalle, $u(-g_0) \in u(\mathcal{B})$ et $u(f_0) \in u(\mathcal{B})$, donc c appartient à $u(\mathcal{B})$.

En conclusion, $] -1, 1[\subset u(\mathcal{B}) \subset] -1, 1[$, donc $u(\mathcal{B}) =] -1, 1[$: $u(\mathcal{B})$ n'est pas un fermé de $\mathbb{R}$, alors que $\mathcal{B}$ est une partie bornée et fermée de $(E, \|\cdot\|_\infty)$. $\square$

17.13. Image non bornée d'une partie fermée et bornée par une application continue.

Nous reprenons une nouvelle fois l'espace vectoriel normé $(E, \|\cdot\|_\infty)$ de l'exemple 17.10 (pages 324 et 325).

3. Le mathématicien hongrois Frigyes Riesz (1880-1956) a démontré que la boule unité fermée d'un espace vectoriel normé est compacte si, et seulement si, cet espace est de dimension finie ; voir [SKAN], chapitre 7, §2.

Nous considérons l'application :

$$\left| \begin{array}{l} \varphi : E \longrightarrow \mathbb{R} \\ f \longmapsto \varphi(f) = \int_0^{1/2} |f(t) - 1| dt + \int_{1/2}^1 |f(t) + 1| dt. \end{array} \right.$$

En distinguant les deux cas : il existe un point c de $[0, 1/2]$ tel que $f(c) \neq 1$; $f(t) = 1$ pour tout point t de $[0, 1/2]$, on prouve, par le même raisonnement qu'à la fin de l'exemple 17.10, que $\varphi(f) > 0$ pour tout vecteur f de E . Rappelons que si x et y sont des nombres réels, $||x| - |y|| \leq |x - y|$ (inégalités triangulaires). On a, quels que soient les vecteurs f et g de E :

$$\begin{aligned} \varphi(f) - \varphi(g) &= \int_0^{1/2} (|f(t) - 1| - |g(t) - 1|) dt + \int_{1/2}^1 (|f(t) + 1| - |g(t) + 1|) dt \\ \text{donc :} \\ |\varphi(f) - \varphi(g)| &\leq \int_0^{1/2} \underbrace{||f(t) - 1| - |g(t) - 1||}_{\leq |f(t) - g(t)|} dt + \int_{1/2}^1 \underbrace{||f(t) + 1| - |g(t) + 1||}_{\leq |f(t) - g(t)|} dt \\ &\leq \int_0^1 \underbrace{|f(t) - g(t)|}_{\leq \|f - g\|_\infty} dt \leq \|f - g\|_\infty, \end{aligned}$$

ce qui montre que φ est une application lipschitzienne de $(E, \|\cdot\|_\infty)$ dans $(\mathbb{R}, |\cdot|)$. Par suite, φ est continue sur E ; or φ ne s'annule pas sur E , donc son inverse $\psi = 1/\varphi$ est une application continue de E dans $\mathbb{R}$.

Nous utilisons de nouveau la suite $(f_n)_{n \geq 2}$ de l'exemple 17.10. Pour tout entier $n \geq 2$, f_n est décroissante sur $[0, 1]$, donc $-1 = f_n(-1) \leq f_n(t) \leq f_n(1) = 1$ pour tout point t de $[0, 1]$ et, en utilisant les résultats des calculs effectués pour obtenir la valeur de $\varphi(f_n)$ dans l'exemple 17.10, il vient :

$$\begin{aligned} \varphi(f_n) &= \int_0^{1/2} |f_n(t) - 1| dt + \int_{1/2}^1 |f_n(t) + 1| dt \\ &= \int_0^{\frac{1}{2} - \frac{1}{n}} 0 dt + \int_{\frac{1}{2} - \frac{1}{n}}^{\frac{1}{2}} (1 - f_n(t)) dt + \int_{\frac{1}{2}}^{\frac{1}{2} + \frac{1}{n}} (1 + f_n(t)) dt + \int_{\frac{1}{2} + \frac{1}{n}}^1 0 dt \\ &= 0 + \frac{1}{n} - \frac{n}{2} \frac{1}{n^2} + \frac{1}{n} - \frac{n}{2} \frac{1}{n^2} + 0 = \frac{1}{n}, \end{aligned}$$

donc $\psi(f_n) = n$. Or $\|f_n\|_\infty = 1$ pour tout entier $n \geq 2$, donc l'image par ψ de la boule unité fermée $\mathcal{B}$ de $(E, \|\cdot\|_\infty)$ n'est pas bornée dans $\mathbb{R}$, alors que $\mathcal{B}$ est fermée et bornée dans E . $\square$

■ Parties convexes d'un espace vectoriel normé

DÉFINITION 17.3. — Si a et b sont des vecteurs d'un espace vectoriel réel, le segment a, b est l'ensemble $[a, b] = \{(1 - t)a + tb \mid t \in [0, 1]\}$ et, si $a \neq b$, l'intervalle ouvert a, b est l'ensemble $]a, b[= [a, b] \setminus \{a, b\} = \{(1 - t)a + tb \mid t \in]0, 1[\}$.

DÉFINITION 17.4. — Une partie C d'un espace vectoriel réel E est convexe, ou C est un convexe de E , si, quels que soient les vecteurs a et b appartenant à C , le segment $[a, b]$ est inclus dans C .

En particulier, les sous-espaces affines d'un espace vectoriel réel E sont convexes.

Dans un espace vectoriel normé, les parties convexes sont connexes (définitions 15.14 et 15.15, page 295) car elles sont connexes par arcs (définition 15.19, page 300). Dans $\mathbb{R}$, les intervalles sont les convexes et aussi les connexes.

17.14. Partie connexe qui n'est pas convexe.

Nous munissons l'espace vectoriel réel $\mathbb{C}$ du module et nous considérons le cercle unité $\mathbb{U} = \{z \in \mathbb{C} \mid |z| = 1\}$ de $\mathbb{C}$. L'application $f : t \mapsto f(t) = e^{it}$ de $\mathbb{R}$ dans $\mathbb{C}$ est continue, $\mathbb{R}$ est connexe et $\mathbb{U} = f(\mathbb{R})$, donc $\mathbb{U}$ est connexe (théorème 15.7). Les complexes $a = -1$ et $b = 1$ sont des vecteurs de $\mathbb{U}$ et $0 = (1 - (1/2))a + (1/2)b$ appartient au segment $[a, b]$ mais pas à $\mathbb{U}$, donc $\mathbb{U}$ n'est pas un convexe de $\mathbb{C}$. $\square$

Dans un espace hilbertien, et à plus forte raison dans un espace euclidien — le rappel des définitions de ces espaces figure dans la suite du présent chapitre —, tout convexe fermé non vide contient un élément de norme minimale et un seul. Montrons que ceci devient faux dans un espace vectoriel normé quelconque, même si on le suppose complet — un espace vectoriel normé est complet si, muni de la distance associée à la norme, c'est un espace métrique complet.

17.15. Partie convexe fermée sans élément de norme minimale.

Nous reprenons l'espace vectoriel normé $(E, \|\cdot\|_\infty)$ des exemples 17.10, 17.11 et 17.13, et la forme linéaire continue u sur E définie dans l'exemple 17.10. L'espace vectoriel normé $(E, \|\cdot\|_\infty)$ est complet — c'est un résultat classique.

L'ensemble $C = \{f \in E \mid u(f) = 1\}$ est convexe — c'est un sous-espace affine — et fermé — car $C = u^{-1}(\{1\})$. L'image par u de la boule unité fermée $\mathcal{B}$ de E est égale à $] -1, 1[$ (exemple 17.12) ; par suite, pour tout $f \in E$, $\|f\|_\infty \leq 1$ entraîne $-1 < u(f) < 1$, d'où l'on déduit que $\|f\|_\infty > 1$ pour tout vecteur f de C .

Nous utilisons de nouveau la suite $(f_n)_{n \geq 2}$ de vecteurs de E de l'exemple 17.10. Nous avons vu que, pour tout entier $n \geq 2$, $\|f_n\|_\infty = 1$ et $u(f_n) = 1 - (1/n)$. Nous introduisons la suite $(h_n)_{n \geq 2}$ de terme général :

$$h_n = \frac{n}{n-1} f_n.$$

Pour tout entier $n \geq 2$, $u(h_n) = 1$, donc h_n appartient à C , et $\|h_n\|_\infty = n/(n-1)$. La suite $(\|h_n\|_\infty)$ converge donc vers 1, ce qui montre que C ne contient aucun élément de norme minimale. $\square$

17.16. Partie convexe fermée possédant une infinité d'éléments de norme minimale.

Nous munissons l'espace vectoriel réel $\mathbb{R}^2$ de la norme « max » :

$$\|\cdot\|_\infty : x = (x_1, x_2) \mapsto \|x\|_\infty = \text{Max}(|x_1|, |x_2|).$$

Nous notons D la droite affine de $\mathbb{R}^2$ d'équation $x_1 = 1$, c'est-à-dire l'ensemble :

$$D = \{x = (x_1, x_2) \mid x_1 = 1\} = \{(1, t) \mid t \in \mathbb{R}\}.$$

L'application $p : x = (x_1, x_2) \mapsto p(x) = x_1$ de $\mathbb{R}^2$ dans $\mathbb{R}$ est continue et on a $D = p^{-1}(\{1\})$, donc D est fermé. Pour tout réel t , $\|(1, t)\|_\infty = \text{Max}(1, |t|) \geq 1$. De plus, pour tout $t \in [-1, 1]$, $\|(1, t)\|_\infty = 1$, donc les vecteurs $(1, t)$ pour t parcourant le segment $[-1, 1]$ sont tous de norme minimale dans D . $\square$

■ Points internes à une partie

Rappelons que si E est un espace topologique et A une partie de E , un point a de E est intérieur à A si A est un voisinage de a dans E ; en particulier, un point a intérieur à A appartient à A .

DÉFINITION 17.5. — Un point a d'une partie A d'un espace vectoriel normé E est interne à A si l'intersection avec A de toute droite affine de E passant par a contient un intervalle ouvert centré en a , ce qui équivaut à l'assertion suivante : pour tout vecteur non nul u de E , il existe un nombre réel $\varepsilon > 0$ tel que, pour tout $t \in]-\varepsilon, \varepsilon[$, $a + tu$ appartient à A .

Si a est un point intérieur à une partie A d'un espace vectoriel normé $(E, \|\cdot\|)$, il existe un réel $\varepsilon_0 > 0$ tel que la boule ouverte de centre a et de rayon ε_0 est incluse dans A , donc, pour tout vecteur non nul u de E , le réel $\varepsilon = \varepsilon_0/\|u\|$ est strictement positif et, pour tout $t \in]-\varepsilon, \varepsilon[$, $a + tu$ appartient à A , ce qui montre que le point a est interne à A . La réciproque est vraie si A est convexe, mais fausse en général.

17.17. Point interne qui n'est pas un point intérieur.

Nous munissons l'espace vectoriel réel $\mathbb{R}^2$ de la norme euclidienne :

$$\begin{cases} \|\cdot\| : \mathbb{R}^2 \longrightarrow \mathbb{R}_+ \\ z = (x, y) \longmapsto \|z\| = \sqrt{x^2 + y^2} \end{cases}$$

et nous posons $B = \{(x, y) \mid 0 < x < 1, 0 < y < 1 \text{ et } x^4 < y < x^2\}$ et $A = \mathbb{C}_{\mathbb{R}^2} B$. Pour tout réel ε tel que $0 < \varepsilon < 1$, $0 < \varepsilon^4 < \varepsilon^3 < \varepsilon^2 < 1$ donc le vecteur $(\varepsilon, \varepsilon^3)$ appartient à B . Par conséquent, toute boule ouverte de centre $0 = (0, 0)$ rencontre le complémentaire B de A dans $\mathbb{R}^2$, donc le point $0 = (0, 0)$ n'est pas intérieur à A .

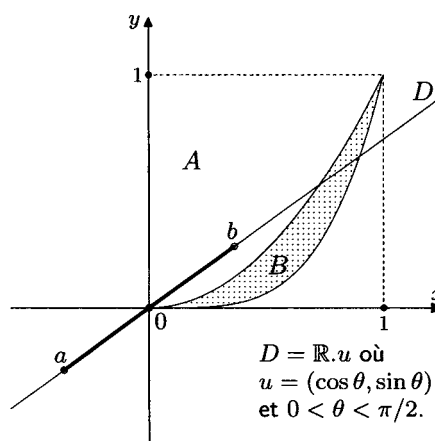
Soit D une droite affine de $\mathbb{R}^2$ passant par le point $0 = (0, 0)$; c'est la droite vectorielle $D = \mathbb{R}u = \{\lambda u \mid \lambda \in \mathbb{R}\}$ pour un vecteur unitaire u d'ordonnée positive ou nulle, vecteur qui s'écrit donc $u = (\cos \theta, \sin \theta)$ où $\theta \in]0, \pi]$.

Si $\pi/2 \leq \theta \leq \pi$, D est incluse dans A , car, pour tout vecteur $z = (x, y)$ de D , $(x \geq 0 \text{ et } y \leq 0)$ ou $(x \leq 0 \text{ et } y \geq 0)$ donc $D \cap B$ est vide — en fait, $D \cap B$ est vide aussi pour $\pi/4 < \theta < \pi/2$ (voir le dessin ci-contre).

Il reste à examiner le cas où $0 < \theta < \pi/2$. Dans ce cas, $\sin \theta > 0$ et $\cos \theta > 0$.

Soit $z = (x, y)$ un vecteur de D . Il existe un réel r tel que $z = ru = (r \cos \theta, r \sin \theta)$. Si $r \leq 0$, $x = r \cos \theta \leq 0$, donc z appartient à A ; sinon, $r > 0$ donc, si $z \in B$, on a $r^4 \cos^4 \theta < r \sin \theta < r^2 \cos^2 \theta$, ce qui équivaut à $r^3 \cos^4 \theta < \sin \theta < r \cos^2 \theta$, d'où l'on déduit que si $r \cos^2 \theta \leq \sin \theta$, alors $z \notin B$ donc z appartient à A . Posons :

$$r_0 = \frac{\sin \theta}{2 \cos^2 \theta} > 0, \quad a = -r_0 u \quad \text{et} \quad b = r_0 u.$$



Pour tout nombre réel r tel que $0 < r < r_0$, $r \cos^2 \theta < r_0 \cos^2 \theta < \sin \theta$ donc z appartient à A . Par suite, l'intervalle ouvert $]a, b[$, centré en $0 = (0, 0)$ (voir le dessin de la page précédente), est inclus dans $D \cap A$.

En conclusion, le point $0 = (0, 0)$ est interne à l'ensemble A . $\square$

■ Isométries

DÉFINITION 17.6. — Une application f d'un espace métrique (E, d_E) dans un espace métrique (F, d_F) est une isométrie si, quels que soient les points x et y de E :

$$d_F(f(x), f(y)) = d_E(x, y).$$

Une isométrie d'un espace métrique compact E dans lui-même est une bijection de E sur E . Si l'espace métrique E n'est pas compact, une isométrie de E dans E est injective, mais n'est pas forcément une bijection de E sur E . Nous allons donner des exemples illustrant cette affirmation, à l'aide d'un espace métrique obtenu en munissant une partie non compacte d'un espace vectoriel normé de la distance induite par la distance associée à la norme, et d'une isométrie de cet espace métrique dans lui-même ; pour un espace vectoriel normé de dimension finie, la partie sera fermée ou bornée, et dans le cas d'une dimension infinie, elle sera bornée et fermée.

17.18. Isométrie d'une partie fermée F de $\mathbb{R}^2$ dans elle-même qui n'est pas une bijection de F sur F .

Nous munissons l'espace vectoriel réel $\mathbb{R}^2$ de la norme euclidienne $\|\cdot\|$ (voir l'exemple précédent 17.17), nous introduisons le demi-plan :

$$F = \mathbb{R} \times [0, +\infty[= \{(x, y) \mid x, y \in \mathbb{R} \text{ et } y \geq 0\}$$

de $\mathbb{R}^2$ et nous considérons l'espace métrique obtenu en munissant F de la distance induite par la distance associée à la norme euclidienne de $\mathbb{R}^2$.

Pour tout vecteur $c = (a, b)$ tel que $a < 0$, la boule ouverte de centre c et de rayon $-a$ est incluse dans $\Omega = \mathbb{R} \times]-\infty, 0[$, donc Ω est un ouvert. Or le complémentaire de F dans $\mathbb{R}^2$ est Ω , donc F est une partie fermée de $\mathbb{R}^2$.

Pour tout vecteur $z = (x, y)$ de F , $y + 1 > y \geq 0$ donc $(x, y + 1)$ appartient à F , ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : F \longrightarrow F \\ z = (x, y) \longmapsto f(z) = (x, y + 1) \end{array} \right.$$

dont il est clair que c'est une isométrie. Cependant un vecteur $z = (x, y)$ de $\mathbb{R}^2$ tel que $0 \leq y < 1$ — il en existe — appartient à F mais n'admet pas d'antécédent par f , donc f n'est pas une bijection de F sur F . $\square$

17.19. Isométrie d'une partie bornée B de $\mathbb{C}$ dans elle-même qui n'est pas une bijection de B sur B .

Nous munissons l'espace vectoriel réel $\mathbb{C}$ du module. Posons $B = \{e^{in} \mid n \in \mathbb{N}\}$. Tous les éléments de B sont de module 1, donc B est une partie bornée de $\mathbb{C}$. Pour tout entier naturel n , $e^{ie^{in}} = e^{i(n+1)}$, donc le produit $e^{ie^{in}}$ appartient à B ,

ce qui justifie l'existence de l'application :

$$\left| \begin{array}{l} f : B \longrightarrow B \\ z \longmapsto f(z) = e^i z, \end{array} \right.$$

restriction à B de la rotation de centre 0 et d'angle 1 dans l'espace euclidien $\mathbb{C}$. Clairement, f est une isométrie de B . Comme $1 = e^{i0}$, 1 appartient à B . Si le nombre complexe 1 admettait un antécédent par f , il existerait un entier naturel n tel que $1 = e^{ie^{in}} = e^{i(n+1)}$, donc un entier relatif k tel que $n+1 = 2k\pi$, et par conséquent π serait rationnel. Ainsi, l'élément 1 de B n'admet pas d'antécédent par f , donc f n'est pas une bijection de B sur B . $\square$

17.20. Isométrie d'un fermé borné $\mathcal{B}$ d'un espace vectoriel normé dans lui-même qui n'est pas une bijection de $\mathcal{B}$ sur $\mathcal{B}$.

Nous munissons l'espace vectoriel réel $E = \mathbb{R}[X]$ de la norme :

$$\|\bullet\| : P = \sum_{k=0}^{+\infty} a_k X^k \mapsto \|P\| = \max_{k \in \mathbb{N}} |a_k| \quad \left(= \max_{0 \leq k \leq p} |a_k| \text{ si } p \in \mathbb{N} \text{ et } \deg(P) \leq p \right)$$

(voir les exemples 17.7 page 322, et 17.11 page 326). L'application :

$$\left| \begin{array}{l} g : E \longrightarrow E \\ P \longmapsto g(P) = XP \end{array} \right.$$

est linéaire et, pour tout vecteur $P = \sum_{k=0}^{+\infty} a_k X^k$ de E , $g(P) = \sum_{\ell=1}^{+\infty} a_{\ell-1} X^\ell$. Si $P = \sum_{k=0}^{+\infty} a_k X^k$ et $Q = \sum_{k=0}^{+\infty} b_k X^k$ sont des vecteurs de E , on a :

$$P - Q = \sum_{k=0}^{+\infty} (a_k - b_k) X^k \text{ et } g(P) - g(Q) = g(P - Q) = \sum_{\ell=1}^{+\infty} (a_{\ell-1} - b_{\ell-1}) X^\ell$$

donc :

$$(1) \quad \|g(P) - g(Q)\| = \max_{\ell \in \mathbb{N}^*} |a_{\ell-1} - b_{\ell-1}| = \max_{k \in \mathbb{N}} |a_k - b_k| = \|P - Q\|.$$

Pour tout $n \in \mathbb{N}$, $\|X^n\| = 1$, donc $\mathcal{B} = \{X^n \mid n \in \mathbb{N}\}$ est une partie bornée de E .

Si $P, Q \in \mathcal{B}$ et si $P \neq Q$, il existe $p, q \in \mathbb{N}$ tels que $P = X^p$, $Q = X^q$ et $p \neq q$, donc $\|P - Q\| = \|X^p - X^q\| = 1$. Soit (P_n) une suite de vecteurs de $\mathcal{B}$ qui converge dans E vers un polynôme L . C'est une suite de Cauchy, donc il existe $N \in \mathbb{N}$ tel que, quels que soient les entiers $k \geq N$ et $\ell \geq N$, $\|P_k - P_\ell\| < 1$, ce qui, puisque $P_k \in \mathcal{B}$ et $P_\ell \in \mathcal{B}$, montre que $P_k = P_\ell$. Ainsi, $P_n = P_N$ pour tout $n \geq N$, donc $L = P_N$; par conséquent, $L \in \mathcal{B}$. Il en résulte que $\mathcal{B}$ est une partie fermée de E .

Pour tout $n \in \mathbb{N}$, $g(X^n) = X \times X^n = X^{n+1}$ appartient à $\mathcal{B}$, donc la restriction f de g à $\mathcal{B}$ est une application de $\mathcal{B}$ dans $\mathcal{B}$, et on déduit de (1) que f est une isométrie de $\mathcal{B}$.

Si X^0 admettait un antécédent par f , il existerait un entier naturel n tel que $X^0 = f(X^n) = X^{n+1}$, en contradiction avec $n+1 \neq 0$. Le vecteur X^0 de $\mathcal{B}$ n'admettant pas d'antécédent, f n'est pas une bijection de $\mathcal{B}$ sur $\mathcal{B}$. $\square$

DÉFINITION 17.7. — Une partie non vide A d'un espace métrique E est dédoublable si A est la réunion disjointe de deux parties A_1 et A_2 de E isométriques à A dans E — ce qui signifie qu'il existe des isométries g_1 et g_2 de l'espace métrique E telles que $A_1 = g_1(A)$ et $A_2 = g_2(A)$.

On démontre qu'aucune partie non vide de l'espace métrique $\mathbb{R}$ n'est dédoublable et qu'aucune partie bornée et non vide d'un plan euclidien n'est dédoublable, ce qui n'est pas le cas si cette partie n'est pas bornée.

17.21. Partie dédoublable d'un plan euclidien.

Nous munissons $\mathbb{C}$ de sa structure canonique de plan euclidien et nous choisissons un réel θ tel que le nombre complexe $u = e^{i\theta}$ est transcendant⁴.

Notons $\mathcal{P}$ l'ensemble des « polynômes à coefficients dans $\mathbb{N}$ », ce qui signifie que $\mathcal{P} = \{P = \sum_{n=0}^{+\infty} a_n X^n \mid a_n \in \mathbb{N} \text{ pour tout } n \in \mathbb{N}\}$, posons $\mathcal{A} = \{P(u) \mid P \in \mathcal{P}\}$ et notons ρ la rotation de centre 0 et d'angle θ et τ la translation de vecteur 1. Alors ρ et τ sont des isométries du plan euclidien $\mathbb{C}$ et on a, pour tout nombre complexe z , $\rho(z) = uz$ et $\tau(z) = z + 1$.

Soit z un élément de $\mathcal{A}$. On a $z = P(u)$ où $P \in \mathcal{P}$, donc $P = \sum_{n=0}^{+\infty} a_n X^n$, $(a_n)_{n \geq 0}$ étant une suite d'entiers naturels nulle à partir d'un certain rang. Si $a_0 = 0$, le polynôme $Q = \sum_{n=1}^{+\infty} a_n X^{n-1}$ appartient à $\mathcal{P}$, donc $Q(u) \in \mathcal{A}$, et comme $P = XQ$, $z = P(u) = uQ(u) = \rho(Q(u))$, d'où l'on déduit que z est un élément de $\mathcal{A}_1$. Si $a_0 \neq 0$, alors $a_0 \geq 1$ donc $R = (a_0 - 1) + \sum_{n=1}^{+\infty} a_n X^n$ appartient à $\mathcal{P}$, ce qui montre que $R(u) \in \mathcal{A}$, et comme $P = 1 + R$, $z = P(u) = 1 + R(u) = \tau(R(u))$ donc z est un élément de $\mathcal{A}_2$. Nous avons ainsi établi que $\mathcal{A} \subset \mathcal{A}_1 \cup \mathcal{A}_2$.

Si $z \in \mathcal{A}_1$, on a $z = \rho(v) = uv$ où $v \in \mathcal{A}$, d'où l'existence d'un élément P de $\mathcal{P}$ tel que $v = P(u)$, et comme $Q = XP$ appartient à $\mathcal{P}$ et que $z = uP(u) = Q(u)$, v appartient à $\mathcal{A}$. Par suite $\mathcal{A}_1$ est inclus dans $\mathcal{A}$. On prouve de la même façon que $\mathcal{A}_2$ est inclus dans $\mathcal{A}$, ce qui montre que $\mathcal{A}_1 \cup \mathcal{A}_2 \subset \mathcal{A}$. Par conséquent $\mathcal{A} = \mathcal{A}_1 \cup \mathcal{A}_2$.

Supposons que $\mathcal{A}_1 \cap \mathcal{A}_2 \neq \emptyset$. Nous choisissons un point z de $\mathcal{A}_1 \cap \mathcal{A}_2$. Il existe des éléments $P = \sum_{n=0}^{+\infty} a_n X^n$ et $Q = \sum_{n=0}^{+\infty} b_n X^n$ de $\mathcal{P}$ tels que $z = \rho(P(u)) = uP(u)$ et $z = \tau(Q(u)) = 1 + Q(u)$. Alors $R = XP - Q - 1$ appartient à $\mathcal{P}$ et $R(u) = 0$. Comme le polynôme R est à coefficients entiers et que u est transcendant, R est le polynôme nul. Par suite $Q = -1 + XP = (-1) + \sum_{n=1}^{+\infty} a_{n-1} X^n$, donc $b_0 = -1$, en contradiction avec l'appartenance de b_0 à $\mathbb{N}$.

En conclusion, $\mathcal{A}_1 \cap \mathcal{A}_2 = \emptyset$, donc $\mathcal{A}$ est dédoublable. $\square$

■ Espaces vectoriels euclidiens et hilbertiens

DÉFINITION 17.8. — Un produit scalaire $(\bullet | \bullet)$ sur un espace vectoriel réel E est une forme bilinéaire symétrique :

$$\left| \begin{array}{l} (\bullet | \bullet) : E \times E \longrightarrow \mathbb{R} \\ (x, y) \longmapsto (x | y) \end{array} \right.$$

définie positive sur E , ce qui signifie que $(\bullet | \bullet)$ est une application bilinéaire symétrique de $E \times E$ dans $\mathbb{R}$ et que $(x | x) > 0$ pour tout vecteur non nul x de l'espace vectoriel E .

4. A l'occasion de sa démonstration de la transcendance de π en 1882, le mathématicien allemand Carl Lindemann (1852-1939) amorce la preuve du théorème suivant : Si α est un nombre réel ou complexe algébrique et si α est différent de zéro, e^α est transcendant. En choisissant un réel algébrique non nul θ — il y en a autant que l'on veut ! —, le nombre complexe $i\theta$ est algébrique et différent de zéro, donc $u = e^{i\theta}$ est transcendant.

DÉFINITION 17.9. — a) Un espace préhilbertien réel est un espace vectoriel réel E muni d'un produit scalaire $(\cdot | \cdot)$, et alors l'application :

$$\left| \begin{array}{l} \|\cdot\| : E \longrightarrow \mathbb{R}_+ \\ x \longmapsto \|x\| = \sqrt{(x | x)} \end{array} \right.$$

est une norme sur E , appelée sa norme euclidienne, et la distance associée à $\|\cdot\|$ s'appelle la distance euclidienne de E .

b) Un espace hilbertien réel est un espace préhilbertien réel complet pour sa distance euclidienne.

On définit de même un espace préhilbertien complexe à l'aide d'une forme sesquilinéaire définie positive⁵, également notée $(\cdot | \cdot)$.

DÉFINITION 17.10. — Si x et y sont des vecteurs d'un espace préhilbertien, x est orthogonal à y si $(x | y) = 0$.

La relation binaire d'orthogonalité sur un espace préhilbertien E étant symétrique, on dit que des vecteurs x et y de E sont orthogonaux pour exprimer que x (ou y) est orthogonal à y (ou x).

THÉORÈME 17.4. — **Théorème de Pythagore**⁶.

Si des vecteurs x et y d'un espace préhilbertien sont orthogonaux, alors :

$$\|x\|^2 + \|y\|^2 = \|x + y\|^2.$$

La réciproque est vraie dans un espace préhilbertien réel, mais devient fausse dans le cas complexe.

17.22. Vecteurs u et v tels que $\|u\|^2 + \|v\|^2 = \|u + v\|^2$ alors qu'ils ne sont pas orthogonaux.

Nous considérons l'espace préhilbertien complexe obtenu en munissant l'espace vectoriel complexe $\mathbb{C}^2$ du produit scalaire complexe canonique $(\cdot | \cdot)$ défini, si $x = (x_1, x_2)$ et $y = (y_1, y_2)$ sont des vecteurs de $\mathbb{C}^2$, par : $(x | y) = \bar{x}_1 y_1 + \bar{x}_2 y_2$. On a donc, pour tout vecteur $x = (x_1, x_2)$ de $\mathbb{C}^2$, $\|x\|^2 = |x_1|^2 + |x_2|^2$.

Posons $u = (1, 1)$ et $v = (i, i)$. Alors $u + v = (1 + i, 1 + i)$, $\|u\|^2 = \|v\|^2 = 1 + 1 = 2$ et $\|u + v\|^2 = |1 + i|^2 + |1 + i|^2 = 2 + 2 = 4$, donc $\|u\|^2 + \|v\|^2 = \|u + v\|^2$.

Or $(u | v) = \bar{1} \times i + \bar{1} \times i = 1 \times i + 1 \times i = 2i \neq 0$, donc les vecteurs u et v ne sont pas orthogonaux⁷. $\square$

DÉFINITION 17.11. — Soit E un espace préhilbertien et A une partie de E .

a) Un vecteur x de E est orthogonal à A si x est orthogonal à tous les vecteurs de A , c'est-à-dire si $(x | y) = 0$ pour tout vecteur y de A .

b) L'orthogonal $A^\perp$ de A est l'ensemble des vecteurs de E orthogonaux à A .

5. Voir [SCHW], chapitre 6, §1.

6. Du nom d'un philosophe et mathématicien grec du VI^e siècle av. J.C.

7. Pour des vecteurs x et y d'un espace préhilbertien complexe, l'égalité $\|x\|^2 + \|y\|^2 = \|x + y\|^2$ est vérifiée si, et seulement si, la partie réelle de $(x | y)$ est nulle.

Si A est une partie d'un espace préhilbertien E , son orthogonal $A^\perp$ est un sous-espace vectoriel de E .

DÉFINITION 17.12. — Des sous-espaces vectoriels F et G d'un espace préhilbertien sont orthogonaux si tout vecteur de F est orthogonal à tout vecteur de G , c'est-à-dire si $(x|y) = 0$ quels que soient les vecteurs x de F et y de G .

THÉORÈME 17.5. — Si E est un espace préhilbertien et F un sous-espace vectoriel de E , les sous-espaces vectoriels F et $F^\perp$ sont orthogonaux et $F \subset (F^\perp)^\perp$, et si de plus F est de dimension finie, alors $E = F \oplus F^\perp$ et $F = (F^\perp)^\perp$.

On voit donc que si F est un sous-espace vectoriel de dimension finie d'un espace préhilbertien E , son orthogonal $F^\perp$ est un supplémentaire de F dans E , que l'on appelle son supplémentaire orthogonal, ce qui permet de définir la projection orthogonale sur F et la symétrie orthogonale par rapport à F .

17.23. Orthogonal réduit à $\{0\}$ d'un sous-espace vectoriel strict.

Nous considérons l'espace vectoriel réel $E = \mathcal{C}([0, 1], \mathbb{R})$ des applications continues de $[0, 1]$ dans $\mathbb{R}$ et l'application :

$$\left| \begin{array}{l} (\bullet|\bullet) : E \times E \longrightarrow \mathbb{R} \\ (f, g) \longmapsto (f|g) = \int_0^1 f(t)g(t) dt. \end{array} \right.$$

Les propriétés de l'intégrale des fonctions à valeurs réelles continues sur un segment montrent que $(\bullet|\bullet)$ est un produit scalaire sur E .

Nous posons $H = \{f \in E \mid f(0) = 0\}$. L'application $\varphi : f \mapsto \varphi(f) = f(0)$ de E dans $\mathbb{R}$ est une forme linéaire non nulle sur E et $H = \text{Ker}(\varphi)$, donc H est un hyperplan vectoriel de E ; par suite H est un sous-espace vectoriel strict de E .

Soit g un vecteur de $H^\perp$. On définit l'application continue $h : t \mapsto h(t) = tg(t)$ de $[0, 1]$ dans $\mathbb{R}$. Alors h est un vecteur de H et g un vecteur de $H^\perp$, donc :

$$0 = (h|g) = \int_0^1 tg(t) \times g(t) dt = \int_0^1 t(g(t))^2 dt.$$

La fonction $t \mapsto t(g(t))^2$ étant continue et positive sur $[0, 1]$ et d'intégrale nulle, on a $t(g(t))^2 = 0$ pour tout point t de $[0, 1]$. On en déduit que $g(t) = 0$ pour tout $t \in]0, 1]$; la continuité de g en 0 et le passage à la limite quand t tend vers 0 à droite montrent que $g(0) = 0$. Finalement, g est le vecteur zéro de E .

En conclusion, $H^\perp = \{0\}$. $\square$

Dans l'exemple précédent 17.23, le sous-espace vectoriel H de l'espace préhilbertien E ne possède pas de supplémentaire orthogonal; en particulier, on ne peut pas définir la projection orthogonale sur H .

17.24. Sous-espace vectoriel F différent de $(F^\perp)^\perp$.

Nous reprenons l'espace préhilbertien E de l'exemple précédent 17.23 et son hyperplan vectoriel H , et nous posons $F = H$. Alors $F^\perp = \{0\}$ donc :

$$(F^\perp)^\perp = E \neq F. \quad \square$$

Si F et G sont des sous-espaces vectoriels d'un espace préhilbertien, $F^\perp + G^\perp$ est inclus dans l'orthogonal $(F \cap G)^\perp$ de $F \cap G$. L'égalité est vraie en dimension finie, mais devient, en général, fausse en dimension infinie⁸.

17.25. Sous-espaces vectoriels F et G tels que $F^\perp + G^\perp \subsetneq (F \cap G)^\perp$.

Nous reprenons l'espace préhilbertien réel $(E, (\cdot | \cdot))$ de l'exemple 17.23, nous posons $F = \{f \in E \mid f(0) = 0\}$ — c'est l'hyperplan H de l'exemple 17.23 — et nous notons G le sous-espace vectoriel de E dont les éléments sont les applications constantes de $[0, 1]$ dans $\mathbb{R}$. Si un vecteur f de E appartient à $F \cap G$, alors, pour tout point t de $[0, 1]$, $f(t) = f(0)$ puisque l'application f est constante, et $f(0) = 0$, donc f est le vecteur zéro de E . Par conséquent $F \cap G = \{0\}$, donc $(F \cap G)^\perp = E$. Comme G est le sous-espace vectoriel de E engendré par l'application constante $\mathbf{1} : t \mapsto \mathbf{1}(t) = 1$ de $[0, 1]$ dans $\mathbb{R}$, un vecteur f de E est orthogonal à G si, et seulement si, il est orthogonal à $\mathbf{1}$, ce qui équivaut à $0 = (f | \mathbf{1}) = \int_0^1 f(t) dt$. Par suite :

$$G^\perp = \left\{ f \in E \mid \int_0^1 f(t) dt = 0 \right\}.$$

Nous avons prouvé dans l'exemple 17.23 que $F^\perp = \{0\}$, d'où l'égalité :

$$F^\perp + G^\perp = G^\perp.$$

Si f_0 est une application continue de $[0, 1]$ dans $\mathbb{R}$, si $f_0(t) \geq 0$ pour tout point t de $[0, 1]$ et s'il existe $t_0 \in [0, 1]$ tel que $f_0(t_0) \neq 0$ — par exemple $f_0 : t \mapsto f_0(t) = t^2$ —, l'intégrale de f_0 sur $[0, 1]$ est strictement positive, donc le vecteur f_0 de E n'appartient pas à $G^\perp = F^\perp + G^\perp$. En conclusion, $F^\perp + G^\perp \subsetneq E = (F \cap G)^\perp$. $\square$

Nous généralisons l'exemple 17.24 à une famille infinie de sous-espaces vectoriels.

17.26. Suite $(F_n)_{n \in \mathbb{N}}$ de sous-espaces vectoriels telle que :

$$\sum_{n \in \mathbb{N}} F_n^\perp \subsetneq \left(\bigcap_{n \in \mathbb{N}} F_n \right)^\perp.$$

Nous utilisons de nouveau l'espace préhilbertien réel $(E, (\cdot | \cdot))$ de l'exemple 17.23. L'ensemble $Q = \mathbb{Q} \cap [0, 1]$ est dénombrable. Nous choisissons une bijection ψ de $\mathbb{N}$ sur Q . En posant $a_n = \psi(n)$ pour tout entier naturel n , on obtient une suite $(a_n)_{n \in \mathbb{N}}$ d'éléments deux à deux distincts de Q tels que $Q = \{a_n \mid n \in \mathbb{N}\}$. Pour tout entier naturel n , l'application $\varphi_n : f \mapsto \varphi_n(f) = f(a_n)$ de E dans $\mathbb{R}$ est une forme linéaire non nulle sur E , donc l'ensemble :

$$F_n = \{f \in E \mid f(a_n) = 0\} = \text{Ker}(\varphi_n)$$

est un hyperplan vectoriel de E . Nous construisons ainsi une suite $(F_n)_{n \in \mathbb{N}}$ de sous-espaces vectoriels de E .

Soit $n \in \mathbb{N}$. Si g est un vecteur de E et si g est orthogonal à F_n , alors, en utilisant l'application continue $h : t \mapsto h(t) = |t - a_n| g(t)$ de $[0, 1]$ dans $\mathbb{R}$, qui est un vecteur de F_n , l'égalité $0 = (h | g)$ et la continuité de g en a_n , on prouve comme dans l'exemple 17.23 que g est le vecteur zéro de E . Par conséquent $F_n^\perp = \{0\}$.

8. Par contre, l'égalité $(F + G)^\perp = F^\perp \cap G^\perp$ est toujours vraie.

La somme $\sum_{n \in \mathbb{N}} F_n^\perp$ de la famille $(F_n^\perp)_{n \in \mathbb{N}}$ est donc égale à $\{0\}$. Nous posons :

$$G = \bigcap_{n \in \mathbb{N}} F_n.$$

Soit f un vecteur de G . Alors l'application f de $[0, 1]$ dans $\mathbb{R}$ s'annule en tous les points de Q . Comme Q est dense dans $[0, 1]$ et que f est continue sur $[0, 1]$, $f(t) = 0$ pour tout point t de $[0, 1]$, donc f est le vecteur zéro de E . Il en résulte que $G = \{0\}$, ce qui montre que $G^\perp = E$. Finalement :

$$\sum_{n \in \mathbb{N}} F_n^\perp = \{0\} \subsetneq E = \left(\bigcap_{n \in \mathbb{N}} F_n \right)^\perp. \quad \square$$

DÉFINITION 17.13. — Un espace vectoriel euclidien est un espace préhilbertien réel de dimension finie non nulle.

Un espace vectoriel normé de dimension finie étant complet, un espace vectoriel euclidien est un espace hilbertien réel.

DÉFINITION 17.14. — Un endomorphisme orthogonal d'un espace vectoriel euclidien E est un endomorphisme f de E qui conserve la norme euclidienne, ce qui signifie que $\|f(x)\| = \|x\|$ pour tout vecteur x de E .

Soit E un espace vectoriel euclidien. Les expressions du produit scalaire à l'aide de la norme montrent qu'un endomorphisme orthogonal de E est un endomorphisme f de E qui conserve le produit scalaire, ce qui signifie que $(f(x) | f(y)) = (x | y)$ quels que soient les vecteurs x et y de E .

En fait, une *application* de E dans E — que l'on ne suppose pas *a priori* linéaire — qui conserve le produit scalaire est automatiquement linéaire, donc c'est un endomorphisme orthogonal de E . Cependant la conservation de la norme euclidienne n'entraîne pas la linéarité.

17.27. Application qui conserve la norme euclidienne mais qui n'est pas linéaire.

Nous munissons l'espace vectoriel réel $\mathbb{R}^2$ de son produit scalaire canonique, défini, si $x = (x_1, x_2)$ et $y = (y_1, y_2)$ sont des vecteurs de $\mathbb{R}^2$ par : $(x | y) = x_1 y_1 + x_2 y_2$. L'espace $\mathbb{R}^2$ devient ainsi un espace vectoriel euclidien de dimension 2 et, pour tout vecteur $x = (x_1, x_2)$ de $\mathbb{R}^2$, $\|x\| = \sqrt{(x_1)^2 + (x_2)^2}$.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R}^2 \\ x = (x_1, x_2) \longmapsto f(x) = \begin{cases} x & \text{si } x_2 \neq 0, \\ -x & \text{si } x_2 = 0. \end{cases} \end{array} \right.$$

Pour tout vecteur x de E , $f(x) = x$ ou $-x$, donc $\|f(x)\| = \|x\|$. Il en résulte que l'application f conserve la norme euclidienne. Or :

$$f((1, 1) + (1, -1)) = f((2, 0)) = -(2, 0) = (-2, 0)$$

et :

$$f((1, 1)) + f((1, -1)) = (1, 1) + (1, -1) = (2, 0),$$

donc f n'est pas linéaire. $\square$

Si E est un espace vectoriel euclidien, les applications de E dans E qui conservent l'orthogonalité sont les similitudes affines de E . En revanche, ceci ne caractérise pas les similitudes vectorielles : une application qui conserve l'orthogonalité peut ne pas être linéaire.

17.28. Application qui conserve l'orthogonalité mais qui n'est pas linéaire.

Nous munissons l'espace vectoriel réel $\mathbb{R}^2$ de son produit scalaire canonique (voir l'exemple précédent 17.26) et nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R}^2 \longrightarrow \mathbb{R}^2 \\ x = (x_1, x_2) \longmapsto f(x) = \begin{cases} x & \text{si } x_1 x_2 \neq 0, \\ (x_2, 0) & \text{si } x_1 = 0, \\ (0, x_1) & \text{si } x_2 = 0. \end{cases} \end{array} \right.$$

Nous notons D_1 et D_2 les « axes » du repère canonique de $\mathbb{R}^2$, ce qui signifie que D_1 est la droite vectorielle engendrée par $e_1 = (1, 0)$ et D_2 la droite vectorielle engendrée par $e_2 = (0, 1)$. La restriction de f à $\mathbb{R}^2 \setminus (D_1 \cup D_2)$ est l'identité de $\mathbb{R}^2 \setminus (D_1 \cup D_2)$, sa restriction à D_1 est une bijection de D_1 sur D_2 et sa restriction à D_2 une bijection de D_2 sur D_1 , donc f est une bijection de $\mathbb{R}^2$ sur $\mathbb{R}^2$. Les deux axes étant orthogonaux, f conserve l'orthogonalité. Or :

$$f((1, 0) + (0, 2)) = f((1, 2)) = (1, 2)$$

et :

$$f((1, 0)) + f((0, 2)) = (0, 1) + (2, 0) = (2, 1),$$

donc f n'est pas linéaire. $\square$

Chapitre 18

Courbes planes

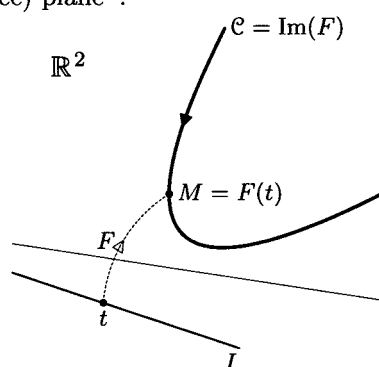
La géométrie plane existe depuis des temps reculés, et les courbes ont été étudiées dès l'Antiquité, qu'il s'agisse des courbes de nature algébrique comme les coniques ou des courbes de nature mécanique comme la conchoïde de Nicomède¹ ou la spirale d'Archimède. L'étude des tangentes aux courbes planes se développe aux XVII^e et XVIII^e siècles parallèlement à celle de la dérivabilité, et la notion de courbe se précise dans la deuxième moitié du XIX^e siècle ; quant aux paradoxes concernant les courbes, ils apparaissent pour la plupart entre 1880 et 1930.

Dans tout le chapitre, nous travaillons dans un plan euclidien (orienté), identifié par le choix d'un repère orthonormal (direct) à $\mathbb{R}^2$ muni de sa structure euclidienne canonique ; l'origine du plan $\mathbb{R}^2$ est le point $O = (0,0)$ et on note $\|\cdot\|$ sa norme euclidienne canonique. Le mot « intervalle » désigne un intervalle de $\mathbb{R}$ d'intérieur non vide et k est un entier naturel ou l'infini.

Nous définissons la notion de courbe (paramétrée) plane².

DÉFINITION 18.1. — Une courbe plane de classe $\mathcal{C}^k$ — ou simplement une courbe de classe $\mathcal{C}^k$ — est un couple $\gamma = (I, F)$ où I est un intervalle et F une application de I dans $\mathbb{R}^2$ de classe $\mathcal{C}^k$ sur I .

DÉFINITION 18.2. — Soit $\gamma = (I, F)$ une courbe de classe $\mathcal{C}^k$. Pour tout $t \in I$, le point de γ de paramètre t est le couple $M = (t, F(t))$, le plus souvent identifié au point géométrique $M = F(t)$, l'application F est le paramétrage de γ et la partie $\mathcal{C} = \text{Im}(F)$ de $\mathbb{R}^2$ est appelée le support de la courbe γ .



La flèche indique le « sens de parcours » de $\mathcal{C}$ lorsque le « paramètre » t décrit l'intervalle I en croissant.

1. Mathématicien grec, qui vécut à Alexandrie au II^e siècle av. J.C. Il construit sa conchoïde dans le but de résoudre le problème de la trisection de l'angle.

2. Cette définition des courbes paramétrées est introduite par Camille Jordan en 1883 dans son *Cours d'Analyse*. Suite à la construction en 1890 par Giuseppe Peano d'une courbe remplissant un carré — voir l'exemple 18.12, pages 144 et 145 —, il impose en 1893 à l'application F définissant la courbe d'être injective.

Nous aurons besoin, pour définir les tangentes et les directions asymptotiques, de la notion de limite en α d'une fonction $D : t \mapsto D(t)$ de $\mathbb{R}$ dans l'ensemble $\mathcal{D}$ des droites vectorielles de $\mathbb{R}^2$, définie au voisinage d'un point α de $\overline{\mathbb{R}}$. Le lecteur trouvera dans [RAM5], § 1.1.2, comment la définir à l'aide d'une topologie sur $\mathcal{D}$. On peut cependant en donner une version plus intuitive.

On prouve d'abord — par exemple à l'aide de la condition nécessaire et suffisante d'égalité dans la relation de Cauchy-Schwarz — un lemme : Si $U : t \mapsto U(t)$ et $V : t \mapsto V(t)$ sont des fonctions de $\mathbb{R}$ dans $\mathbb{R}^2$ définies au voisinage de α , si $U(t)$ admet pour limite un vecteur non nul X de $\mathbb{R}^2$ quand t tend vers α , si la famille $(U(t), V(t))$ est liée pour t au voisinage de α et si $V(t)$ admet une limite Y quand t tend vers α , alors la famille (X, Y) est liée.

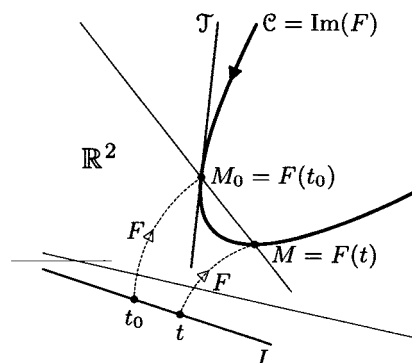
Soit $\alpha \in \overline{\mathbb{R}}$, une fonction $D : t \mapsto D(t)$ de $\mathbb{R}$ dans l'ensemble des droites vectorielles de $\mathbb{R}^2$, définie au voisinage de α , et L une droite vectorielle de $\mathbb{R}^2$. On dit que $D(t)$ admet pour limite L quand t tend vers α s'il existe une fonction $U : t \mapsto U(t)$ de $\mathbb{R}$ dans $\mathbb{R}^2$, définie au voisinage de α , et un vecteur directeur X de L tels que $(U(t))$ est, pour t au voisinage de α , une base de $D(t)$ et $U(t)$ admet pour limite X quand t tend vers α . On déduit alors du lemme l'unicité de l'éventuelle limite de $D(t)$ quand t tend vers α et de cette définition que si $U : t \mapsto U(t)$ est une fonction de $\mathbb{R}$ dans $\mathbb{R}^2$ définie au voisinage de α , si, pour t au voisinage de α , $(U(t))$ est une base de $D(t)$ et si $U(t)$ admet une limite X non nulle quand t tend vers α , alors $D(t)$ admet pour limite quand t tend vers α la droite vectorielle $\underline{L}$ engendrée par X .

■ Tangentes et points d'inflexion

Soit $\gamma = (I, F)$ une courbe de classe $\mathcal{C}^k$, t_0 un point de l'intervalle I et M_0 le point de γ de paramètre t_0 .

DÉFINITION. — Soit $\mathcal{T}$ une droite affine de $\mathbb{R}^2$ passant par le point M_0 . La droite $\mathcal{T}$ est la tangente à γ en M_0 si le point $M = F(t)$ est, pour t au voisinage de t_0 , différent de $M_0 = F(t_0)$ et si la direction de la droite affine $(M_0 F(t))$ admet pour limite la direction de la droite $\mathcal{T}$ quand t tend vers t_0 .

On peut parler de la tangente à γ en M_0 en raison de l'unicité de la limite d'une droite vectorielle.



Les vecteurs dérivés successifs de F et la formule de Taylor-Young permettent dans les cas favorables de prouver l'existence de la tangente et de la déterminer.

THÉORÈME 18.1. — Si F est dérivable en t_0 et si $F'(t_0)$ est différent du vecteur zéro, la droite $\mathcal{T}$ passant par M_0 et dirigée par $F'(t_0)$ est la tangente γ en M_0 .

THÉORÈME 18.2. — Si $k \geq 2$, si $p \in \mathbb{N}$ et $2 \leq p \leq k$, si le vecteur $F^{(i)}(t_0)$ est nul pour tout entier i tel que $1 \leq i < p$ et si $F^{(p)}(t_0)$ est différent du vecteur zéro, la droite $\mathcal{T}$ passant par M_0 et dirigée par $F^{(p)}(t_0)$ est la tangente γ en M_0 .

Un cas particulier important est celui où le paramètre est l'abscisse x du point M .

DÉFINITION 18.3. — Si I est un intervalle et f une application de I dans $\mathbb{R}$, de classe $\mathcal{C}^k$ sur I , la courbe d'équation cartésienne $y = f(x)$ est la courbe $\gamma = (I, F)$ de classe $\mathcal{C}^k$ définie par l'application :

$$\left| \begin{array}{l} F : I \longrightarrow \mathbb{R}^2 \\ x \longmapsto F(x) = (x, f(x)). \end{array} \right.$$

THÉORÈME 18.3. — Soit I un intervalle, f une application de I dans $\mathbb{R}$, de classe $\mathcal{C}^k$ sur I , et γ la courbe d'équation cartésienne $y = f(x)$. Si x_0 appartient à I , si $y_0 = f(x_0)$ et si f est dérivable en x_0 , la courbe γ admet en $M_0 = (x_0, y_0)$ une tangente $\mathcal{T}$, et $\mathcal{T}$ est la droite passant par M_0 et dirigée par le vecteur $(1, f'(x_0))$, c'est-à-dire la droite d'équation $y = y_0 + f'(x_0)(x - x_0)$.

Pour l'étude détaillée de l'allure d'une courbe au voisinage d'un point (points d'inflexion, points de rebroussement...), voir par exemple [RAM5] (chapitre 1).

Une courbe γ peut avoir une tangente en son point M_0 de paramètre t_0 sans que le paramétrage F soit dérivable en t_0 .

18.1. Courbe dont le paramétrage n'est pas dérivable en un point mais qui admet une tangente en ce point.

Nous considérons l'application, clairement continue sur $\mathbb{R}$:

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ t \longmapsto f(t) = \begin{cases} 0 & \text{si } t = 0, \\ t\left(2 + \sin \frac{1}{t}\right) & \text{si } t \neq 0 \end{cases} \end{array} \right.$$

et la courbe $\gamma = (\mathbb{R}, F)$ de classe $\mathcal{C}^0$ où :

$$\left| \begin{array}{l} F : \mathbb{R} \longrightarrow \mathbb{R}^2 \\ t \longmapsto F(t) = (f(t), f(t)). \end{array} \right.$$

On a, pour tout nombre réel $t \neq 0$, $Qf(0, t) = \frac{f(t) - f(0)}{t - 0} = \frac{f(t)}{t} = 2 + \sin \frac{1}{t}$.

A l'aide de la suite de terme général $1/((\pi/2) + n\pi)$, on prouve que le quotient $Qf(0, t)$ n'admet pas de limite quand t tend vers 0. Par conséquent f n'est pas dérivable en 0, donc la fonction F n'est pas dérivable en 0. Le point de γ de paramètre 0 est $F(0) = (0, 0) = O$. Pour tout nombre réel $t \neq 0$, $2 + \sin(1/t) > 0$, donc $F(t) \neq O$ et la droite affine $(OF(t))$ est dirigée par le vecteur constant $U = (1, 1)$, donc la droite $\mathcal{T}$ passant par O et dirigée par U , c'est-à-dire la droite d'équation $y = x$, est la tangente à γ en son point O de paramètre 0. $\square$

Dans l'exemple précédent 18.1, la fonction F n'est pas dérivable en t_0 . Montrons que F peut être de classe $\mathcal{C}^\infty$ sans que la courbe admette une tangente en t_0 .

18.2. Courbe de classe $\mathcal{C}^\infty$ n'admettant pas de tangente en un point.

Nous avons vu dans l'exemple 9.25 (page 178) que l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ e^{-\frac{1}{x^2}} & \text{si } x \neq 0 \end{cases} \end{array} \right.$$

est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$ et que $f^{(n)}(0) = 0$ pour tout entier naturel n . Il en résulte que $\gamma_1 = (\mathbb{R}, F)$ et $\gamma_2 = (\mathbb{R}, G)$, pour les applications :

$$\left| \begin{array}{l} F : \mathbb{R} \longrightarrow \mathbb{R}^2 \\ t \longmapsto F(t) = (f(t), 0) \end{array} \right. \text{ et } \left| \begin{array}{l} G : \mathbb{R} \longrightarrow \mathbb{R}^2 \\ t \longmapsto G(t) = \begin{cases} (0, f(t)) & \text{si } t < 0, \\ (f(t), 0) & \text{si } t \geq 0, \end{cases} \end{array} \right.$$

sont des courbes de classe $\mathcal{C}^\infty$. Le point de γ_1 de paramètre 0 est $O = (0, 0)$ ainsi que celui de γ_2 . Pour tout réel $t \neq 0$, $F(t) \neq O$ et la droite affine $(OF(t))$ est dirigée par le vecteur constant $U = (1, 0)$, donc la droite (Ox) est la tangente à γ_1 en son point de paramètre 0. Pour tout réel $t \neq 0$, $G(t) \neq O$ et la droite affine $(OG(t))$ est dirigée par le vecteur constant $V = (0, 1)$ si $t < 0$ et par le vecteur constant $U = (1, 0)$ si $t > 0$, donc, si γ_2 admettait une tangente en $O = (0, 0)$, ce serait à la fois (Oy) et (Ox) . Par conséquent la courbe γ_2 n'admet pas de tangente en son point de paramètre 0. $\square$

On peut, grâce à la fonction f de l'exemple précédent 18.2, construire de manière analogue des courbes admettant en leur point $O = (0, 0)$ de paramètre 0 un point d'inflexion ou un point de rebroussement, sans que ceux-ci puissent être déterminés par les dérivées successives. Ainsi, si l'on considère les applications :

$$\left| \begin{array}{l} h_1 : \mathbb{R} \longrightarrow \mathbb{R} \\ t \longmapsto h_1(t) = \begin{cases} 0 & \text{si } t = 0, \\ f(t) \sin \frac{1}{t} & \text{si } t \neq 0, \end{cases} \\ h_2 : \mathbb{R} \longrightarrow \mathbb{R} \\ t \longmapsto h_2(t) = \begin{cases} 0 & \text{si } t = 0, \\ f(t) \cos \frac{1}{t} & \text{si } t \neq 0, \end{cases} \end{array} \right.$$

et :

$$\left| \begin{array}{l} H : \mathbb{R} \longrightarrow \mathbb{R}^2 \\ t \longmapsto H(t) = (h_1(t), h_2(t)), \end{array} \right.$$

la courbe $(\mathbb{R}, H)$ est une spirale dont le point de paramètre t tend vers $O = (0, 0)$ quand t tend vers 0 et atteint $(0, 0)$ en $t = 0$, sans tangente bien que H soit de classe $\mathcal{C}^\infty$ en 0. Si l'on considère l'application :

$$\left| \begin{array}{l} K : \mathbb{R} \longrightarrow \mathbb{R}^2 \\ t \longmapsto K(t) = \begin{cases} (-t, -f(t)) & \text{si } t < 0, \\ (t, f(t)) & \text{si } t \geq 0, \end{cases} \end{array} \right.$$

la courbe $(\mathbb{R}, K)$ admet un point de rebroussement de première espèce en son point $O = (0, 0)$ de paramètre 0, pour l'application :

$$\left| \begin{array}{l} L : \mathbb{R} \longrightarrow \mathbb{R}^2 \\ t \longmapsto L(t) = \begin{cases} (-t, 3 - f(t)) & \text{si } t < 0, \\ (t, 3 + f(t)) & \text{si } t \geq 0, \end{cases} \end{array} \right.$$

la courbe $(\mathbb{R}, L)$ admet un point de rebroussement de deuxième espèce en son point $O = (0, 0)$ de paramètre 0 et pour l'application :

$$\left| \begin{array}{l} M : \mathbb{R} \longrightarrow \mathbb{R}^2 \\ t \longmapsto M(t) = \begin{cases} (t, -f(t)) & \text{si } t < 0, \\ (t, f(t)) & \text{si } t \geq 0, \end{cases} \end{array} \right.$$

la courbe $(\mathbb{R}, M)$ admet un point de rebroussement de deuxième espèce en son point $O = (0, 0)$ de paramètre 0.

18.3. Courbe traversant une infinité de fois sa tangente en un point.

Nous avons vu dans l'exemple 9.9 (page 166) que l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ x^2 \sin \frac{1}{x} & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

est dérivable sur $\mathbb{R}$ et que $f'(0) = 0$. La tangente $\mathcal{T}$ à la courbe γ d'équation $y = f(x)$ en son point $O = (0, 0)$ de paramètre 0 est donc la droite passant par O et dirigée par le vecteur $U = (1, 0)$; par suite $\mathcal{T} = (0x)$. Or f change de signe en chaque point $x_k = 1/(k\pi)$ où $k \in \mathbb{Z} \setminus \{0\}$, donc la courbe γ traverse une infinité de fois sa tangente au voisinage de O . $\square$

Si la courbe d'équation $y = f(x)$ reste d'un même côté de la tangente au voisinage du point étudié, on parle d'allure normale et, si elle la traverse en ce point, on dit qu'il s'agit d'un point d'inflexion. L'exemple précédent 18.3 montre qu'il est possible que l'on n'ait ni une allure normale, ni un point d'inflexion.

Pour une courbe γ d'équation cartésienne $y = f(x)$ — voir la définition 18.3 et le théorème 18.3 (page 340) — et un point x_0 intérieur à l'intervalle de définition de la fonction f , nous précisons ce que signifie : le point $M_0 = (x_0, f(x_0))$ de paramètre x_0 est un point d'inflexion de la courbe γ .

DÉFINITION 18.4. — Soit γ la courbe d'équation cartésienne $y = f(x)$ où I est un intervalle, f une application continue de I dans $\mathbb{R}$ et x_0 un point intérieur à I . Le point $M_0 = (x_0, y_0) = (x_0, f(x_0))$ de paramètre x_0 est un point d'inflexion de γ si la fonction f est dérivable en x_0 et s'il existe un nombre réel $\alpha > 0$ tel que $f(x) \leq y_0 + f'(x_0)(x - x_0)$ pour tout $x \in]x_0 - \alpha, x_0[$ et $f(x) \geq y_0 + f'(x_0)(x - x_0)$ pour tout $x \in]x_0, x_0 + \alpha[$, ou $f(x) \geq y_0 + f'(x_0)(x - x_0)$ pour tout $x \in]x_0 - \alpha, x_0[$ et $f(x) \leq y_0 + f'(x_0)(x - x_0)$ pour tout $x \in]x_0, x_0 + \alpha[$ — ce qui, puisque la droite $\mathcal{T}$ d'équation $y = y_0 + f'(x_0)(x - x_0)$ est la tangente à γ en M_0 , signifie que le point de paramètre x est au-dessous de la tangente pour tout $x \in]x_0 - \alpha, x_0[$ et au-dessus de la tangente pour tout $x \in]x_0, x_0 + \alpha[$, ou l'inverse.

On dispose alors des deux théorèmes suivants.

THÉORÈME 18.4. — Si la fonction f est dérivable sur l'intervalle I , si x_0 est un point intérieur à I , si $f'(x_0) = 0$ et si la fonction dérivée première f' garde un signe constant sur un voisinage $]x_0 - r, x_0 + r[$ de x_0 inclus dans I , le point $M_0 = (x_0, y_0) = (x_0, f(x_0))$ de paramètre x_0 est un point d'inflexion de la courbe γ d'équation cartésienne $y = f(x)$.

THÉORÈME 18.5. — Si f est deux fois dérivable sur l'intervalle I , si x_0 est un point intérieur à I et si sa fonction dérivée seconde f'' s'annule en x_0 en changeant de signe, le point $M_0 = (x_0, y_0) = (x_0, f(x_0))$ de paramètre x_0 est un point d'inflexion de la courbe γ d'équation cartésienne $y = f(x)$.

Les réciproques de ces deux théorèmes sont fausses.

18.4. Inflexion au point de paramètre 0 d'une courbe d'équation cartésienne $y = f(x)$, la fonction dérivée f' s'annulant en 0 mais ne gardant pas un signe constant au voisinage de 0.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ x^3 \left(2 + \sin \frac{1}{x} \right) & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

La fonction f est de classe $\mathcal{C}^1$ sur les intervalles $]-\infty, 0[$ et $]0, +\infty[$ et, pour tout réel $x \neq 0$:

$$f'(x) = 3x^2 \left(2 + \sin \frac{1}{x} \right) - x \cos \frac{1}{x}.$$

On a, pour tout réel $x \neq 0$, $|x(2 + \sin(1/x))| \leq 3|x|$ et $|x \cos(1/x)| \leq |x|$, donc :

$$\lim_{x \rightarrow 0} \left(x \left(2 + \sin \frac{1}{x} \right) \right) = 0 \quad \text{et} \quad \lim_{x \rightarrow 0} \left(x \cos \frac{1}{x} \right) = 0.$$

Par conséquent, $f'(x)$ admet pour limite 0 quand x tend vers 0, donc f est de classe $\mathcal{C}^1$ sur $\mathbb{R}$ et $f'(0) = 0$. Nous étudions la courbe γ d'équation cartésienne $y = f(x)$, qui est donc de classe $\mathcal{C}^1$. Comme $f'(0) = 0$, la tangente à γ en son point $O = (0, 0)$ de paramètre 0 est la droite (Ox) , d'équation $y = 0$. Pour tout réel $x \neq 0$, $2 + \sin(1/x) > 0$, donc $f(x) < 0$ pour tout $x < 0$ et $f(x) > 0$ pour tout $x > 0$. Le point O de paramètre 0 est donc un point d'inflexion de γ .

Nous introduisons la suite $(x_n)_{n \geq 1}$ de terme général $x_n = \frac{1}{n\pi}$. On a, pour tout entier $n \geq 1$:

$$f'(x_n) = \frac{6}{n^2\pi^2} - \frac{(-1)^n}{n\pi} = \frac{(-1)^{n+1}}{n\pi} \left(1 - (-1)^n \frac{6}{n\pi} \right).$$

Comme $1 < 6/\pi < 2$, on voit que, pour tout entier $n \geq 2$, $6/(n\pi) < 1$. Il en résulte que, pour tout entier $n \geq 2$, $f'(x_n)$ est du signe de $(-1)^{n+1}$. Tout voisinage de 0 à droite contient tous les x_n pour n suffisamment grand, donc f' ne garde un signe constant sur aucun voisinage de 0 à droite. De même on prouve, à l'aide de la suite $(y_n)_{n \geq 1} = (-x_n)_{n \geq 1}$, que f' ne garde un signe constant sur aucun voisinage de 0 à gauche. Ainsi, f' ne garde de signe constant sur aucun voisinage de 0, alors que le point O de paramètre 0 est un point d'inflexion de γ . $\square$

18.5. Inflexion au point de paramètre 0 d'une courbe d'équation cartésienne $y = f(x)$, la fonction f'' dérivée d'ordre 2 de f s'annulant en 0 sans changer de signe.

Nous considérons l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} 0 & \text{si } x = 0, \\ x^5 \left(2 + \sin \frac{1}{x} \right) & \text{si } x \neq 0. \end{cases} \end{array} \right.$$

La fonction f est de classe $\mathcal{C}^2$ sur les intervalles $]-\infty, 0[$ et $]0, +\infty[$ et, pour tout réel $x \neq 0$:

$$f'(x) = 5x^4 \left(2 + \sin \frac{1}{x} \right) - x^3 \cos \frac{1}{x}$$

et :

$$f''(x) = 20x^3 \left(2 + \sin \frac{1}{x} \right) - 5x^2 \cos \frac{1}{x} - 3x^2 \cos \frac{1}{x} - x \sin \frac{1}{x}.$$

On a, pour tout nombre réel $x \neq 0$, $|x(2 + \sin(1/x))| \leq 3|x|$, $|x \sin(1/x)| \leq |x|$ et $|x \cos(1/x)| \leq |x|$, donc les trois expressions $x(2 + \sin(1/x))$, $x \sin(1/x)$ et $x \cos(1/x)$ admettent pour limite 0 quand x tend vers 0. Par conséquent $f'(x)$ et $f''(x)$ admettent pour limite 0 quand x tend vers 0, donc la fonction f est de classe $\mathcal{C}^2$ sur $\mathbb{R}$, $f'(0) = 0$ et $f''(0) = 0$. Nous étudions la courbe γ d'équation cartésienne $y = f(x)$, qui est donc de classe $\mathcal{C}^2$. On prouve comme dans l'exemple précédent 18.4 que la tangente à la courbe γ en son point $O = (0, 0)$ de paramètre 0 est la droite (Ox) et que le point de paramètre 0 est un point d'inflexion de γ .

L'expression de $f''(x)$ pour tout $x \neq 0$ montre que $f''(x) = -x \sin \frac{1}{x} + \frac{o}{x \rightarrow 0}(x)$.

Il existe donc un réel $\alpha > 0$ tel que, pour tout réel x différent de 0, $|x| < \alpha$ entraîne $|f''(x) + x \sin(1/x)| < x/2$. Nous introduisons la suite $(x_n)_{n \geq 1}$ de terme général :

$$x_n = \frac{1}{n\pi + \frac{\pi}{2}}.$$

La suite (x_n) converge vers 0, ce qui justifie le choix d'un entier naturel n_0 tel que $0 < x_n < \alpha$ pour tout entier $n \geq n_0$. De plus, $\sin(1/x_n) = (-1)^n$ pour tout $n \in \mathbb{N}$. Il en résulte que, pour tout entier $n \geq n_0$:

$$\left| f''(x_n) + \frac{(-1)^n}{n\pi + \pi/2} \right| < \frac{1}{2(n\pi + \pi/2)}$$

donc :

$$-\frac{1}{2} < \left(n\pi + \frac{\pi}{2} \right) f''(x_n) + (-1)^n < \frac{1}{2},$$

ce qui montre que :

$$\begin{cases} -\frac{3}{2} < \left(n\pi + \frac{\pi}{2} \right) f''(x_n) < -\frac{1}{2} & \text{si } n \text{ est pair,} \\ \text{et } \frac{1}{2} < \left(n\pi + \frac{\pi}{2} \right) f''(x_n) < \frac{3}{2} & \text{si } n \text{ est impair,} \end{cases}$$

d'où l'on déduit que $f''(x_n) < 0$ si n est pair et $f''(x_n) > 0$ si n est impair. Tout voisinage de 0 à droite contient tous les x_n pour n suffisamment grand, donc f'' ne garde un signe constant sur aucun voisinage de 0 à droite. De même on prouve, à l'aide de la suite $(y_n)_{n \geq 1} = (-x_n)_{n \geq 1}$, que f'' ne garde un signe constant sur aucun voisinage de 0 à gauche. Ainsi, f'' ne garde de signe constant sur aucun voisinage de 0, alors que le point de paramètre 0 est un point d'inflexion de γ . $\square$

■ Directions asymptotiques et asymptotes

Soit $\gamma = (I, F)$ une courbe de classe $\mathcal{C}^k$. Nous supposons que β appartient à $\overline{\mathbb{R}}$, que β est la borne supérieure ou la borne inférieure de l'intervalle I dans $\overline{\mathbb{R}}$ et que β n'appartient pas à I ; par suite $I =]\alpha, \beta[$, $[\alpha, \beta[$, $]\beta, \alpha[$ ou $]\beta, \alpha]$.

On démontre que si $A \in \mathbb{R}^2$, la distance de A à $F(t)$ tend vers $+\infty$ quand t tend vers β si, et seulement si, la distance de $O = (0, 0)$ à $F(t)$ tend vers $+\infty$ quand t tend vers β . On dit que γ admet une *branche infinie* quand t tend vers β si, pour tout $A \in \mathbb{R}^2$, la distance de A à $F(t)$ tend vers $+\infty$ quand t tend vers β , ce qui signifie que le point $M = F(t)$ « s'éloigne à l'infini » de tous les points du plan quand t tend vers β . De plus γ admet une *branche infinie* quand t tend vers β si, et seulement si, la distance de O à $F(t)$ tend vers $+\infty$ quand t tend vers β .

Supposons que γ admette une *branche infinie* quand t tend vers β . On prouve que

si $A \in \mathbb{R}^2$ et si la direction de la droite affine $(AF(t))$ admet une limite L quand t tend vers β , alors, pour tout $B \in \mathbb{R}^2$, la direction de la droite $(BF(t))$ admet pour limite L quand t tend vers β . On dit que la droite vectorielle L de $\mathbb{R}^2$ est la *direction asymptotique* de γ quand t tend vers β si, pour tout point A de $\mathbb{R}^2$, la direction de la droite $(AF(t))$ admet pour limite L quand t tend vers β . De plus la droite vectorielle L est la direction asymptotique de γ quand t tend vers β si, et seulement si, la direction de $(OF(t))$ admet pour limite L quand t tend vers β . Supposons que γ admette une branche infinie quand t tend vers β . On dit que la droite affine $\mathcal{A}$ de $\mathbb{R}^2$ est *asymptote* à γ quand t tend vers β si la distance du point $F(t)$ à la droite $\mathcal{A}$ admet pour limite 0 quand t tend vers β . On démontre l'unicité d'une éventuelle asymptote quand t tend vers β —qui est ainsi *l'asymptote* de γ quand t tend vers β —et que si la droite $\mathcal{A}$ est asymptote à γ quand t tend vers β , la direction de $\mathcal{A}$ est la direction asymptotique de γ quand t tend vers β .

Une croyance populaire affirme qu'une courbe s'approche de plus en plus de son asymptote sans jamais la rencontrer ; ceci est faux.

18.6. Courbe qui traverse une infinité de fois son asymptote.

Nous considérons les applications :

$$\left| \begin{array}{l} f :]0, +\infty[\longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \frac{\sin x}{x} \end{array} \right. \quad \text{et} \quad \left| \begin{array}{l} F :]0, +\infty[\longrightarrow \mathbb{R}^2 \\ x \longmapsto F(x) = (x, f(x)) \end{array} \right.$$

et la courbe $\gamma = (]0, +\infty[, F)$, de classe $\mathcal{C}^\infty$.

Pour tout $x > 0$, la distance de $O = (0, 0)$ à $F(x)$ est $\delta(x) = \sqrt{x^2 + (f(x))^2} \geq x$; comme $\lim_{(x \rightarrow +\infty)} \delta(x) = +\infty$, γ admet une branche infinie quand x tend vers $+\infty$. Pour tout nombre réel $x > 0$, la distance de $F(x)$ à la droite (Ox) est égale à $|f(x)| \leq 1/x$, donc cette distance admet pour limite 0 quand x tend vers $+\infty$. Par suite la droite (Ox) est l'asymptote de γ quand x tend vers $+\infty$. De plus, pour tout entier $k \geq 1$, $f(k\pi) = 0$ et f change de signe au point de γ d'abscisse $k\pi$, donc la courbe traverse son asymptote en tous les points $(k\pi, 0)$ pour $k \in \mathbb{N}^*$. $\square$

18.7. Branche infinie sans direction asymptotique.

Nous introduisons l'application :

$$\left| \begin{array}{l} F : [0, +\infty[\longrightarrow \mathbb{R} \\ t \longmapsto F(t) = (e^t \cos t, e^t \sin t) \end{array} \right.$$

et la courbe $\gamma = ([0, +\infty[, F)$, de classe $\mathcal{C}^\infty$. Pour tout réel $t \geq 0$, la distance de $O = (0, 0)$ à $F(t)$ est e^t , donc γ admet une branche infinie quand t tend vers $+\infty$. Supposons que γ admette une direction asymptotique L quand t tend vers $+\infty$. Il existe un réel $A \geq 0$, une fonction $U : t \mapsto U(t) = (u(t), v(t))$ de $\mathbb{R}$ dans $\mathbb{R}^2$ définie sur $]A, +\infty[$ et un vecteur directeur $X = (a, b)$ de L tels que $(U(t))$ est, pour tout $t > A$, une base de la direction $D(t)$ de la droite $(OF(t))$ et $U(t)$ admet pour limite X quand t tend vers $+\infty$. Pour tout réel $t > A$, $(\cos t, \sin t) \in D(t)$, donc il existe un réel $\varphi(t)$ tel que $\cos t = \varphi(t)u(t)$ et $\sin t = \varphi(t)v(t)$. On a, pour tout réel $t > A$, $1 = |\varphi(t)|\sqrt{(u(t))^2 + (v(t))^2}$, donc $|\varphi(t)|$ admet pour limite le réel $\lambda = 1/\sqrt{a^2 + b^2} > 0$ quand t tend vers $+\infty$. Par suite, $|\cos t|$ admet pour limite $\lambda|a|$ quand t tend vers $+\infty$; or, à l'aide des suites (x_n) et (y_n) de termes généraux

$x_n = 2n\pi$ et $y_n = (\pi/2) + x_n$, on voit que $|\cos t|$ n'admet pas de limite quand t tend vers $+\infty$. En conclusion, la courbe γ n'admet pas de direction asymptotique quand t tend vers $+\infty$. $\square$

18.8. Courbe qui admet une direction asymptotique mais pas d'asymptote.

Nous introduisons l'application :

$$\left| \begin{array}{l} F :]-1, +\infty[\longrightarrow \mathbb{R}^2 \\ t \longmapsto F(t) = \left(\frac{e^t}{t+1}, \frac{te^t}{t+1} \right) \end{array} \right.$$

et la courbe $\gamma = (]-1, +\infty[, F)$, de classe $\mathcal{C}^\infty$.

Pour tout réel $t > -1$, la distance de $O = (0, 0)$ au point $F(t)$ est :

$$\delta(t) = \frac{e^t}{t+1} \sqrt{1+t^2} \underset{t \rightarrow +\infty}{\sim} \frac{e^t}{t} \times t = e^t,$$

donc $\lim_{(t \rightarrow +\infty)} \delta(t) = +\infty$. Par suite γ admet une branche infinie quand t tend vers $+\infty$. En factorisant les coordonnées de $F(t)$ par $(te^t)/(t+1)$, on voit que, pour tout réel $t > 0$, la droite $(OF(t))$ est dirigée par le vecteur $U(t) = (1/t, 1)$. Comme $U(t)$ admet pour limite $(0, 1)$ quand t tend vers $+\infty$, la direction de la droite (Oy) est la direction asymptotique de γ quand t tend vers $+\infty$.

Supposons que γ admette une asymptote $\mathcal{A}$ quand t tend vers $+\infty$. La direction de $\mathcal{A}$ est la direction asymptotique, donc il existe un nombre réel a tel que $\mathcal{A}$ est la droite d'équation $x = a$. Pour tout réel $t > -1$, la distance de $F(t)$ à $\mathcal{A}$ est $|e^t/(t+1) - a|$, donc $|e^t/(t+1) - a|$ admet pour limite 0 quand t tend vers $+\infty$. Par suite $\lim_{(t \rightarrow +\infty)} e^t/(t+1) = a$, en contradiction avec $\lim_{(t \rightarrow +\infty)} e^t/(t+1) = +\infty$. En conclusion, γ n'admet pas d'asymptote quand t tend vers $+\infty$. $\square$

18.9. Autre courbe admettant une direction asymptotique mais pas d'asymptote.

Nous considérons l'application :

$$\left| \begin{array}{l} F : \mathbb{R} \longrightarrow \mathbb{R}^2 \\ x \longmapsto F(x) = (x, \sin x) \end{array} \right.$$

et la courbe $\gamma = (\mathbb{R}, F)$, de classe $\mathcal{C}^\infty$. Pour tout nombre réel $x > 0$, la distance de $O = (0, 0)$ à $F(x)$ est $\delta(x) = \sqrt{x^2 + \sin^2 x} \geq x$. On voit que $\lim_{(x \rightarrow +\infty)} \delta(x) = +\infty$, donc γ admet une branche infinie quand x tend vers $+\infty$. Pour tout réel $x > 0$, la droite $(OF(x))$ est dirigée par le vecteur $(x, \sin x)$, donc aussi par le vecteur $U(x) = (1, (\sin x)/x)$. Or, pour tout $x > 0$, $|(\sin x)/x| \leq 1/x$, donc le vecteur $U(x)$ admet pour limite $(1, 0)$ quand x tend vers $+\infty$, ce qui montre que la direction de la droite (Ox) est la direction asymptotique de γ quand x tend vers $+\infty$.

Supposons que γ admette une asymptote $\mathcal{A}$ quand x tend vers $+\infty$. La direction de $\mathcal{A}$ est la direction asymptotique, donc il existe un réel a tel que $\mathcal{A}$ est la droite d'équation $y = a$. Pour tout x , la distance de $F(x)$ à $\mathcal{A}$ est $|\sin x - a|$, qui admet donc pour limite 0 quand x tend vers $+\infty$; par conséquent sinus admet pour limite a en $+\infty$. Or, à l'aide de la suite $(u_n)_{n \geq 0}$ de terme général $u_n = (\pi/2) + n\pi$, on voit que sinus n'admet pas de limite en $+\infty$, donc γ n'admet pas d'asymptote quand x tend vers $+\infty$. $\square$

■ Longueur d'une courbe

Nous étudions ici des courbes $\gamma = (I, F)$ de classe $\mathcal{C}^k$ pour lesquelles l'intervalle I est un segment $[a, b]$ où a et b sont des nombres réels tels que $a < b$.

Soit $\gamma = ([a, b], F)$ une telle courbe. Si $\sigma = (x_0, x_1, \dots, x_n)$ est une subdivision du segment $[a, b]$, ce qui signifie que $(x_i)_{0 \leq i \leq n}$ est une famille de réels tels que $a = x_0 < x_1 < \dots < x_{n-1} < x_n = b$, on pose :

$$L(\gamma, \sigma) = \sum_{i=0}^{n-1} \|F(x_{i+1}) - F(x_i)\|$$

(voir le dessin ci-contre).

On note $\Delta[a, b]$ l'ensemble des subdivisions du segment $[a, b]$ et on pose :

$$\mathbb{L}(\gamma) = \{L(\gamma, \sigma) \mid \sigma \in \Delta[a, b]\}.$$

L'ensemble $\mathbb{L}(\gamma)$ est appelé l'ensemble des longueurs des lignes polygonales inscrites dans γ — voir le dessin. Plus la subdivision σ va être « fine », plus $L(\gamma, \sigma)$ va augmenter et, intuitivement, si la courbe γ a une longueur, se rapprocher de cette longueur, qui doit donc être la borne supérieure de $\mathbb{L}(\gamma)$ dans $\mathbb{R}$.

DÉFINITION 18.5. — La courbe $\gamma = ([a, b], F)$ est rectifiable si l'ensemble $\mathbb{L}(\gamma)$ est majoré dans $\mathbb{R}$ et, si γ est rectifiable, la longueur de γ est le nombre réel positif :

$$\ell(\gamma) = \sup(\mathbb{L}(\gamma)).$$

Si une courbe $\gamma = ([a, b], F)$ n'est pas rectifiable, $\sup(\mathbb{L}(\gamma)) = +\infty$, ce qui permet de dire que γ est de longueur infinie.

THÉORÈME 18.6. — Si $\gamma = ([a, b], F)$ est un courbe de classe $\mathcal{C}^k$ et si $k \geq 1$, γ est rectifiable et sa longueur est :

$$\ell(\gamma) = \int_a^b \|F'(t)\| dt.$$

18.10. Courbe de classe $\mathcal{C}^0$ qui n'est pas rectifiable.

Nous reprenons l'application, définie dans l'exemple 10.12 (page 203) :

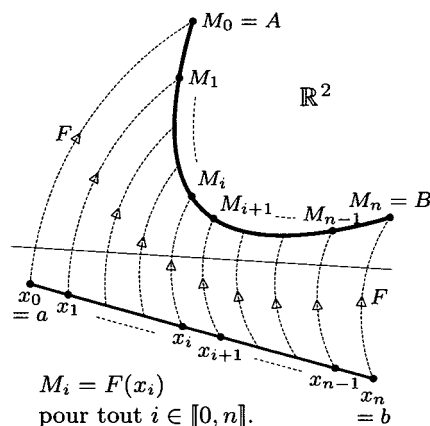
$$\left| \begin{array}{l} f : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f(x) = \begin{cases} x \sin \frac{\pi}{x} & \text{si } x \neq 0, \\ 0 & \text{si } x = 0. \end{cases} \end{array} \right.$$

Elle est continue sur $[0, 1]$. Nous lui associons la courbe γ d'équation cartésienne $y = f(x)$, c'est-à-dire la courbe $\gamma = ([0, 1], F)$ de classe $\mathcal{C}^0$ pour l'application :

$$\left| \begin{array}{l} F : [0, 1] \longrightarrow \mathbb{R}^2 \\ x \longmapsto F(x) = (x, f(x)). \end{array} \right.$$

Si $\sigma = (x_0, x_1, \dots, x_n)$ est une subdivision de $[0, 1]$, on a, pour tout $i \in \llbracket 0, n-1 \rrbracket$:

$$\|F(x_{i+1}) - F(x_i)\| = \sqrt{(x_{i+1} - x_i)^2 + (f(x_{i+1}) - f(x_i))^2} \geq |f(x_{i+1}) - f(x_i)|$$



donc :

$$L(\gamma, \sigma) = \sum_{i=0}^{n-1} \|F(x_{i+1}) - F(x_i)\| \geq \sum_{i=0}^{n-1} |f(x_{i+1}) - f(x_i)| = v(f, \sigma).$$

Or nous avons vu dans l'exemple 10.12 que la fonction f n'est pas à variation bornée sur le segment $[0, 1]$, ce qui signifie que l'ensemble des $v(f, \sigma)$ pour σ parcourant l'ensemble $\Delta[0, 1]$ des subdivisions de $[0, 1]$ n'est pas majoré dans $\mathbb{R}$. Par conséquent l'ensemble $\mathbb{L}(\gamma)$ des $L(\gamma, \sigma)$ pour $\sigma \in \Delta[0, 1]$ n'est pas majoré dans $\mathbb{R}$, donc γ n'est pas rectifiable. $\square$

18.11. Courbe de classe $\mathcal{C}^0$ dont la restriction à tout segment inclus dans l'intervalle de définition n'est pas rectifiable.

Nous reprenons l'application continue f de $\mathbb{R}$ dans $\mathbb{R}$ qui a été introduite dans l'exemple 10.13 (pages 204 et 205) et nous associons à f la courbe γ d'équation cartésienne $y = f(x)$, c'est-à-dire la courbe $\gamma = (\mathbb{R}, F)$ de classe $\mathcal{C}^0$ définie par l'application $F : x \mapsto F(x) = (x, f(x))$ de $\mathbb{R}$ dans $\mathbb{R}^2$.

Soit a et b des nombres réels tels que $a < b$. Nous avons vu dans l'exemple 10.13 que f n'est pas à variation bornée sur le segment $[a, b]$. Le même raisonnement que dans l'exemple précédent 18.10 montre alors que la courbe $(\gamma_{[a,b]}, F_{[a,b]})$, de classe $\mathcal{C}^0$, n'est pas rectifiable. $\square$

Soit a et b des réels tels que $a < b$ et f une application continue de $[a, b]$ dans $\mathbb{R}$, auxquels nous associons la courbe γ d'équation cartésienne $y = f(x)$, c'est-à-dire la courbe $\gamma = ([a, b], F)$ de classe $\mathcal{C}^0$ pour l'application :

$$\left| \begin{array}{l} F : [a, b] \longrightarrow \mathbb{R}^2 \\ x \longmapsto F(x) = (x, f(x)). \end{array} \right.$$

A toute subdivision $\sigma = (x_0, x_1, \dots, x_n)$ de $[a, b]$ sont associés les réels positifs :

$$L(\gamma, \sigma) = \sum_{i=0}^{n-1} \|F(x_{i+1}) - F(x_i)\| \quad \text{et} \quad v(f, \sigma) = \sum_{i=0}^{n-1} |f(x_{i+1}) - f(x_i)|.$$

On a, pour tout vecteur $U = (u_1, u_2)$ de $\mathbb{R}^2$, $|u_2| \leq \|U\| \leq |u_1| + |u_2|$, ce qui permet de prouver que, pour toute subdivision σ de $[a, b]$:

$$(1) \quad v(f, \sigma) \leq L(\gamma, \sigma) \leq (b - a) + v(f, \sigma).$$

Il en résulte que la courbe γ est rectifiable si, et seulement si, f est à variation bornée sur $[a, b]$. Cependant, si f est à variation bornée sur $[a, b]$, la longueur $\ell(\gamma)$ de γ n'est pas, en général, la variation totale $V_{[a,b]}(f) = \sup\{v(f, \sigma) \mid \sigma \in \Delta[a, b]\}$ de f sur $[a, b]$; (1) montre seulement que $V_{[a,b]}(f) \leq \ell(\gamma) \leq (b - a) + V_{[a,b]}(f)$.

Considérons un intervalle I et une suite (γ_n) de courbes de classe $\mathcal{C}^k$, de terme général $\gamma_n = (I, F_n)$ où F_n est une application de classe $\mathcal{C}^k$ de I dans $\mathbb{R}^2$. Si la suite de fonctions (F_n) converge sur I vers une application continue F de I dans $\mathbb{R}^2$, on obtient à la limite la courbe $\gamma = (I, F)$, de classe au moins $\mathcal{C}^0$. La convergence uniforme de la suite (F_n) sur I justifie l'existence de la fonction limite F et prouve sa continuité sur I , d'où l'existence de la courbe limite $\gamma = (I, F)$. Notons que dans les ouvrages de vulgarisation, la courbe limite est seulement suggérée

intuitivement. Si I est un segment, la longueur $\ell(\gamma)$ de la courbe limite n'est pas en général la limite de la suite $(\ell(\gamma_n))$ des longueurs des courbes γ_n comme le montrent les deux exemples suivants.

18.12. Suite (γ_n) de courbes telle que la longueur de sa courbe limite γ n'est pas la limite de la suite $(\ell(\gamma_n))$ des longueurs des courbes γ_n .

Nous introduisons comme dans l'exemple 9.6 (pages 161 et 162) l'application $g : x \mapsto g(x) = d(x, \mathbb{Z})$ (distance de x à $\mathbb{Z}$) de $\mathbb{R}$ dans $\mathbb{R}$, continue et de période 1 sur $\mathbb{R}$ — la fonction g est représentée graphiquement à la page 161. On a, pour tout réel x , $0 \leq g(x) \leq 1/2$. Nous associons à tout entier $n \geq 1$ l'application :

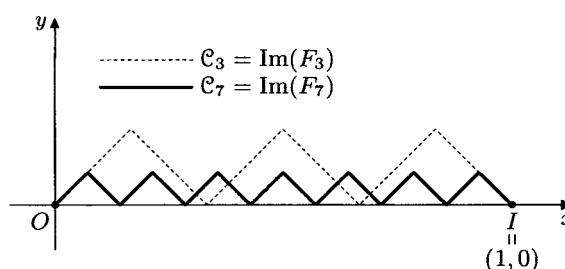
$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \frac{1}{n} g(nx) \end{array} \right.$$

et la courbe γ_n d'équation cartésienne $y = f_n(x)$, c'est-à-dire $\gamma_n = ([0, 1], F_n)$, de classe $\mathcal{C}^0$, pour l'application :

$$\left| \begin{array}{l} F_n : [0, 1] \longrightarrow \mathbb{R}^2 \\ x \longmapsto F_n(x) = (x, f_n(x)). \end{array} \right.$$

Pour tout entier $n \geq 1$, le support $\mathcal{C}_n$ de la courbe γ_n est une ligne brisée formée de $2n$ segments de longueur $\sqrt{2}/(2n)$, donc la longueur de γ_n est $\ell(\gamma_n) = \sqrt{2}$.

On a, pour tout entier $n \geq 1$ et tout $x \in [0, 1]$:



$$\|F_n(x) - (x, 0)\| = |f_n(x)| \leq \frac{1}{2n},$$

donc la suite (F_n) converge uniformément sur le segment $[0, 1]$ vers l'application $F : x \mapsto F(x) = (x, 0)$ de $[0, 1]$ dans $\mathbb{R}^2$. On en déduit que la courbe limite de la suite (γ_n) est la courbe $\gamma = ([0, 1], F)$, dont le support est le segment $[O, I]$ de $\mathbb{R}^2$ — pour le point $I = (1, 0)$ — et dont la longueur est égale à 1. $\square$

18.13. Suite (γ_n) de courbes telle que la longueur de sa courbe limite γ est finie alors que la suite $(\ell(\gamma_n))$ des longueurs des courbes γ_n tend vers $+\infty$.

Nous considérons l'application continue g de $\mathbb{R}$ dans $\mathbb{R}$ définie dans l'exemple 9.6 (pages 161 et 162) et utilisée dans l'exemple précédent 18.12. Nous associons à tout entier $n \geq 1$ l'application :

$$\left| \begin{array}{l} f_n : [0, 1] \longrightarrow \mathbb{R} \\ x \longmapsto f_n(x) = \frac{1}{\sqrt{n}} g(nx) \end{array} \right.$$

et la courbe γ_n d'équation cartésienne $y = f_n(x)$, c'est-à-dire $\gamma_n = ([0, 1], F_n)$, de classe $\mathcal{C}^0$, pour l'application $F_n : x \mapsto F_n(x) = (x, f_n(x))$ de $[0, 1]$ dans $\mathbb{R}^2$.

Pour tout entier $n \geq 1$, le support $\mathcal{C}_n$ de la courbe γ_n est une ligne brisée formée

de $2n$ segments de longueur :

$$\sqrt{\frac{1}{4n^2} + \frac{1}{4n}} = \frac{1}{2n} \sqrt{1+n},$$

donc la longueur de la courbe γ_n est égale à $\sqrt{1+n}$. Par conséquent la suite $(\ell(\gamma_n))$ tend vers $+\infty$. On prouve comme dans l'exemple précédent 18.12 que la suite (F_n) converge uniformément sur $[0, 1]$ vers l'application $F : x \mapsto F(x) = (x, 0)$ de $[0, 1]$ dans $\mathbb{R}$, ce qui montre que la courbe limite de la suite (γ_n) est la courbe $\gamma = ([0, 1], F)$, de support le segment $[O, I]$ et de longueur 1. $\square$

■ Courbes remarquables

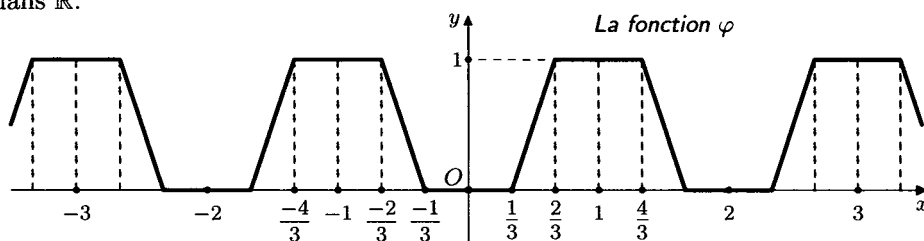
Peu après que Camille Jordan (voir la note 2, page 338) ait défini sa notion de courbe à l'aide de l'image continue d'un intervalle, Giuseppe Peano montre en 1890 qu'avec cette définition, une courbe peut remplir tout le carré $[0, 1] \times [0, 1]$. Ceci semble défier le bon sens — on imagine une courbe comme un fil de fer — et montre que pour bien des applications, il est nécessaire d'être plus restrictif sur la définition pour obtenir une notion de courbe correspondant à notre intuition.

18.14. Courbe de Peano.

Nous considérons la fonction φ à valeurs dans $\mathbb{R}$ définie sur $[-1, 1]$ par :

$$\varphi(t) = \begin{cases} 0 & \text{si } |t| \leq \frac{1}{3}, \\ 3|t| - 1 & \text{si } \frac{1}{3} < |t| < \frac{2}{3}, \\ 1 & \text{si } \frac{2}{3} \leq |t| \leq 1 \end{cases}$$

et prolongée à $\mathbb{R}$ en lui imposant d'être de période 2. Les définitions de $\varphi|_{[-1,1]}$ coïncidant aux bornes des intervalles de définition, la fonction φ est continue sur $[-1, 1]$; de plus, $\varphi(1) = \varphi(-1) (= 1)$, donc φ est une application continue de $\mathbb{R}$ dans $\mathbb{R}$.



Nous introduisons les suites $(u_n)_{n \geq 1}$ et $(v_n)_{n \geq 1}$ de fonctions dont les termes généraux sont les applications continues de $\mathbb{R}$ dans $\mathbb{R}$:

$$u_n : t \mapsto u_n(t) = \frac{\varphi(4^{2n}t)}{2^n} \quad \text{et} \quad v_n : t \mapsto v_n(t) = \frac{\varphi(4^{2n-1}t)}{2^n}.$$

Comme $0 \leq \varphi(x) \leq 1$ pour tout réel x , on a, pour tout réel t et tout $n \in \mathbb{N}^*$:

$$0 \leq u_n(t) \leq \frac{1}{2^n} \quad \text{et} \quad 0 \leq v_n(t) \leq \frac{1}{2^n}$$

donc les séries de fonctions $\sum_n u_n$ et $\sum_n v_n$ convergent uniformément sur $\mathbb{R}$,

ce qui justifie l'existence des applications :

$$f : t \mapsto f(t) = \sum_{n=1}^{+\infty} u_n(t) \text{ et } g : t \mapsto g(t) = \sum_{n=1}^{+\infty} v_n(t)$$

et prouve que f et g sont continues sur $\mathbb{R}$. De plus $\sum_{n=1}^{+\infty} 1/2^n = 1$, donc on a $0 \leq f(t) \leq 1$ et $0 \leq g(t) \leq 1$ pour tout nombre réel t . Nous considérons enfin l'application :

$$\left| \begin{array}{l} F : [0, 1] \longrightarrow \mathbb{R}^2 \\ t \longmapsto F(t) = (f(t), g(t)) \end{array} \right.$$

et la courbe $\gamma = ([0, 1], F)$, de classe $\mathcal{C}^0$. Pour $t \in [0, 1]$, $0 \leq f(t), g(t) \leq 1$, donc $F(t) \in [0, 1] \times [0, 1]$. Soit $M = (x, y)$ un point du carré $[0, 1] \times [0, 1]$. Alors $0 \leq x \leq 1$ et $0 \leq y \leq 1$. Nous voulons prouver que M appartient au support de la courbe γ . Nous considérons les développements illimités en base 2 de x et de y : $x = 0, x_1 x_2 x_3 \dots x_n \dots$ et $y = 0, y_1 y_2 y_3 \dots y_n \dots$; c'est le développement propre pour un point de $[0, 1[$ et, pour 1, du développement impropre $0, 111 \dots 1 \dots$. Nous définissons la suite $(a_n)_{n \geq 1}$ de chiffres en base 2 en posant, pour tout entier $p \geq 1$, $a_{2p} = x_p$ et $a_{2p-1} = y_p$ et nous introduisons le point t de $[0, 1]$ défini par :

$$t = \sum_{n=1}^{+\infty} \frac{a_n}{4^n} = 0, a_1 a_2 a_3 \dots a_n \dots \text{ (en base 4)}.$$

Nous calculons $F(t) = (f(t), g(t))$. Soit k un entier ≥ 1 . On a :

$$4^k t = \underbrace{\sum_{n=1}^{k-1} a_n 4^{k-n}}_{= m_k} + a_k + \underbrace{\sum_{n=k+1}^{+\infty} \frac{a_n}{4^{n-k}}}_{= \alpha_k},$$

m_k est un entier pair, $0 \leq \alpha_k \leq \sum_{n=k+1}^{+\infty} \frac{1}{4^{n-k}} = \sum_{p=1}^{+\infty} \frac{1}{4^p} = \frac{1}{3}$ et $\varphi(t) = \varphi(a_k + \alpha_k)$.

Si $a_k = 0$, on a $0 \leq a_k + \alpha_k \leq 1/3$ donc $\varphi(4^k t) = 0 = a_k$, et, si $a_k = 1$, alors $1 \leq a_k + \alpha_k \leq 4/3$ donc $\varphi(4^k t) = 1 = a_k$; ainsi, dans tous les cas, $\varphi(4^k t) = a_k$. Il en résulte que :

$$f(t) = \sum_{n=1}^{+\infty} \frac{\varphi(4^{2n} t)}{2^n} = \sum_{n=1}^{+\infty} \frac{a_{2n}}{2^n} = \sum_{n=1}^{+\infty} \frac{x_n}{2^n} = x$$

et :

$$g(t) = \sum_{n=1}^{+\infty} \frac{\varphi(4^{2n-1} t)}{2^n} = \sum_{n=1}^{+\infty} \frac{a_{2n-1}}{2^n} = \sum_{n=1}^{+\infty} \frac{y_n}{2^n} = y,$$

ce qui prouve que $M = (x, y) = F(t)$, donc que M appartient au support $\mathcal{C}$ de la courbe γ . En conclusion, le support $\mathcal{C}$ de la courbe γ est égal au carré $[0, 1] \times [0, 1]$: la courbe γ « remplit » tout le carré $[0, 1] \times [0, 1]$. $\square$

Nous construisons une courbe définie par une fonction continue, qui envoie l'ensemble de Cantor, d'intérieur vide et de mesure nulle, sur le carré $[0, 1] \times [0, 1]$. Cette courbe, comme celle de Peano, n'est évidemment pas rectifiable.

18.15. Courbe de Schoenberg³.

Nous reprenons la fonction φ définie sur $\mathbb{R}$ dans l'exemple précédent 18.14 et nous introduisons les suites $(u_n)_{n \geq 0}$ et $(v_n)_{n \geq 0}$ de fonctions dont les termes généraux sont les applications continues de $\mathbb{R}$ dans $\mathbb{R}$:

$$u_n : t \mapsto u_n(t) = \frac{\varphi(3^{2n}t)}{2^{n+1}} \text{ et } v_n : t \mapsto v_n(t) = \frac{\varphi(3^{2n+1}t)}{2^{n+1}}.$$

Comme $0 \leq \varphi(x) \leq 1$ pour tout réel x , on a, pour tout réel t et tout entier naturel n , $0 \leq u_n(t) \leq 1/2^{n+1}$ et $0 \leq v_n(t) \leq 1/2^{n+1}$, donc les séries de fonctions $\sum_n u_n$ et $\sum_n v_n$ convergent uniformément sur $\mathbb{R}$, ce qui justifie l'existence des applications :

$$f : t \mapsto f(t) = \sum_{n=0}^{+\infty} u_n(t) \text{ et } g : t \mapsto g(t) = \sum_{n=0}^{+\infty} v_n(t)$$

et prouve que f et g sont continues sur $\mathbb{R}$. De plus $\sum_{n=0}^{+\infty} 1/2^{n+1} = 1$, donc on a $0 \leq f(t) \leq 1$ et $0 \leq g(t) \leq 1$ pour tout nombre réel t . Nous considérons enfin l'application :

$$\left| \begin{array}{l} F : [0, 1] \longrightarrow \mathbb{R}^2 \\ t \longmapsto F(t) = (f(t), g(t)) \end{array} \right.$$

et la courbe $\gamma = ([0, 1], F)$, de classe $\mathcal{C}^0$. Pour $t \in [0, 1]$, $0 \leq f(t), g(t) \leq 1$, donc $F(t) \in [0, 1] \times [0, 1]$. Nous prouvons que l'image par F de l'ensemble de Cantor C étudié dans l'exemple 5.30 (pages 96 et 97) est le carré $[0, 1] \times [0, 1]$.

Soit $M = (x, y)$ un point de $[0, 1] \times [0, 1]$. Alors $0 \leq x \leq 1$ et $0 \leq y \leq 1$. Nous voulons prouver que M appartient au support de la courbe γ . Nous considérons les développements illimités en base 2 de x et de y : $x = 0, x_1 x_2 x_3 \dots x_n \dots$ et $y = 0, y_1 y_2 y_3 \dots y_n \dots$; il s'agit du développement propre pour un point de $[0, 1[$ et, pour 1, du développement impropre $0, 111 \dots 1 \dots$. Nous définissons la suite $(a_n)_{n \geq 1}$ de chiffres en base 3 en posant, pour tout $p \in \mathbb{N}^*$, $a_{2p} = 2x_p$ et $a_{2p-1} = 2y_p$, et nous posons $t = \sum_{n=1}^{+\infty} a_n/3^n = 0, a_1 a_2 a_3 \dots a_n \dots$ (en base 3), qui appartient à $[0, 1]$. Pour tout n , $a_n \neq 1$, donc t appartient à l'ensemble C . Comme φ est de période 2, que $\varphi(x) = 0$ pour $x \in [-1/3, 1/3]$ et que $\varphi(x) = 1$ pour $x \in [2/3, 4/3]$, on voit que, pour tout entier naturel k et tout point x de $[k - (1/3), k + (1/3)]$, $\varphi(x) = 0$ si k est pair et $\varphi(x) = 1$ si k est impair.

Soit $n \in \mathbb{N}^*$. On a, en base 3, $3^{2n}t = a_1 a_2 \dots a_{2n-1} a_{2n}, a_{2n+1} a_{2n+2} \dots$. L'entier naturel $k = a_1 a_2 \dots a_{2n-1} a_{2n}$ est pair car tous ses chiffres le sont. Si $a_{2n+1} = 0$, alors $0, a_{2n+1} a_{2n+2} \dots$ appartient à $[0, 1/3]$, donc $3^{2n}t$ appartient à $[k, k + (1/3)]$, ce qui montre que $\varphi(3^{2n}t) = 0$; si $a_{2n+1} = 2$, alors $0, a_{2n+1} a_{2n+2} \dots$ appartient à $[2/3, 1]$ donc $3^{2n}t$ appartient à $[k + (2/3), k + 1] = [(k + 1) - (1/3), k + 1]$, d'où l'on déduit que $\varphi(3^{2n}t) = 1$. Ainsi, dans les deux cas, $\varphi(3^{2n}t) = x_{n+1}$. On prouve de même que $\varphi(3^{2n+1}t) = y_{n+1}$. Par conséquent :

$$f(t) = \sum_{n=0}^{+\infty} \frac{x_{n+1}}{2^{n+1}} = \sum_{n=1}^{+\infty} \frac{x_n}{2^n} = x \text{ et } g(t) = \sum_{n=1}^{+\infty} \frac{y_{2n+1}}{2^{n+1}} = \sum_{n=1}^{+\infty} \frac{y_n}{2^n} = y,$$

ce qui prouve que $M = (x, y) = F(t)$, donc que M appartient à l'image de C par F . En conclusion, l'image par F de l'ensemble de Cantor est le carré $[0, 1] \times [0, 1]$. $\square$

3. Cette courbe a été introduite par le mathématicien d'origine roumaine Isaac Schoenberg (1903-1990).

On peut, comme dans l'exemple 9.6 (pages 161 et 162), démontrer que les fonctions f et g de l'exemple précédent 18.15 ne sont dérivables en aucun point.

On peut également, si $q \in \mathbb{N}$ et $q \geq 3$, envoyer par une application continue F de $[0, 1]$ dans $\mathbb{R}^q$ l'ensemble de Cantor sur $[0, 1]^q$, en considérant, pour tout $k \in \llbracket 1, q \rrbracket$, l'application continue de $[0, 1]$ dans $\mathbb{R}$:

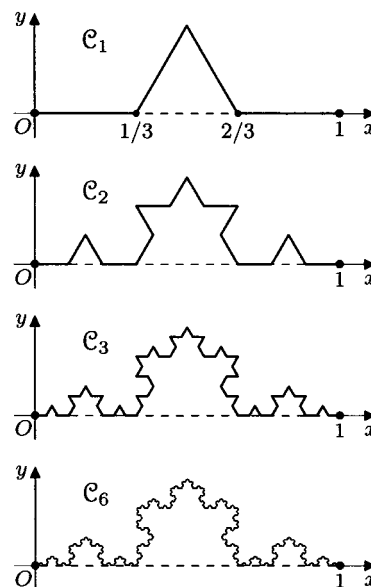
$$f_k : t \mapsto f_k(t) = \sum_{n=0}^{+\infty} \frac{f(3^{qn+k-1}t)}{2^{n+1}}$$

et en introduisant l'application $F : t \mapsto F(t) = (f_1(t), \dots, f_q(t))$ de $[0, 1]$ dans $\mathbb{R}^q$.

18.16. Courbe de Von Koch.

Nous considérons la suite $(\mathcal{C}_n)$ de lignes brisées construite par récurrence de la manière suivante : $\mathcal{C}_0$ est le segment $[O, I]$ où $O = (0, 0)$ et $I = (1, 0)$, $\mathcal{C}_1$ est la ligne brisée $[O, A, B, C, I] = [O, A] \cup [A, B] \cup [B, C] \cup [C, I]$ où $A = (1/3, 0)$ et $C = (2/3, 0)$ et où B est l'image de C dans la rotation de centre A et d'angle $\pi/3$, et, pour tout $n \in \mathbb{N}^*$, si $\mathcal{C}_n = [A_0, A_1, \dots, A_{4^n}] = [A_0, A_1] \cup [A_1, A_2] \cup \dots \cup [A_{4^n-1}, A_{4^n}]$, on obtient la ligne brisée $\mathcal{C}_{n+1}$ en remplaçant, pour tout $i \in \llbracket 1, 4^i \rrbracket$, $[A_{i-1}, A_i]$ par la ligne brisée $[A_{i-1}, B_i, C_i, D_i, A_i] = [A_{i-1}, B_i] \cup [B_i, C_i] \cup [C_i, D_i] \cup [D_i, A_i]$ où B_i est le barycentre $\frac{2}{3}A_{i-1} + \frac{1}{3}A_i$, D_i le barycentre $\frac{1}{3}A_{i-1} + \frac{2}{3}A_i$ et B_i l'image de D_i dans la rotation de centre B_i et d'angle $\pi/3$. La courbe de Von Koch⁴ est la « courbe limite » de la suite $(\mathcal{C}_n)_{n \in \mathbb{N}}$ de lignes brisées. Il nous faut donner un sens précis à cette notion de limite. Pour ceci nous construisons, par récurrence sur $n \in \mathbb{N}$, une suite $(\gamma_n)_{n \geq 0}$ de courbes de classe $\mathcal{C}^0$, de terme général $\gamma_n = ([0, 1], F_n)$ où F_n est une application continue de $[0, 1]$ dans $\mathbb{R}^2$, le support de γ_n étant, pour tout $n \in \mathbb{N}$, la ligne brisée $\mathcal{C}_n$. On débute par l'application $F_0 : t \mapsto F_0(t) = (t, 0)$ de $[0, 1]$ dans $\mathbb{R}^2$; le support de $\gamma_0 = ([0, 1], F_0)$ est $\mathcal{C}_0$ et sa longueur est égale à 1.

La courbe γ_1 est obtenue par l'application F_1 définie par : $F_1(t) = F_0(t)$ pour $t \in [0, 1/3[$ ou $t \in [2/3, 1]$; $F_1(t) = (t, (t-1/3)\sqrt{3})$ pour $t \in [1/3, 1/2[$; $F_1(t) = (t, (2/3-t)\sqrt{3})$ pour $t \in [1/2, 2/3[$. La famille $\tau_1 = (0, \frac{1}{3}, \frac{1}{2}, \frac{2}{3}, 1)$ est une subdivision de $[0, 1]$, on a $O = F_1(0)$, $A = F_1(\frac{1}{3})$, $B = F_1(\frac{1}{2})$, $C = F_1(\frac{2}{3})$ et $I = F_1(1)$, $\mathcal{C}_1$ est le support de γ_1 , la longueur de γ_1 est $4/3$ et $\|F_1(t) - F_0(t)\| \leq \sqrt{3}/6$ pour tout point t de $[0, 1]$, cette valeur étant atteinte pour $t = 1/2$. Soit $n \in \mathbb{N}^*$. Nous supposons connue la courbe $\gamma_n = ([0, 1], F_n)$, de support la ligne brisée $\mathcal{C}_n = [A_0, A_1, \dots, A_{4^n}]$, vérifiant ce qui suit : il existe une subdivision $\tau_n = (t_0, t_1, \dots, t_{4^n})$ de $[0, 1]$ telle que $F_n(t_i) = A_i$ pour tout $i \in \llbracket 0, 4^n \rrbracket$, la restriction de F_n à $[t_{i-1}, t_i]$ est affine pour tout $i \in \llbracket 1, 4^n \rrbracket$, la longueur de γ_n est $(4/3)^n$ et $\|F_n(t) - F_{n-1}(t)\| \leq \sqrt{3}/(2 \times 3^{n-1})$ pour tout $t \in [0, 1]$. Pour tout $i \in \llbracket 1, 4^n \rrbracket$, nous posons



4. Cette courbe a été introduite par le mathématicien suédois Helge Von Koch en 1906.

$\beta_i = (t_i - t_{i-1})/3$ et $u_i = F_n(t_i) - F_n(t_{i-1})$, et nous désignons par v_i l'unique vecteur unitaire tel que $((1/\|u_i\|)u_i, v_i)$ est une base orthonormale directe de $\mathbb{R}^2$. Nous construisons alors l'application F_{n+1} de $[0, 1]$ dans $\mathbb{R}^2$ ainsi : pour tout $i \in \llbracket 1, 4^n \rrbracket$, $F_{n+1}(t) = F_n(t)$ pour $t \in [t_{i-1}, t_{i-1} + \beta_i[$ ou $t \in [t_{i-1} + 2\beta_i, t_i[$; $F_{n+1}(t) = F_n(t) + (\sqrt{3}/3^n)(t - (t_{i-1} + \beta_i))v_i$ pour $t \in [t_{i-1} + \beta_i, t_{i-1} + (3/2)\beta_i[$; $F_{n+1}(t) = F_n(t) + (\sqrt{3}/3^n)(t_{i-1} + 2\beta_i - t)v_i$ pour $t \in [t_{i-1} + (3/2)\beta_i, t_{i-1} + 2\beta_i[$; ajoutons que si $i = 4^n$, on remplace $[t_{i-1} + 2\beta_i, t_i[$ par le segment $[t_{i-1} + 2\beta_i, t_i]$. On obtient ainsi la courbe $\gamma_{n+1} = ([0, 1], F_{n+1})$, dont on prouve aisément qu'elle possède au rang $n + 1$ les mêmes propriétés que γ_n au rang n ; en particulier, le support de γ_{n+1} est $\mathcal{C}_{n+1}$, sa longueur est $(4/3)^{n+1}$ et, pour tout point t de $[0, 1]$:

$$\|F_{n+1}(t) - F_n(t)\| \leq \frac{\sqrt{3}}{2 \times 3^n}.$$

Soit $n \in \mathbb{N}$. On a, pour tout entier $p \geq 1$ et tout point t du segment $[0, 1]$:

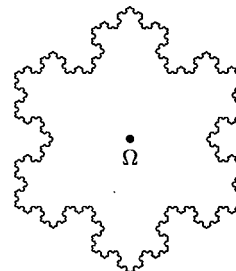
$$\begin{aligned} \|F_{n+p}(t) - F_n(t)\| &= \left\| \sum_{k=1}^p (F_{n+k}(t) - F_{n+k-1}(t)) \right\| \leq \sum_{k=1}^p \|F_{n+k}(t) - F_{n+k-1}(t)\| \\ &\leq \sum_{k=1}^p \frac{\sqrt{3}}{2} \times \frac{1}{3^{n+k}} = \frac{\sqrt{3}}{2} \times \frac{1}{3^{n+1}} \times \frac{1 - \frac{1}{3^p}}{1 - \frac{1}{3}} \leq \frac{\sqrt{3}}{4 \times 3^n} = \alpha_n. \end{aligned}$$

La majoration obtenue est indépendante de p et de t , et la suite (α_n) converge vers 0 donc, $(\mathbb{R}^2, \|\cdot\|)$ étant un espace vectoriel normé complet, le critère de Cauchy uniforme montre que la suite (F_n) converge uniformément sur $[0, 1]$ vers une application continue F de $[0, 1]$ dans $\mathbb{R}^2$. Par conséquent la suite (γ_n) de courbes de classe $\mathcal{C}^0$ admet pour limite la courbe $\gamma = ([0, 1], F)$, de classe $\mathcal{C}^0$, et la courbe de Von Koch est cette courbe limite γ . La représentation graphique du support $\mathcal{C}_6$ de γ_6 faite dans le dernier dessin de la page précédente donne une idée du support $\mathcal{C}$ de la courbe γ .

Soit $n \in \mathbb{N}^*$. Pour tout indice $i \in \llbracket 0, 4^n \rrbracket$, $F_{n+1}(t_i) = F_n(t_i)$ — par construction de l'application F_{n+1} —, on voit de même, de proche en proche, que, pour tout entier $p \geq 1$, $F_{n+p}(t_i) = F_n(t_i)$, et le passage à la limite quand p tend vers l'infini montre que $F(t_i) = F_n(t_i)$. Par suite, avec les notations introduites à la page 347, au début du paragraphe « Longueur d'une courbe », $L(\tau_n, \gamma)$ est la longueur de γ_n , donc $L(\tau_n, \gamma) = (4/3)^n$. La suite $(L(\tau_n, \gamma))$ tend donc vers $+\infty$, ce qui prouve que la courbe de Von Koch γ n'est pas rectifiable : sa longueur est infinie. On démontre également qu'aucun sous-arc de γ n'est rectifiable et qu'ainsi la distance sur γ entre deux points distincts quelconques de cette courbe est infinie. $\square$

18.17. Courbe fermée de longueur infinie, englobant une surface finie.

Nous introduisons le point $\Omega = (1/2, -\sqrt{3}/6)$ et nous faisons subir au support $\mathcal{C}$ de la courbe de Von Koch γ les rotations de centre Ω et d'angles $(2\pi)/3$ et $-(2\pi)/3$. Les trois arcs obtenus se ressoldent et l'on obtient alors le flocon de Von Koch, courbe fermée de longueur infinie, qui englobe une surface finie. $\square$



Chapitre 19

Probabilités

On considère souvent que le calcul des probabilités commence au milieu du XVI^e siècle avec les échanges épistolaires entre Pierre de Fermat et Blaise Pascal. Peu après, des ouvrages de Christiaan Huygens, puis de Jacques Bernoulli, fondent les premiers éléments de la théorie. Pierre Siméon Laplace publie en 1812 la théorie analytique des probabilités, ouvrage dans lequel il présente la synthèse des nouvelles approches de cette théorie et où il utilise les nouveaux outils mathématiques développés au XVIII^e siècle. Les premières théories sur les processus stochastiques se développent vers 1900, en particulier avec Andreï Markov. Il faut attendre les travaux d'Andreï Kolmogorov pour aboutir, en 1929, à la formalisation de la théorie que nous utilisons de nos jours ; il y fait figurer les notions établis et étudiées par Henri Lebesgue, Emile Borel et Johann Radon dans le domaine de l'intégrale abstraite, en particulier les tribus et la théorie de la mesure.

Pour modéliser une expérience aléatoire, on introduit un ensemble Ω représentant toutes les résultats possibles. On entend par événements certaines parties de l'ensemble Ω . A chaque événement, on associe une probabilité d'apparition, c'est-à-dire un réel appartenant au segment $[0, 1]$ vérifiant certaines propriétés naturelles. Il est en général impossible, si Ω est infini, d'attribuer à chaque partie une probabilité en respectant ces propriétés, et l'on est amené à se restreindre à un sous ensemble de l'ensemble des parties de Ω , appelé une tribu sur Ω . Aussi pose-t-on les définitions suivantes.

DÉFINITION 19.1. — Une tribu sur un ensemble Ω est une famille $\mathcal{T}$ de parties de Ω vérifiant les trois propriétés suivantes :

- (I) Ω appartient à $\mathcal{T}$.
- (II) Pour tout élément A de $\mathcal{T}$, son complémentaire $\complement_A \Omega$ appartient à $\mathcal{T}$.
- (III) Pour toute suite $(A_n)_{n \in \mathbb{N}}$ d'éléments de Ω , $\bigcup_{n \in \mathbb{N}} A_n$ appartient à $\mathcal{T}$.

DÉFINITION 19.2. — Si $\mathcal{E}$ est une partie d'un ensemble Ω , la tribu sur Ω engendrée par $\mathcal{E}$ est la plus petite modulo l'inclusion des tribus sur Ω contenant $\mathcal{E}$.

DÉFINITION 19.3. — Si $\mathcal{T}$ est une tribu sur un ensemble Ω , une probabilité sur la tribu $\mathcal{T}$ est une application p de $\mathcal{T}$ dans le segment $[0, 1]$ telle que :

- (i) $p(\Omega) = 1$
- (ii) Pour toute suite $(A_n)_{n \in \mathbb{N}}$ d'éléments deux à deux disjoints de $\mathcal{T}$, la série $\sum_n p(A_n)$ converge et :

$$p\left(\bigcup_{n \in \mathbb{N}} A_n\right) = \sum_{n=0}^{+\infty} p(A_n).$$

Un tel triplé $(\Omega, \mathcal{T}, p)$ est appelé un espace probabilisé.

La tribu sur $\mathbb{R}$ engendrée par l'ensemble des intervalles ouverts est appelée la tribu borélienne¹ et se note $\mathcal{B}(\mathbb{R})$; ses éléments sont les boréliens.

■ Événements indépendants

DÉFINITION 19.4. — a) Les événements A et B sont indépendants si :

$$p(A \cap B) = p(A)p(B).$$

- b) Si n est un entier ≥ 2 , les événements $A_1, A_2, \dots, A_n$ sont indépendants dans leur ensemble si :

$$p(A_1 \cap A_2 \cap \dots \cap A_n) = p(A_1) \times p(A_2) \times \dots \times p(A_n).$$

19.1. Événements deux à deux indépendants qui ne sont pas indépendants dans leur ensemble.

On lance un dé tétraédrique avec quatre faces numérotées de 1 à 4. On considère les événements A obtenir 1 ou 2, B obtenir 1 ou 3 et C obtenir 2 ou 3. Chacun de ces événements a une probabilité de survenir égale à $1/2$. Chacune des intersections $A \cap B$, $A \cap C$ et $B \cap C$ correspond respectivement à obtenir 1, 2 et 3. Ils ont une probabilité de :

$$\frac{1}{4} = \left(\frac{1}{2}\right)^2.$$

Les événements sont donc indépendants deux à deux. Cependant $A \cap B \cap C = \emptyset$ est un événement de probabilité $0 \neq (1/2)^3 = 1/8$. $\square$

Si on se laisse guider par l'intuition de la notion d'indépendance, on pourrait s'attendre à ce qu'un événement indépendant de deux autres eux-mêmes indépendants entre eux, soit indépendant de leur intersection.

19.2. Événement A indépendant d'événements B et C eux-mêmes indépendants entre eux, mais qui n'est pas indépendant de leur intersection.

Nous reprenons l'exemple précédent 19.1. L'événement A est indépendant de B et de C . Or $B \cap C = \{3\}$ donc $p(B \cap C) = 1/4$ et $A \cap B \cap C = \emptyset$, donc :

$$p(A \cap (B \cap C)) = 0 \neq \frac{1}{8} = p(A)p(B \cap C).$$

De plus B et C sont indépendants. $\square$

1. Du nom du mathématicien français Emile Borel (1871-1956).

■ Espérance mathématique et variance

Une expérience aléatoire est en général intéressante car il en découle un bénéfice ou une perte, c'est-à-dire que l'on associe à chaque résultat un nombre réel. Ceci amène à la définition suivante.

DÉFINITION 19.5. — Soit $(\Omega, \mathcal{T}, p)$ un espace probabilisé.

- a) Une variable aléatoire est une application X de Ω dans $\mathbb{R}$ telle que, pour tout borélien B de $\mathbb{R}$, l'image réciproque $X^{-1}(B)$ de B par X appartient à $\mathcal{T}$.
- b) Une variable aléatoire est discrète si l'ensemble $X(\Omega)$ des valeurs prises par X est fini ou dénombrable.
- c) Une densité de probabilité est une fonction f positive, continue par morceaux et intégrable sur $\mathbb{R}$ telle que :

$$\int_{-\infty}^{+\infty} f(t) dt = 1.$$

Une variable aléatoire X admet f pour densité si on a, quels que soient les réels a et b tels que $a \leq b$:

$$p(X \in]a, b]) = \int_a^b f(t) dt.$$

Soit f une densité de probabilité. On pose, quels que soient les nombres réels a et b tels que $a \leq b$:

$$p_f(]a, b]) = \int_a^b f(t) dt.$$

Cette fonction se prolonge de manière unique en une probabilité sur $(\mathbb{R}, \mathcal{B}(\mathbb{R}))$. Il est alors clair que $\text{Id}_{\mathbb{R}}$ est une variable aléatoire de densité f .

Inversement, soit X une variable aléatoire. On définit une probabilité sur $(\mathbb{R}, \mathcal{B}(\mathbb{R}))$ en posant, pour tout borélien B :

$$p_X(B) = p(X^{-1}(B)).$$

La variable aléatoire X définit une probabilité sur $\mathbb{R}$ muni de la tribu borélienne, appelée la loi de probabilité de X .

En fait, on définit souvent une variable aléatoire sans donner explicitement l'expérience aléatoire associée, mais uniquement sa loi de probabilité.

DÉFINITION 19.6. — a) Une variable aléatoire discrète X , admettant pour image l'ensemble $\{x_n \mid n \in I\}$ où I est un ensemble fini ou dénombrable — en général, $I = \mathbb{N}$ ou $[0, m]$ pour un entier naturel m —, admet une espérance mathématique si la série $\sum_n x_n p(X = x_n)$ converge et, si cette série converge, l'espérance mathématique de X est la somme :

$$E(X) = \sum_{n \in I} x_n p(X = x_n).$$

- b) Une variable aléatoire continue X de densité de probabilité f admet une espérance mathématique si la fonction $t \mapsto t f(t)$ est intégrable sur $\mathbb{R}$ et, si c'est le cas, l'espérance mathématique de X est l'intégrale :

$$E(X) = \int_{-\infty}^{+\infty} t f(t) dt$$

L'espérance mathématique peut ne pas exister.

19.3. Variable aléatoire discrète n'ayant pas d'espérance.

Nous considérons la variable aléatoire discrète X , prenant ses valeurs dans l'ensemble $\{2^n \mid n \in \mathbb{N}^*\}$ et définie par $p(X=2^n) = 1/2^n$. On a :

$$\sum_{n=1}^{+\infty} \frac{1}{2^n} = 1$$

donc il s'agit bien d'une loi de probabilité. Cependant $2^n p(X=2^n) = 1$ pour tout entier $n \geq 1$, donc la série $\sum_n 2^n p(X=2^n)$ diverge. Par conséquent X n'a pas d'espérance². $\square$

19.4. Variable aléatoire à densité n'ayant pas d'espérance.

Considérons une variable aléatoire X de densité de probabilité f définie sur $\mathbb{R}$ par l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \longrightarrow \mathbb{R} \\ t \longmapsto f(t) = \frac{1}{\pi(1+t^2)}. \end{array} \right.$$

Comme :

$$f(t) \underset{t \rightarrow +\infty}{\sim} \frac{1}{\pi} \times \frac{1}{t^2}$$

et que f est paire, f est intégrable sur $\mathbb{R}$ par le critère de Riemann³. De plus l'application :

$$\left| \begin{array}{l} F : \mathbb{R} \longrightarrow \mathbb{R} \\ t \longmapsto F(t) = \frac{1}{\pi} \arctan t \end{array} \right.$$

est une primitive de f sur $\mathbb{R}$, donc :

$$\int_{-\infty}^{+\infty} f(t) dt = \frac{1}{\pi} \left(\lim_{+\infty} \arctan - \lim_{-\infty} \arctan \right) = \frac{1}{\pi} \left(\frac{\pi}{2} - \frac{-\pi}{2} \right) = 1.$$

La fonction f est donc une densité de probabilité. Cependant :

$$t f(t) \underset{t \rightarrow +\infty}{\sim} \frac{1}{\pi} \times \frac{1}{t}$$

donc la fonction $t \mapsto t f(t)$ n'est pas intégrable sur $\mathbb{R}$: l'espérance de X n'est pas définie. $\square$

DÉFINITION 19.7. — La variance $V(X)$ d'une variable aléatoire X est, si elle existe, l'espérance de la variable aléatoire $(X - E(X))^2$.

On remarque que, pour avoir une variance, une variable aléatoire doit déjà avoir une espérance. La variance de X existe si, et seulement si, $E(X^2)$ existe ; on a alors dans ce cas $V(X) = E(X^2) - (E(X))^2$.

2. Ce résultat, connu sous le nom de paradoxe de Saint-Petersbourg, a été énoncé par Nicolas Bernoulli (1695-1726) sous la forme suivante : Un joueur mise une somme d'argent pour avoir le droit de participer au jeu suivant. On lance une pièce de monnaie. Si la pièce tombe sur face, le jeu s'arrête, sinon le joueur gagne un euro et rejoue. Les lancers suivants sont identiques mais le gain éventuel du joueur double de valeur chaque fois. Quelle mise maximale le joueur doit-il accepter pour participer à ce jeu ? Comme l'espérance mathématique est infinie, il devrait logiquement être d'accord pour n'importe quelle somme !

3. La loi de probabilité définie par cette fonction est connue sous le nom de loi de Cauchy.

19.5. Variable aléatoire discrète ayant une espérance mais n'ayant pas de variance.

Nous introduisons la suite $(a_n)_{n \geq 1}$ de terme général $a_n = \frac{4}{n(n+1)(n+2)}$.

Nous considérons la variable aléatoire discrète X , prenant ses valeurs dans $\mathbb{N}^*$, définie par $p(X=n) = a_n$. On a, pour tout entier $n \geq 1$:

$$a_n = \frac{2}{n} - \frac{4}{n+1} + \frac{2}{n+2}.$$

On en déduit, en décalant les indices, que, pour tout entier $n \geq 1$:

$$\sum_{k=1}^n a_k = \sum_{k=1}^n \frac{2}{k} - \sum_{k=2}^{n+1} \frac{4}{k} + \sum_{k=3}^{n+2} \frac{2}{k} = 1 - \frac{4}{n+1} + \frac{2}{n+1} + \frac{2}{n+2},$$

terme général d'une suite qui converge vers 1. Nous avons ainsi défini une loi de probabilité. Comme :

$$na_n \underset{n \rightarrow \infty}{\sim} \frac{4}{n^2}$$

la série $\sum(na_n)$ converge, donc X possède une espérance. Cependant :

$$n^2 a_n \underset{n \rightarrow +\infty}{\sim} \frac{4}{n},$$

terme général d'une série divergente, donc $E(X^2)$ n'est pas défini, ce qui montre que X ne possède pas de variance. $\square$

19.6. Variable aléatoire à densité ayant une espérance mais n'ayant pas de variance.

Nous considérons la variable aléatoire X , à valeurs dans $\mathbb{R}^+$, de densité de probabilité f définie par l'application :

$$\left| \begin{array}{l} f : \mathbb{R} \rightarrow \mathbb{R} \\ t \mapsto f(t) = \begin{cases} 0 & \text{si } t < 0, \\ \frac{2}{(1+t)^3} & \text{si } t \geq 0. \end{cases} \end{array} \right.$$

On vérifie que :

$$\int_0^{+\infty} f(t) dt = \left[-\frac{1}{(1+t)^2} \right]_{t=0}^{+\infty} = 1,$$

et comme $f \geq 0$, c'est bien une densité de probabilité. L'espérance de X existe et :

$$E(X) = \int_0^{+\infty} t f(t) dt = \int_0^{+\infty} \left(\frac{2}{(1+t)^2} - \frac{2}{(1+t)^3} \right) dt = 2 - 1 = 1.$$

Or :

$$t^2 f(t) \underset{t \rightarrow +\infty}{\sim} \frac{2}{t}$$

donc la fonction $t \mapsto t^2 f(t)$ n'est pas intégrable sur $[0, +\infty[$. On en déduit que X ne possède pas de variance. $\square$

DÉFINITION 19.8. — Des variables aléatoires X et Y sont indépendantes si, quels que soient les éléments A et B de la tribu borélienne $\mathcal{B}(\mathbb{R})$:

$$p((X \in A) \cap (Y \in B)) = p(X \in A)p(Y \in B).$$

DÉFINITION 19.9. — La covariance d'un couple (X, Y) de variables aléatoires est le nombre réel :

$$\text{cov}(X, Y) = E((X - E(X))(Y - E(Y))) = E(XY) - E(X)E(Y).$$

Des variables aléatoires ne sont pas corrélées si leur covariance est nulle.

Des variables aléatoires indépendantes ne sont pas corrélées, mais la réciproque est fausse.

19.7. Variables aléatoires qui ne sont ni indépendantes, ni corrélées.

Nous considérons une variable aléatoire X , prenant seulement trois valeurs, et définie par $p(X = 1) = p(X = 0) = p(X = -1) = 1/3$, et nous posons $Y = X^2$. Alors $E(X) = (-1) \times (1/3) + 0 \times (1/3) + 1 \times (1/3) = 0$. Comme X ne prend que des valeurs x telles que $x^3 = x$, on a $XY = X^3 = X$ donc $E(XY) = E(X) = 0$. On en déduit que $E(XY) = E(X)E(Y) = 0$ ce qui prouve que les variables aléatoires X et Y ne sont pas corrélées. De plus $X = 1$ entraîne $Y = X^2 = 1$, donc $(X = 1) \cap (Y = 1) = (X = 1)$, d'où l'on déduit que :

$$p((X = 1) \cap (Y = 1)) = \frac{1}{3} \text{ et } p(X^2 = 1) = p(X = 1) + p(X = -1) = \frac{2}{3}.$$

Ainsi :

$$p(X = 1)p(Y = 1) = \frac{1}{3} \times \frac{2}{3} = \frac{2}{9} \neq \frac{1}{3} = p((X = 1) \cap (Y = 1)),$$

ce qui montre que les variables aléatoires X et Y ne sont pas indépendantes. $\square$

■ Convergence des suites de variables aléatoires

Lorsqu'une expérience aléatoire se reproduit indéfiniment, on est conduit à introduire une suite de variables aléatoires $(X_n)_{n \in \mathbb{N}}$. L'étude du comportement de cette suite lorsque n augmente nous amène à considérer la convergence de ce type de suite. En fonction de ce que l'on étudie, différents modes de convergence sont utilisés. Nous montrerons qu'ils ne sont jamais équivalents.

Dans les définitions qui suivent, les variables aléatoires sont définies sur un espace probabilisé $(\Omega, \mathcal{T}, p)$.

DÉFINITION 19.10. — a) Une suite (X_n) de variables aléatoires converge presque sûrement vers la variable aléatoire X si :

$$p\left(\left\{\omega \in \Omega \mid \lim_{n \rightarrow \infty} X_n(\omega) = X(\omega)\right\}\right) = 1.$$

b) Une suite (X_n) de variables aléatoires converge en probabilité vers la variable aléatoire X si, pour tout nombre réel $\varepsilon > 0$:

$$\lim_{n \rightarrow \infty} p(|X_n - X| > \varepsilon) = 0.$$

c) Une suite (X_n) de variables aléatoires converge en moyenne vers la variable aléatoire X si :

$$\lim_{n \rightarrow \infty} E(|X_n - X|) = 0,$$

et elle converge en moyenne quadratique vers la variable aléatoire X si :

$$\lim_{n \rightarrow \infty} E((X_n - X)^2) = 0.$$

DÉFINITION 19.11. — La fonction de répartition F_X d'une variable aléatoire X est l'application :

$$\left| \begin{array}{l} F_X : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto F_X(x) = p(X < x). \end{array} \right.$$

Si X est une variable aléatoire, sa fonction de répartition F_X est une application croissante de $\mathbb{R}$ dans $\mathbb{R}$, $\lim_{x \rightarrow -\infty} F_X(x) = 0$ et $\lim_{x \rightarrow +\infty} F_X(x) = 1$.

DÉFINITION 19.12. — Soit $(X_n)_{n \in \mathbb{N}}$ une suite de variables aléatoires, X une variable aléatoire, F la fonction de répartition de X et, pour tout $n \in \mathbb{N}$, F_n celle de X_n . La suite (X_n) converge en loi vers X si, en tout point t où F est continue :

$$\lim_{n \rightarrow \infty} F_n(t) = F(t).$$

Remarquons que cette dernière définition ne dépend que des lois des variables aléatoires. Pour la convergence en loi, il est important d'affirmer que la suite (F_n) des fonctions de répartition converge vers une fonction de répartition F comme le montre l'exemple suivant.

19.8. Suite (F_n) de fonctions de répartition convergeant vers une fonction F qui n'est pas une fonction de répartition.

Nous considérons la suite $(F_n)_{n \in \mathbb{N}}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} F_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto F_n(x) = \begin{cases} 0 & \text{si } x \leq n, \\ x - n & \text{si } x \in]n, n+1[, \\ 1 & \text{si } x \geq n+1. \end{cases} \end{array} \right.$$

Il est clair que, pour tout entier naturel n , F_n est la fonction de répartition d'une variable aléatoire suivant la loi uniforme sur $[n, n+1]$. Pourtant, pour tout nombre réel x , $F_n(x) = 0$ dès que $n \geq N_x + 1$ où N_x est la partie entière de x , ce qui montre que la suite $(F_n(x))_{n \in \mathbb{N}}$ converge vers 0. Par conséquent la suite (F_n) converge simplement vers l'application nulle, qui n'est pas une fonction de répartition. $\square$

Si X est une variable aléatoire à densité f et si, pour tout n , X_n est une variable aléatoire à densité f_n , la convergence simple la suite de (f_n) vers f entraîne celle de (F_{X_n}) vers F_X , donc entraîne la convergence en loi de la suite (X_n) vers X . La réciproque est fausse.

19.9. Suite (X_n) où X_n est, pour tout n , une variable aléatoire à densité (f_n) , qui converge en loi vers une variable aléatoire X à densité f , alors que la suite (f_n) ne converge pas vers f .

Nous introduisons la suite $(F_n)_{n \in \mathbb{N}^*}$ d'applications de $\mathbb{R}$ dans $\mathbb{R}$, de terme général :

$$\left| \begin{array}{l} F_n : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto F_n(x) = \begin{cases} 0 & \text{si } x < 0, \\ x - \frac{\sin(2n\pi x)}{2n\pi} & \text{si } x \in [0, 1], \\ 1 & \text{si } x > 1. \end{cases} \end{array} \right.$$

Soit $n \in \mathbb{N}^*$. Comme $F_n(0) = 0$ et $F_n(1) = 0$, F_n est continue sur les intervalles fermés $]-\infty, 0]$, $[0, 1]$ et $[1, +\infty[$, donc F_n est continue sur $\mathbb{R}$. De plus F_n est dérivable sur $]0, 1[$ et, pour tout point x de $]0, 1[$, $F_n'(x) = 1 - \cos(2n\pi x) \geq 0$, donc F_n est une application croissante de $\mathbb{R}$ dans $\mathbb{R}$. Ainsi F_n est la fonction de répartition d'une variable aléatoire X_n à densité f_n , où $f_n(x) = 0$ pour $x \in \mathbb{R} \setminus [0, 1]$ et $f_n(x) = 1 - \cos(2n\pi x)$ pour $x \in [0, 1]$. La suite (F_n) converge simplement vers l'application :

$$\left| \begin{array}{l} F : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto F(x) = \begin{cases} 0 & \text{si } x < 0, \\ x & \text{si } x \in [0, 1], \\ 1 & \text{si } x > 1 \end{cases} \end{array} \right.$$

et F est la fonction de répartition de la loi uniforme sur $[0, 1]$. Cependant, pour un point x de $]0, 1[$, la suite $(f_n(x))$ n'admet pas de limite. $\square$

Les différents types de convergence ne sont pas équivalents. Seules les implications du diagramme ci-dessous, dans lequel CV est une abréviation de «convergence», sont vraies⁴ :

$$\begin{array}{c} (\text{CV presque sûre}) \implies (\text{CV en probabilité}) \implies (\text{CV en loi}) \\ \uparrow \\ (\text{CV en moyenne quadratique}) \implies (\text{CV en moyenne}) \end{array}$$

Les exemples qui suivent montrent qu'aucune autre implication n'est vraie, excepté celles obtenues par transitivité.

19.10. Suite de variables aléatoires qui converge en loi mais ne converge pas en probabilité.

Considérons une suite (X_n) de variables aléatoires deux à deux indépendantes et suivant une loi de Bernoulli de paramètre p . Elles ont la même fonction de répartition, donc (X_n) converge en loi vers une variable aléatoire de Bernoulli X .

On a, quels que soient les entiers naturels distincts m et n :

$$p\left(\left\{|X_n - X_m| > \frac{1}{2}\right\}\right) = p\left(\left(\{X_n = 1\} \cap \{X_m = 0\}\right) \cup \left(\{X_n = 0\} \cap \{X_m = 1\}\right)\right) = 2p(1-p).$$

Soit X une variable aléatoire quelconque. On a :

$$\left\{|X_n - X_m| > \frac{1}{2}\right\} \subset \left\{|X_n - X| > \frac{1}{4}\right\} \cup \left\{|X_m - X| > \frac{1}{4}\right\}.$$

Si la suite (X_n) convergeait en probabilité vers X , les limites à l'infini de :

$$p\left(\left\{|X_n - X| > \frac{1}{4}\right\}\right) \text{ et } p\left(\left\{|X_m - X| > \frac{1}{4}\right\}\right)$$

seraient nulles, donc *a fortiori* celle de :

$$p\left(\left\{|X_n - X_m| > \frac{1}{2}\right\}\right),$$

ce qui contredit le fait que cette dernière probabilité est égale à $2p(1-p) > 0$. $\square$

4. Voir [DELM], chapitre V, § V.1 et § V.2.

Pour la convergence en probabilité, en moyenne ou presque sûre, il y a équivalence entre la convergence de (X_n) vers X et celle de $(X_n - X)$ vers 0. C'est faux pour la convergence en loi comme le montre l'exemple suivant.

19.11. Suite (X_n) qui converge en loi vers X alors que $(X_n - X)$ ne converge pas en loi vers 0.

Nous considérons une variable aléatoire X définie par $p(X=1) = p(X=-1) = 1/2$ et $p(X=a) = 0$ si $a \notin \{-1, 1\}$ et nous posons $X_n = -X$ pour tout $n \in \mathbb{N}$. Il est clair que X_n suit, pour tout n , la même loi que X , donc la suite (X_n) converge en loi vers X . En revanche $X_n - X = -2X$ ne dépend pas de l'entier naturel n et n'est pas nulle, donc $(X_n - X)$ ne tend vers 0. $\square$

19.12. Suite de variables aléatoires à densité qui converge en loi vers une variable aléatoire discrète.

Nous choisissons une application F de $\mathbb{R}$ dans $\mathbb{R}$ croissante et de classe $\mathcal{C}^1$ telle que $\lim_{x \rightarrow -\infty} F(x) = 0$ et $\lim_{x \rightarrow +\infty} F(x) = 1$.

Citons par exemple l'application $F : x \mapsto F(x) = \frac{1}{2} + \frac{1}{\pi} \arctan x$ de $\mathbb{R}$ dans $\mathbb{R}$.

Pour tout entier $n \geq 1$ nous introduisons l'application :

$$\left| \begin{array}{l} F_n : \mathbb{R} \longrightarrow \mathbb{R} \\ t \longmapsto F_n(t) = F(nt) \end{array} \right.$$

et nous appelons X_n une variable aléatoire de fonction de répartition F_n . Alors la suite $(F_n(t))$ converge vers 1 pour tout réel $t > 0$ et vers 0 pour tout réel $t < 0$. Ainsi la suite de fonction (F_n) converge simplement vers l'application :

$$\left| \begin{array}{l} G : \mathbb{R} \longrightarrow \mathbb{R} \\ t \longmapsto G(t) = \begin{cases} 0 & \text{si } t < 0, \\ F(0) & \text{si } t = 0, \\ 1 & \text{si } t > 0. \end{cases} \end{array} \right.$$

Excépté en 0 où elle n'est pas continue, la fonction G est égale à la fonction de répartition d'une variable aléatoire qui prend presque sûrement la valeur 0. $\square$

19.13. Suite de variables aléatoires discrètes qui converge en loi vers une variable aléatoire à densité.

Nous notons E la fonction «partie entière» et nous associons, à tout entier $n \geq 1$, une variable aléatoire X_n suivant une loi uniforme sur :

$$A_n = \left\{ \frac{1}{n}, \frac{2}{n}, \dots, 1 \right\}.$$

Pour tout entier $n \geq 1$, la fonction de répartition de X_n est l'application :

$$\left| \begin{array}{l} F_n : \mathbb{R} \longrightarrow \mathbb{R} \\ t \longmapsto F_n(t) = \begin{cases} 0 & \text{si } x < 0, \\ \frac{1}{n} E(nx) & \text{si } x \in [0, 1], \\ 1 & \text{si } x > 1. \end{cases} \end{array} \right.$$

Soit x un nombre réel. On a, pour tout entier $n \geq 1$, $nx - 1 < E(nx) \leq nx$, donc

$x - (1/n) < F_n(x) \leq x$. Il en résulte que $\lim_{(n \rightarrow \infty)} F_n(x) = x$. La suite (F_n) converge donc simplement sur $\mathbb{R}$ vers l'application :

$$F : \mathbb{R} \longrightarrow \mathbb{R} \\ x \longmapsto F(x) = \begin{cases} 0 & \text{si } x < 0, \\ x & \text{si } x \in [0, 1], \\ 1 & \text{si } x > 1. \end{cases}$$

En dérivant pour $x \neq 0$ et $x \neq 1$, on reconnaît la densité d'une variable aléatoire suivant une loi uniforme sur $[0, 1]$. $\square$

19.14. Suite de variables aléatoires qui converge en probabilité mais ne converge pas presque sûrement.

Nous introduisons la suite $(a_n)_{n \geq 1}$ de terme général $a_n = \frac{n(n-1)}{2}$.

Pour tout entier $n \geq 1$, $a_{n+1} = a_n + n$ et $a_n = \sum_{i=1}^{n-1} i$. La suite (a_n) est strictement croissante et tend vers $+\infty$. Par conséquent, pour tout $\ell \in \mathbb{N}$, il existe un couple (n, q) et un seul tel que $\ell = a_n + q$, $n \in \mathbb{N}^*$, $q \in \mathbb{N}$ et $0 \leq q < n$, et nous associons alors à ℓ l'intervalle :

$$I_\ell = \left[\frac{q}{n}, \frac{q+1}{n} \right[.$$

On a $0 = a_1 + 0$ et $I_0 = [0, 1[$, $1 = a_2 + 0$ et $I_1 = [0, 1/2[$, $2 = a_2 + 1$ et $I_2 = [1/2, 1[$, $3 = a_3 + 0$ et $I_3 = [0, 1/3[$, $4 = a_3 + 1$ et $I_4 = [1/3, 2/3[$, etc.

Nous considérons l'expérience aléatoire de tirer au hasard un élément de $I_0 = [0, 1[$. Plus précisément, nous munissons I_0 de sa tribu borélienne et de la mesure de Lebesgue induite, notée m . Ceci définit un espace probabilisé puisque $m(I_0) = 1$.

A tout entier naturel ℓ , que l'on écrit $\ell = a_n + q$ où $n \in \mathbb{N}^*$, $q \in \mathbb{N}$ et $0 \leq q < n$, nous associons la variable aléatoire X_ℓ définie par : $X_\ell(\omega) = 1$ si $\omega \in I_\ell$ et $X_\ell(\omega) = 0$ si $\omega \in I_1 \setminus I_\ell$; alors X_ℓ suit une loi de Bernoulli de paramètre $m(I_\ell) = 1/n$.

Quand ℓ tend vers l'infini, n aussi, donc $p(X_\ell = 1) = p(X_\ell \neq 0) = 1/n$, qui tend vers 0 quand ℓ , et donc n , tend vers l'infini. On voit ainsi, en désignant par Z la variable aléatoire nulle que, pour tout réel $\varepsilon > 0$, $p(|X_\ell - Z| > \varepsilon)$ tend vers 0 quand ℓ tend vers l'infini. La suite (X_ℓ) converge donc en probabilité vers $Z = 0$. Soit $\omega \in [0, 1[$. Nous prouvons que la suite $(X_\ell(\omega))$ n'admet pas de limite.

Soit n un entier ≥ 1 . Nous notons q_n la partie entière de $n\omega$; alors $\omega \in \left[\frac{q_n}{n}, \frac{1+q_n}{n} \right[$.

Nous posons $\ell_n = a_n + q_n$, d'où l'on déduit que $\omega \in I_{\ell_n}$ donc que $X_{\ell_n}(\omega) = 1$, nous choisissons $r_n \in \llbracket 0, n-1 \rrbracket$ tel que $r_n \neq q_n$ (c'est possible dès que $n \geq 2$) et nous posons $m_n = \ell_n + r_n$; alors $\omega \notin I_{m_n}$ donc $X_{m_n}(\omega) = 0$.

Pour tout n , $a_n \leq \ell_n < a_{n+1}$ et $a_n \leq m_n < a_{n+1}$, donc les suites (ℓ_n) et (r_n) sont clairement strictement croissantes. Comme les sous-suites $(X_{\ell_n}(\omega))$ et $(X_{m_n}(\omega))$ de la suite $(X_\ell(\omega))$ tendent respectivement vers 1 et 0, la suite $(X_\ell(\omega))$ n'admet pas de limite. Il n'y a donc pas convergence presque sûre de la suite (X_ℓ) . $\square$

19.15. Suite de variables aléatoires qui converge en moyenne quadratique — donc en moyenne — mais ne converge pas presque sûrement.

Nous reprenons les hypothèses et les notations de l'exemple précédent 19.14. Pour tout $\ell \in \mathbb{N}^*$, X_ℓ suit une loi de Bernoulli de paramètre $1/n$ où (n, q) est le couple tel que $\ell = a_n + q$, $n \in \mathbb{N}^*$, $q \in \mathbb{N}$ et $0 \leq q < n$. Nous démontrons que

la suite (X_ℓ) converge vers $X = 0$ en moyenne et en moyenne quadratique.

Pour tout entier naturel ℓ , $E(X_\ell - X) = E(X_\ell) = 1/n$, qui tend vers 0 quand ℓ , donc n , tend vers l'infini. Par conséquent (X_ℓ) tend vers X en moyenne. Pour tout $\ell \in \mathbb{N}$, X_ℓ ne prend pour valeurs que 0 et 1, donc $(X_\ell - X)^2 = X_\ell^2 = X_\ell$, ce qui montre que $E((X_\ell - X))^2 = E(X_\ell) = 1/n$, qui tend vers 0 quand ℓ , donc n , tend vers l'infini. Ainsi (X_ℓ) tend vers X en moyenne quadratique⁵. $\square$

19.16. Suite de variables aléatoires convergeant presque sûrement mais qui ne converge pas en moyenne — donc en moyenne quadratique.

Comme dans les exemples précédents 19.14 et 19.15, nous considérons l'expérience aléatoire de tirer au hasard un élément de $I_0 = [0, 1[$. Plus précisément, nous munissons I_0 de sa tribu borélienne et de la mesure de Lebesgue induite, notée m . Ceci définit un espace probabilisé puisque $m(I_0) = 1$.

Nous associons à tout entier $n \geq 1$ la variable aléatoire X_n telle que $X_n(\omega) = n$ si $\omega \in]0, 1/n[$ et $X_n(\omega) = 0$ sinon. On a $X_n(0) = 0$ pour tout n et, si $\omega \in]0, 1[$, il existe un entier $N \geq 1$ tel que $0 < 1/N < \omega$, et alors $X_n(\omega) = 0$ pour tout entier $n \geq N$. On en déduit que (X_n) converge presque sûrement vers 0.

Par ailleurs, pour tout entier $n \geq 1$:

$$E(X_n) = 0 \times \left(1 - \frac{1}{n}\right) + n \times \frac{1}{n} = 1$$

qui ne tend pas vers 0. On en déduit que la suite (X_n) ne tend pas vers 0 en moyenne. De même, pour tout entier $n \geq 1$:

$$E(X_n^2) = 0 \times \left(1 - \frac{1}{n}\right) + n^2 \times \frac{1}{n} = n$$

qui tend vers $+\infty$. On en déduit de même que la suite (X_n) ne tend pas vers 0 en moyenne quadratique. $\square$

19.17. Suite de variables aléatoires qui converge en moyenne mais ne converge pas en moyenne quadratique.

Nous reprenons l'exemple précédent 19.16 et nous posons, pour tout entier $n \geq 1$, $Y_n(\omega) = \sqrt{n}$ si $\omega \in]0, 1/n[$ et 0 sinon. Alors, pour tout entier $n \geq 1$:

$$E(X_n) = 0 \times \left(1 - \frac{1}{n}\right) + \sqrt{n} \times \frac{1}{n} = \frac{1}{\sqrt{n}}$$

donc la suite $(E(X_n))$ converge vers 0. En revanche, pour tout entier $n \geq 1$:

$$E(X_n^2) = 0 \times \left(1 - \frac{1}{n}\right) + n \times \frac{1}{n} = 1.$$

On en déduit que (X_n) tend vers 0 en moyenne mais ne tend pas vers 0 en moyenne quadratique. $\square$

5. On aurait pu se dispenser, dans l'exemple précédent 19.14, de prouver la convergence en probabilité de (X_ℓ) vers 0 ; en effet, la convergence en moyenne quadratique entraîne la convergence en probabilité.



Bertrand Hauchecorne enseigne les mathématiques en classes préparatoires au lycée Pothier d'Orléans. Il a déjà publié plusieurs ouvrages parmi lesquels *Des mathématiciens de A à Z* avec Daniel Suratteau et *Les mots et les maths*.

À l'aide de plus de 500 contre-exemples choisis dans tous les domaines des mathématiques, cet ouvrage montre, au-delà de ses côtés divertissants, la valeur mathématique et la vertu pédagogique du contre-exemple.

Cette nouvelle édition est très largement enrichie. L'aspect définitions et théorèmes a pris du corps ; de nombreux graphiques, des références bibliographiques et des notes historiques ont été ajoutés. En outre l'amélioration de la qualité de la mise en page et de l'impression facilite sa lecture.

Cet ouvrage permettra aux étudiants d'approfondir l'enseignement de mathématiques qu'ils reçoivent, à ceux qui préparent le concours du CAPES ou de l'agrégation d'enrichir une leçon, aux enseignants de trouver des thèmes d'exercices ou de problèmes. Plus généralement il intéressera tous ceux qui veulent approfondir leur réflexion sur les notions de définition, d'hypothèse ou de théorème. Il apportera surtout bien du plaisir à ceux dont la curiosité mathématique est toujours en éveil.



Mathieu

www.editions-ellipses.fr