

Analyse appliquée
mat 2466

Yvan SAINT-AUBIN
université de montréal
septembre 2021

*M. Cauchy annonce que, pour se conformer au vœu du Conseil,
il ne s'attachera plus à donner, comme il a fait jusqu'à présent,
des démonstrations parfaitement rigoureuses.*

Conseil d'instruction de l'École Polytechnique,
24 novembre 1825

Préface

Quels étaient les motifs du Conseil de l'Éducation pour modérer les efforts du Baron Cauchy vers la parfaite rigueur mathématique ? Le Conseil doutait-il de la nécessité de cette rigueur ? de son caractère pédagogique ? ou de sa validité même ? Chacun voudra y aller de ses hypothèses.

Il est clair que le débat est encore ouvert puisqu'on enseigne aujourd'hui les mathématiques à tous les niveaux de rigueur. Des ensembles de "recettes" que constituent plusieurs cours aux disciplines utilisant les mathématiques aux enseignements les plus abstraits où le moindre petit lemme est justifié : on en trouve pour tous les goûts. Il est cependant rare de trouver des cours où les divers chapitres sont traités avec des niveaux de rigueur différents. Les présentes notes ont choisi ce parti. Elles oscillent entre une présentation rigoureuse de résultats mathématiques et une discussion intuitive des outils et applications sous-jacents au sujet. Le premier "style" sera aisément reconnu avec ses définitions, théorèmes et preuves. Dans le second, j'ai quand même essayé de mettre en relief les hypothèses cruciales et les étapes qui demanderaient des justifications moins naïves.

Ce choix pourra déplaire à plusieurs : aux futurs mathématiciens qui désiraient voir la discipline sérieusement et une fois pour toutes et aux physiciens, par exemple, qui aimeraient plutôt se concentrer sur les applications spectaculaires de l'analyse de Fourier. J'espère cependant qu'elle pourra rallier quelques lecteurs. Les théorèmes qui sont démontrés le sont car ils forment le coeur du sujet et qu'ils indiquent la voie à venir. Et le nombre des exemples traités plus rapidement révèlent la richesse du domaine d'applications.

Le matériel présenté dans ces notes se trouve dans plusieurs livres. Je dois cependant payer tribut à une source en particulier, Thomas W. Körner, dont le livre *Fourier Analysis* paru aux Presses de l'Université de Cambridge, constitue un chef-d'oeuvre de pédagogie. La présentation mathématique est tirée de ce livre. Enfin c'est avec plaisir que je remercie Ervig Lapalme : il a lu ces notes et fait les exercices. Ses nombreux commentaires et suggestions ont été grandement appréciés.

Yvan Saint-Aubin (1996)

Je remercie Simon St-Amant et Maxime Fortier Bourque pour leurs suggestions et les coquilles qu'ils ont relevées.

Y.S.-A. (2024)

Table des matières

Préface	v
Table des matières	vii
1 Séries de Fourier	1
1.1 Définition et exemples	1
1.2 Séries d'une parité donnée	9
1.3 Séries de période L	15
1.4 Séries complexes	18
1.5 La corde vibrante	24
1.5.1 La corde plombée	24
1.5.2 La corde vibrante	30
2 Convergence des séries de Fourier	39
2.1 Rappel d'analyse	39
2.1.1 La continuité	40
2.1.2 La convergence d'une suite et d'une série	41
2.1.3 L'intégrale de Riemann	43
2.1.4 Convergence uniforme	45
2.2 Le théorème de Fejér	47
2.3 Convergence ponctuelle	58
2.4 Rappel d'algèbre linéaire	63
2.5 Convergence en moyenne	72
3 Problème de Sturm-Liouville et fonctions spéciales	85
3.1 Rappel sur les équations différentielles	85
3.2 Séparation de variables	90
3.3 Le problème de Sturm-Liouville	98
3.4 Polynômes de Legendre	105
3.5 Fonctions de Bessel	123

4	La transformation de Fourier	145
4.1	Un passage à la limite et exemples de calcul	145
4.2	Quelques théorèmes	157
4.3	La convolution	169
4.4	Une application : l'âge de la terre selon Lord Kelvin	177
A	Fonctions eulériennes	187
A.1	Les fonctions Γ et $B(p, q)$	187
A.2	La formule de Stirling	194
	Bibliographie	201
	Solutions et suggestions pour les exercices	203
	Index	209

Chapitre 1

Séries de Fourier

A la fin du XVIII^e siècle, la mécanique classique était sur de bonnes assises et les physiciens commençaient à se tourner vers des problèmes d'une difficulté plus grande. Les équations de Newton régissant la mécanique et, en particulier, la gravitation, mènent à des équations différentielles ordinaires, c'est-à-dire à des équations ne faisant intervenir que des dérivées par rapport à une seule variable, le temps. Les équations qui attirent les physiciens en cette fin de siècle sont des équations aux dérivées partielles qui font intervenir des dérivées par rapport à plusieurs variables indépendantes, telles les équations de la chaleur, de la lumière ou des ondes, de l'électricité et du magnétisme, etc. La naissance du domaine des mathématiques appelé « séries de Fourier » ou encore « analyse harmonique » peut être fixée à 1807, date à laquelle Fourier présente son mémoire à l'Académie des Sciences de Paris sur la théorie de la chaleur. Plusieurs noms sont également attachés à la naissance de ce chapitre des mathématiques : D. Bernoulli, Euler, d'Alembert, Dirichlet, Lagrange, etc.

1.1 Définition et exemples

Dans cette première section, nous introduisons les séries de Fourier et en présentons des exemples simples. Leur utilité est restreinte aux fonctions périodiques.

Soit $f : \mathbb{R} \rightarrow \mathbb{R}$ une fonction définie sur les nombres réels et à valeurs réelles. S'il existe un nombre réel $L > 0$ tel que $f(x + L) = f(x)$ pour tout $x \in \mathbb{R}$, alors la fonction f est dite périodique de période L . Remarquons que si f est périodique de période L , alors elle est également de période $2L, 3L, \dots$. Le plus petit L tel que f soit périodique de période L est la *période* de f .

Les séries de Fourier permettent d'écrire les fonctions périodiques sous la forme d'une somme infinie. Ecrire une fonction comme une somme infinie ne devrait pas vous être inconnu. En effet vous connaissez déjà les séries des puissances d'une

variable

$$f(x) = \sum_{n=0}^{\infty} \frac{f^{(n)}(b)}{n!} (x-b)^n$$

qui caractérisent le comportement d'une fonction f autour d'un point b . Par exemple

$$\sin x = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \dots = \sum_{n=1}^{\infty} \frac{(-1)^{n+1} x^{2n-1}}{(2n-1)!}$$

est le développement en série de Taylor de la fonction sinus autour du point $x = 0$. Les séries de Taylor exigent que les fonctions soient C^∞ dans un voisinage du point autour duquel le développement est fait, c'est-à-dire que les dérivées de tous les ordres existent, et que les termes de la somme tendent vers 0 suffisamment vite pour que la série ait un sens. Même avec ces conditions le rayon de convergence peut être 0. Parmi les grands avantages des séries de Taylor est que les premiers termes ont une interprétation géométrique : le premier terme $f(b)$ est la valeur de f en b , le terme $f^{(1)}(b) = f'(b)$ est la pente de f en b , le troisième terme $f''(b)$ caractérise la parabole en b , etc.

Joseph, baron Fourier
(1768-1830)

Fourier ne cherchait pas une meilleure méthode d'approximer les fonctions. Il cherchait à trouver des solutions à l'équation différentielle qui décrit la chaleur. Pour une raison qui jouera un rôle important par la suite, il fut amené à considérer les sommes de fonctions trigonométriques dont la période est un multiple de 2π :

$$\frac{a_0}{2} + \sum_{m=1}^{\infty} (a_m \cos mx + b_m \sin mx).$$

Ces sommes, si elles ont un sens, c'est-à-dire si elles convergent, sont aussi de période 2π . Supposons en fait que ces sommes convergent uniformément vers une fonction f :

$$f(x) = \frac{a_0}{2} + \sum_{m=1}^{\infty} (a_m \cos mx + b_m \sin mx).$$

La convergence uniforme des séries assure la continuité de f et permet d'exprimer les nombres a_m , $m \in \mathbb{N} \cup \{0\}$ et b_m , $m \in \mathbb{N}$ en termes de la fonction f .¹ (Notons que les a_m et b_m jouent ici le rôle des $f^{(n)}(b)$ dans le développement en série de Taylor : ce sont les *coefficients* du développement en série de Fourier.) Nous allons tout d'abord calculer l'intégrale suivante :

$$\begin{aligned} \int_{-\pi}^{\pi} f(x) \cos nx \, dx &= \int_{-\pi}^{\pi} \frac{a_0}{2} \cos nx \, dx \\ &+ \sum_{m=1}^{\infty} \int_{-\pi}^{\pi} (a_m \cos mx \cos nx + b_m \sin mx \cos nx) \, dx. \end{aligned}$$

1. Dans ces notes l'ensemble $\mathbb{N}$ est l'ensemble $\mathbb{N} = \{1, 2, 3, \dots\}$ des entiers positifs et ne contient pas le nombre 0.

Pour simplifier les intégrales au membre de droite, les identités trigonométriques suivantes seront utiles

$$\begin{aligned}\cos((m \pm n)x) &= \cos mx \cos nx \mp \sin mx \sin nx \\ \sin((m \pm n)x) &= \pm \cos mx \sin nx + \sin mx \cos nx\end{aligned}$$

et donc

$$\begin{aligned}\cos mx \cos nx &= \frac{1}{2}(\cos((m+n)x) + \cos((m-n)x)) \\ \sin mx \cos nx &= \frac{1}{2}(\sin((m+n)x) + \sin((m-n)x)) \\ \sin mx \sin nx &= \frac{1}{2}(\cos((m-n)x) - \cos((m+n)x)).\end{aligned}$$

L'intégrale ci-dessus devient ainsi

$$\begin{aligned}&= \frac{a_0}{2} \int_{-\pi}^{\pi} \cos nx \, dx + \sum_{m=1}^{\infty} \left\{ \int_{-\pi}^{\pi} \frac{1}{2} a_m (\cos((m+n)x) + \cos((m-n)x)) \, dx \right. \\ &\quad \left. + \int_{-\pi}^{\pi} \frac{1}{2} b_m (\sin((m+n)x) + \sin((m-n)x)) \, dx \right\}\end{aligned}$$

Notons maintenant que

$$\int_{-\pi}^{\pi} \cos nx \, dx = \int_{-\pi}^{\pi} \sin nx \, dx = 0, \quad \text{si } n \geq 1,$$

et

$$\int_{-\pi}^{\pi} \cos 0x \, dx = \int_{-\pi}^{\pi} dx = 2\pi, \quad \text{si } n = 0.$$

Donc l'intégrale peut être simplifiée en

$$\begin{aligned}&= \frac{a_0}{2} \delta_{n,0} 2\pi + \frac{1}{2} \sum_{m=1}^{\infty} 2\pi \delta_{m-n,0} a_m \\ &= a_n \pi\end{aligned}$$

où le delta de Kronecker $\delta_{i,j}$ a été utilisé. Ce delta est une fonction de deux entiers i et j qui vaut 1 si ces entiers coïncident et vaut 0 autrement. Donc :

$$a_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \cos nx \, dx, \quad n \geq 0.$$

De la même façon on trouve :

$$b_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \sin nx \, dx, \quad n \geq 1.$$

Nous venons donc de montrer que si

séries de Fourier

$$f(x) = \frac{a_0}{2} + \sum_{m=1}^{\infty} (a_m \cos mx + b_m \sin mx).$$

alors les a_m et b_m sont donnés comme ci-dessus. (L'hypothèse de convergence uniforme a été utilisée lorsque nous avons interverti les deux processus limites que sont la somme $\sum_{m=1}^{\infty}$ et l'intégrale $\int_{-\pi}^{\pi}$.)

Nous avons supposé ci-dessus que les a_m et b_m étaient connus et que les sommes où elles apparaissent convergeaient uniformément. Supposons maintenant que le point de départ soit une fonction $f : [-\pi, \pi) \rightarrow \mathbb{R}$. Contrairement aux contraintes imposées sur la fonction f pour qu'une série de Taylor puisse être calculée, les coefficients a_m et b_m peuvent être calculés en autant que les intégrales ci-dessus existent, par exemple si f est continue. Ainsi f pourrait être discontinue en un certain nombre (fini) de points. Elle pourrait même n'avoir de dérivée en aucun point du domaine $D = [-\pi, \pi)$ en autant que l'intégrale existe. Évidemment il n'est plus évident qu'à partir des a_m et b_m ainsi obtenus, il soit possible de reconstruire f . Voici donc les questions fondamentales qui nous occuperont dans les sections suivantes. Supposons donné une fonction f et les a_m et b_m calculés à l'aide des intégrales ci-dessus.

- Est-ce que les sommes partielles $s_m = a_0/2 + \sum_{n=1}^m (a_n \cos nx + b_n \sin nx)$ tendent vers $f(x)$ en tout point $x \in [-\pi, \pi)$? Sous quelles conditions?
- Peut-on reconstruire f en tout point $x \in D$ à partir de a_m et b_m , à l'aide des sommes partielles ou autrement?
- Peut-on mesurer la « distance » entre les sommes partielles s_m et la fonction originale f ? Cette distance tend-elle vers 0 lorsque m tend vers l'infini?

Mais avant de répondre à ces questions voici quelques exemples !

Exemple 1. Soit la fonction en créneaux définie par :

$$f(x) = \begin{cases} 1, & 0 \leq x < \pi, \\ -1, & -\pi \leq x < 0. \end{cases}$$

si $x \in [-\pi, \pi)$. Pour définir f sur tout l'axe réel, nous imposons la condition de périodicité : $f(x + 2\pi) = f(x)$. Ainsi $f : \mathbb{R} \rightarrow \mathbb{R}$. Calculons les coefficients de Fourier de cette fonction f . Le coefficient a_0 est nul. En effet :

$$a_0 = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) dx = \frac{1}{\pi} \left(\int_{-\pi}^0 -1 dx + \int_0^{\pi} +1 dx \right) = \frac{1}{\pi} (-\pi + \pi) = 0.$$

Remarquons que πa_0 est, par définition, l'aire sous la courbe de la fonction sur l'intervalle $[-\pi, \pi)$. Nous aurions donc pu nous épargner ce calcul.

Les a_n , $n > 0$, se calculent similairement :

$$\begin{aligned} a_n &= \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \cos nx \, dx \\ &= \frac{1}{\pi} \left(\int_{-\pi}^0 -\cos nx \, dx + \int_0^{\pi} +\cos nx \, dx \right) \end{aligned}$$

et puisque la fonction $\cos$ est paire, le changement de variable $y = -x$ dans la première intégrale donne

$$\begin{aligned} &= \frac{1}{\pi} \left(\int_{\pi}^0 -\cos ny(-dy) + \int_0^{\pi} \cos nx \, dx \right) \\ &= 0. \end{aligned}$$

Le fait que tous les a_n soient nuls vient du fait que la fonction f considérée ici est impaire. Nous y reviendrons à la prochaine section.

Enfin les b_n sont obtenus comme suit :

$$\begin{aligned} b_n &= \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \sin nx \, dx \\ &= \frac{1}{\pi} \left(\int_{-\pi}^0 -\sin nx \, dx + \int_0^{\pi} \sin nx \, dx \right) \end{aligned}$$

et à nouveau par le changement de variable $y = -x$ dans la première intégrale

$$\begin{aligned} &= \frac{1}{\pi} \left(\int_{\pi}^0 -\sin(-ny)(-dy) + \int_0^{\pi} \sin nx \, dx \right) \\ &= \frac{2}{\pi} \int_0^{\pi} \sin nx \, dx \\ &= -\frac{2}{\pi} \frac{\cos nx}{n} \Big|_0^{\pi} \\ &= \frac{2}{\pi n} (1 - \cos n\pi) \\ &= \frac{2}{n\pi} (1 - (-1)^n) \\ &= \begin{cases} \frac{4}{n\pi}, & \text{si } n \text{ est impair} \\ 0, & \text{si } n \text{ est pair.} \end{cases} \end{aligned}$$

Ainsi :

$$f(x) \stackrel{?}{=} \sum_{\substack{n=1 \\ n \text{ impair}}}^{\infty} \frac{4}{n\pi} \sin nx.$$

Un point d'interrogation a été placé sur le signe égal puisqu'à ce point, rien ne nous permet d'affirmer que cette égalité soit valide. En effet, il est clair qu'elle est fautive en certains points. Par exemple, en $x = 0$, tous les termes de la somme sont nuls puisque $\sin n0 = 0, \forall n \in \mathbb{Z}$. Cependant, en $x = 0$, la fonction f prend la valeur 1. Malgré cette observation l'accord peut être constaté sur le graphique suivant où les premières sommes partielles $s_m = \sum_{n=1, n \text{ impair}}^m 4 \sin nx / n\pi$ sont tracées. La figure 1.1 donne les graphes pour les sommes s_m avec $m = 1, 9, 21$ et 51 qui correspondent respectivement aux sommes partielles avec les 1, 5, 11 et 26 premiers termes non-nuls. On remarquera que l'accord devient meilleur lorsque m

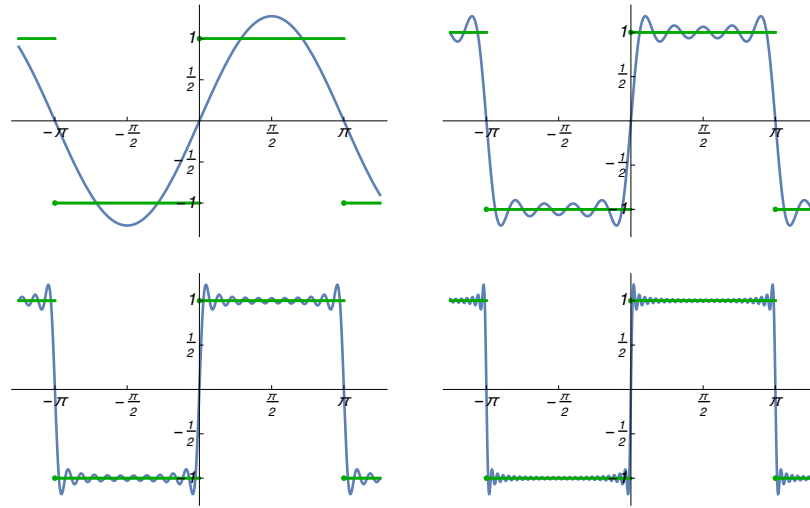


Figure 1.1 – Les premières sommes partielles s_m (en bleu) et la fonction f (en vert) de l'exemple 1 : sur la ligne supérieure s_1 et s_9 , sur l'inférieure s_{21} et s_{51} .

croît même si, proche de $x = 0$, deux pics semblent persister. Cette persistance n'est pas fortuite. Elle est connue sous le nom de phénomène de Gibbs et accompagne les sommes partielles de toutes les fonctions possédant des sauts. Nous reviendrons à ce phénomène dans un exercice de la section 4.1.

Exemple 2. Soit la fonction $f : \mathbb{R} \rightarrow \mathbb{R}$ définie par

$$f(x) = x^2, \quad -\pi \leq x < \pi$$

et

$$f(x + 2\pi) = f(x).$$

Le calcul des coefficients se fait en utilisant les expressions ci-dessus. Le coefficient a_0 est calculé séparément :

$$a_0 = \frac{1}{\pi} \int_{-\pi}^{\pi} x^2 dx = \frac{1}{\pi} \frac{x^3}{3} \Big|_{-\pi}^{\pi} = \frac{1}{\pi} \left(\frac{\pi^3}{3} - \left(-\frac{\pi^3}{3} \right) \right) = \frac{2\pi^2}{3}.$$

Et maintenant les autres $a_n, n > 0$:

$$\begin{aligned} a_n &= \frac{1}{\pi} \int_{-\pi}^{\pi} x^2 \cos nx dx \\ &= \frac{1}{\pi} \left(x^2 \frac{\sin nx}{n} \Big|_{-\pi}^{\pi} - 2 \int_{-\pi}^{\pi} x \frac{\sin nx}{n} dx \right) \end{aligned}$$

1.1. Définition et exemples _____ 7

où cette première étape a été franchie par intégration par parties. Le premier terme s'annule puisque $\sin n\pi = 0, \forall n \in \mathbb{N}$. En intégrant à nouveau par parties, on obtient :

$$= -\frac{2}{\pi n} \left(-x \frac{\cos nx}{n} \Big|_{-\pi}^{\pi} + \int_{-\pi}^{\pi} \frac{\cos nx}{n} dx \right).$$

Cette dernière intégrale s'annule et le résultat est donc :

$$\begin{aligned} &= \frac{2}{\pi n^2} (\pi \cos n\pi - (-\pi) \cos(-n\pi)) \\ &= \frac{4(-1)^n}{n^2}. \end{aligned}$$

(Êtes-vous convaincu que $\cos n\pi = (-1)^n$? Assurez-vous en !) Un calcul similaire donne :

$$b_n = 0.$$

(Pouvez-vous trouver un argument simple pour obtenir les b_n sans effort ?) La série de Fourier ainsi calculée est donc :

$$x^2 \stackrel{?}{=} \frac{\pi^2}{3} + 4 \sum_{n=1}^{\infty} \frac{(-1)^n}{n^2} \cos nx.$$

$$\begin{aligned} f(x) &= \frac{a_0}{2} + \sum_{m=1}^{\infty} (a_m \cos mx + b_m \sin mx) \\ a_m &= \frac{1}{\pi} \int_{-\pi}^{\pi} \cos mx f(x) dx, \quad m \geq 0 \\ b_m &= \frac{1}{\pi} \int_{-\pi}^{\pi} \sin mx f(x) dx, \quad m \geq 1 \end{aligned}$$

formules de base
des séries de Fourier

1*. Dans les exercices ci-dessous (et dans tout le cours), certaines intégrales **Exercices** reviendront plusieurs fois. Les calculer.

$$\begin{aligned} A_{mn} &= \frac{1}{\pi} \int_{-\pi}^{\pi} \cos mx \cos nx dx, \\ B_{mn} &= \frac{1}{\pi} \int_{-\pi}^{\pi} \sin mx \cos nx dx, \\ C_{mn} &= \frac{1}{\pi} \int_{-\pi}^{\pi} \sin mx \sin nx dx, \end{aligned}$$

Attention au cas $m = n = 0$!

2*. Trouver les coefficients de Fourier a_n et b_n pour les fonctions suivantes. Elles sont données sur l'intervalle $(-\pi, \pi]$ et prolongées périodiquement par $f(x+2\pi) = f(x)$.

- (a) $f(x) = 1$
- (b) $f(x) = x$
- (c) $f(x) = \cos x$
- (d) $f(x) = \sin 3x$
- (e) $f(x) = \sin \frac{x}{2}$
- (f) $f(x) = |\sin x|$
- (g) $f(x) = \exp(ax)$ où a est une constante réelle.
- (h) $f(x) = \cosh(ax)$.

3*. Soit $f : \mathbb{R} \rightarrow \mathbb{R}$ une fonction continue. Lesquelles des fonctions suivantes sont périodiques ? Avec quelle période ?

- (a) $a(x) = f(\cos x)$
- (b) $b(x) = \cos(f(x))$
- (c) $c(x) = f(\cos(x) + \cos(x + \sqrt{2}))$
- (d) $d(x) = f(\cos(x) + \cos(\frac{x}{2}))$
- (e) $e(x) = \sum_{i=1}^n \cos \frac{x}{i}$, pour $n \in \mathbb{N}$
- (f) $f(x) = \sum_{i=1}^n \cos ix$, pour $n \in \mathbb{N}$

4. Soit $f : \mathbb{R} \rightarrow \mathbb{R}$ une fonction périodique de période 2π . Est-il possible de remplacer les formules obtenues ci-dessous pour les coefficients de Fourier par les relations :

$$a_m = \frac{1}{\pi} \int_{\alpha-\pi}^{\alpha+\pi} \cos mx \, f(x) \, dx, \quad m \geq 0$$

$$b_m = \frac{1}{\pi} \int_{\beta-\pi}^{\beta+\pi} \sin mx \, f(x) \, dx, \quad m \geq 1$$

où α et β sont des nombres réels quelconques ?

5*. (a) Montrer que, pour chaque entier $n \geq 0$, la fonction $\tilde{f}_n : [-\pi, \pi) \setminus \{0\} \rightarrow \mathbb{R}$ donnée par

$$\tilde{f}_n(x) = \frac{\sin((n + \frac{1}{2})x)}{\sin \frac{x}{2}}$$

peut être étendue en une fonction f_n sur tout $[-\pi, \pi)$ qui soit continue.

(b) Donner le développement en série des fonctions f_n .

6. Montrer que

$$-\ln\left(2\sin\frac{x}{2}\right) = \sum_{n=1}^{\infty} \frac{\cos nx}{n}$$

sur $D = (0, 2\pi)$. Note : Le calcul de a_0 est de loin le plus difficile. Le changement de variable $t = \frac{x}{2}$ dans l'intégrale permet de réécrire

$$\ln 2 \sin \frac{x}{2} = \ln 2 + \ln \sin \frac{x}{2} = \ln 2 + \ln \sin t = \ln 2 + \ln 2 \sin \frac{t}{2} \cos \frac{t}{2}.$$

1.2 Séries d'une parité donnée

Un domaine $D \subset \mathbb{R}$ est *symétrique par rapport à l'origine* ou simplement *symétrique* si $x \in D$ implique que $-x \in D$. Une fonction définie sur un domaine D symétrique est dite *paire* si

$$f(-x) = f(x), \quad \text{pour tout } x \in D.$$

Elle est dite *impaire* si

$$f(-x) = -f(x), \quad \text{pour tout } x \in D.$$

fonction impaire

Une fonction peut être ni paire ni impaire.

Parmi les exemples les plus simples de fonctions paires sont les monômes x^{2n} , $n \in \mathbb{N}$, et les fonctions $\cos nx$, $n \in \mathbb{Z}$. Les monômes x^{2n+1} , $n \in \mathbb{N}$, et $\sin nx$, $n \in \mathbb{Z}$, sont impaires. Les graphes de fonctions paires ou impaires sont faciles à reconnaître. La partie à droite de l'axe vertical du graphe d'une fonction paire est l'image miroir par rapport à cet axe de la partie à gauche. (Voir Figure 1.2 (a).) Pour les fonctions impaires, la partie à droite de l'axe vertical est la réflexion par rapport à l'origine de la partie à gauche. (Voir Figure 1.2 (b).)

Les fonctions f paires ont la propriété

$$\int_{-I}^I f(x) dx = 2 \int_0^I f(x) dx$$

alors que les fonctions g impaires satisfont

$$\int_{-I}^I g(x) dx = 0$$

en autant que l'intervalle d'intégration soit compris dans le domaine de définition de ces fonctions. (Ces deux relations devraient être évidentes.) Notons également que le produit de deux fonctions de même parité est paire alors que le produit de deux fonctions de parité différente est impaire. (Exercice !) On en déduit donc que

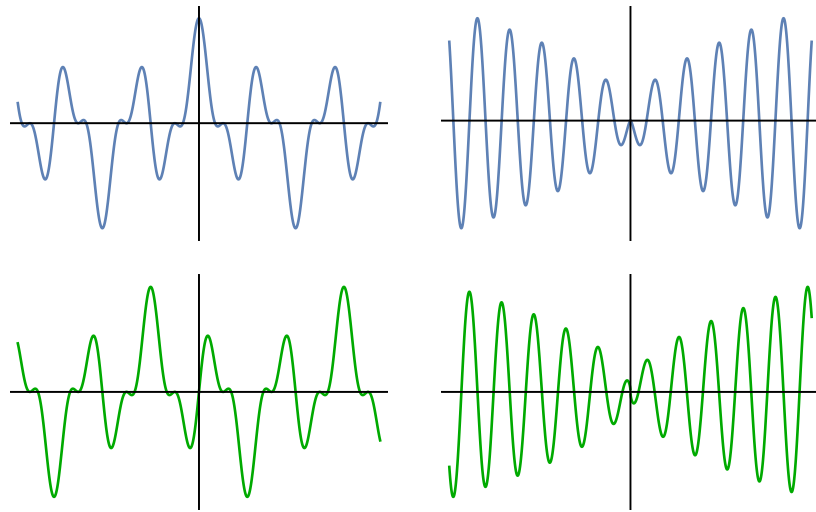


Figure 1.2 – Graphe de fonctions paires (sur la ligne supérieure) et impaires (sur l'inférieure). Celles à gauche sont périodiques, celles à droite ne le sont pas.

la série de Fourier d'une fonction f paire ne contient que des termes $\cos nx$ car tous ses coefficients de Fourier b_n sont nuls

$$\int_{-\pi}^{\pi} f(x) \sin nx \, dx = 0$$

puisque le produit $(f(x) \sin nx)$ est impair. Ainsi

$$p(x) = \frac{a_0}{2} + \sum_{m=1}^{\infty} a_m \cos mx$$

si p est paire. Similairement, si la fonction i est impaire, ses coefficients a_n seront nuls et

$$i(x) = \sum_{m=1}^{\infty} b_m \sin mx.$$

La connaissance de la parité d'une fonction simplifie donc le calcul de son développement en série de Fourier.

prolongement
d'une parité donnée

Les propriétés précédentes permettent également de faire des développements en série de sinus ou en série de cosinus uniquement. Supposons que le domaine de la fonction f soit l'intervalle $[0, \pi]$. (Ainsi, la fonction f n'est définie en *aucun* des points de $[-\pi, 0)$.) Il est possible de construire deux fonctions p et i sur l'intervalle $[-\pi, \pi]$, la première étant paire, la seconde impaire, et qui coïncident avec f sur l'intervalle $(0, \pi]$. Il suffit de définir

$$p(x) = \begin{cases} f(x), & \text{si } x \geq 0 \\ f(-x), & \text{si } x < 0 \end{cases}$$

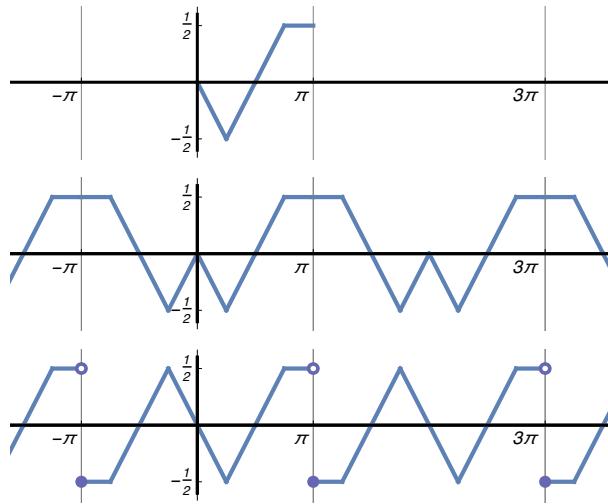


Figure 1.3 – Le graphe d'une fonction f sur $[0, \pi]$ (en haut) et de ses prolongements périodiques pair (au centre) et impair (en bas). (Pourquoi y a-t-il des points marqués sur le graphique du bas et seulement sur celui-là ?)

pour la fonction paire p et

$$i(x) = \begin{cases} f(x), & \text{si } x > 0 \\ 0, & \text{si } x = 0 \\ -f(-x), & \text{si } x < 0 \end{cases}$$

pour la fonction impaire i . La Figure 1.3 donne un exemple de prolongement pair et impair d'une fonction f définie sur l'intervalle $[0, \pi]$.

Puisque le prolongement pair p de f est une fonction paire, son développement en série de Fourier ne contiendra que les termes constant et en cosinus. Quant au développement du prolongement impair i , il ne contiendra que les termes en sinus. Quoique fort différents, ces développements, en autant qu'ils convergent, devraient converger vers la même fonction sur l'intervalle $(0, \pi)$.

Exemple. Soit la fonction $f : [0, \pi] \rightarrow \mathbb{R}$ définie par

$$f(x) = \begin{cases} \cos x, & 0 \leq x \leq \frac{\pi}{2} \\ 0, & \frac{\pi}{2} < x \leq \pi. \end{cases}$$

Nous allons développer cette fonction premièrement en une série de cosinus puis en une série de sinus.

Pour le développement en série de cosinus, nous prolongeons f en une fonction

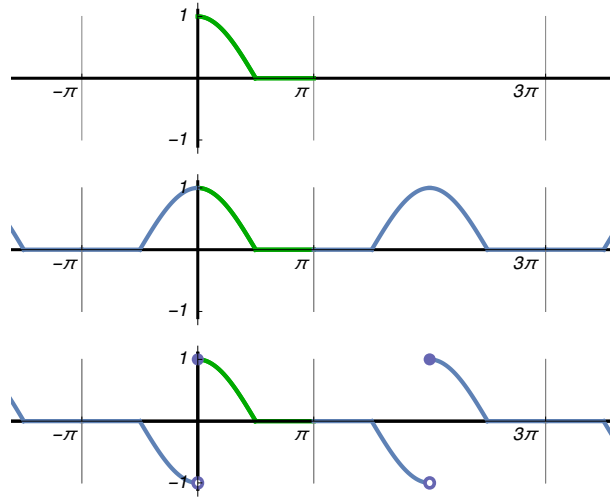


Figure 1.4 – La fonction f (en haut) et ses prolongements pair (au centre) et impair (en bas).

paire p comme nous l'avons indiqué ci-dessus :

$$p(x) = \begin{cases} \cos x, & |x| \leq \frac{\pi}{2} \\ 0, & \frac{\pi}{2} < |x| \leq \pi. \end{cases}$$

Les seuls coefficients à calculer sont les a_n .

$$\begin{aligned} a_0 &= \frac{2}{\pi} \int_0^\pi p(x) dx = \frac{2}{\pi} \int_0^{\frac{\pi}{2}} \cos x dx = \frac{2}{\pi} \sin x \Big|_0^{\frac{\pi}{2}} = \frac{2}{\pi}, \\ a_1 &= \frac{2}{\pi} \int_0^\pi p(x) \cos x dx = \frac{2}{\pi} \int_0^{\frac{\pi}{2}} \cos^2 x dx = \frac{2}{\pi} \left(\frac{1}{2} \frac{\pi}{2} \right) = \frac{1}{2}, \end{aligned}$$

et pour $n \geq 2$:

$$\begin{aligned} a_n &= \int_0^\pi p(x) \cos nx dx = \frac{2}{\pi} \int_0^{\frac{\pi}{2}} \cos x \cos nx dx \\ &= \frac{2}{\pi} \int_0^{\frac{\pi}{2}} \frac{1}{2} (\cos(n-1)x + \cos(n+1)x) dx \\ &= \frac{1}{\pi} \left(\frac{\sin(n+1)x}{n+1} + \frac{\sin(n-1)x}{n-1} \right) \Big|_0^{\frac{\pi}{2}}. \end{aligned}$$

Puisque

$$\frac{\sin(n+1)\frac{\pi}{2}}{n+1} = \begin{cases} 0, & \text{si } n \text{ est impair,} \\ \frac{(-1)^{\frac{n}{2}}}{n+1}, & \text{si } n \text{ est pair,} \end{cases}$$

alors, si n est pair :

$$a_n = \frac{1}{\pi} \left(\frac{(-1)^{\frac{n}{2}}}{n+1} - \frac{(-1)^{\frac{n}{2}}}{n-1} \right) = \frac{(-1)^{\frac{n}{2}}}{\pi} \frac{n-1-n-1}{n^2-1} = -\frac{2(-1)^{\frac{n}{2}}}{\pi(n^2-1)}.$$

(Pourquoi le calcul de a_1 a-t-il été distingué du calcul de a_n pour n quelconque ?)

Donc

$$f(x) \stackrel{?}{=} \frac{1}{\pi} + \frac{1}{2} \cos x - \frac{2}{\pi} \sum_{\substack{n=2 \\ n \text{ pair}}}^{\infty} \frac{(-1)^{\frac{n}{2}}}{n^2-1} \cos nx.$$

Remarquons que si cette égalité tient, le membre droit doit s'annuler pour tout $\frac{\pi}{2} \leq |x| \leq \pi$ et doit suivre le $\cos x$ pour $|x| \leq \frac{\pi}{2}$.

Pour le développement en série de sinus, f est prolongée en la fonction i :

$$i(x) = \begin{cases} 0, & -\pi < x < -\frac{\pi}{2}, \\ -\cos x, & -\frac{\pi}{2} \leq 0 < 0, \\ \cos x, & 0 \leq x \leq \frac{\pi}{2}, \\ 0, & \frac{\pi}{2} < x \leq \pi. \end{cases}$$

Maintenant le calcul se limite aux coefficients b_n .

$$b_1 = \frac{2}{\pi} \int_0^{\frac{\pi}{2}} \cos x \sin x \, dx = \frac{2}{\pi} \frac{1}{2} \int_0^{\frac{\pi}{2}} \sin 2x \, dx = -\frac{1}{\pi} \frac{\cos 2x}{2} \Big|_0^{\frac{\pi}{2}} = \frac{1}{\pi}$$

et si $n \geq 2$:

$$\begin{aligned} b_n &= \frac{2}{\pi} \int_0^{\frac{\pi}{2}} \cos x \sin nx \, dx \\ &= \frac{2}{\pi} \frac{1}{2} \int_0^{\frac{\pi}{2}} (\sin(n+1)x + \sin(n-1)x) \, dx \\ &= -\frac{1}{\pi} \left(\frac{\cos(n+1)x}{n+1} + \frac{\cos(n-1)x}{n-1} \right) \Big|_0^{\frac{\pi}{2}} \\ &= \begin{cases} \frac{2}{\pi} \frac{n}{n^2-1}, & n \text{ pair,} \\ \frac{2}{\pi} \frac{n + (-1)^{(n+1)/2}}{n^2-1}, & n \text{ impair.} \end{cases} \end{aligned}$$

La dernière étape vous demandera peut-être de prendre un crayon... Ainsi

$$i(x) \stackrel{?}{=} \frac{1}{\pi} \sin x + \frac{2}{\pi} \sum_{\substack{n=2 \\ n \text{ pair}}}^{\infty} \frac{n}{n^2 - 1} \sin nx + \frac{2}{\pi} \sum_{\substack{n=3 \\ n \text{ impair}}}^{\infty} \frac{n + (-1)^{(n+1)/2}}{n^2 - 1} \sin nx.$$

Exercices

7*. Identifier les fonctions paires, les impaires et celles qui n'ont pas de parité donnée.

- (a) $a(x) = \cos x^2$
- (b) $b(x) = \sin(-x^3)$
- (c) $c(x) = \exp(-x^2) \tan x$
- (d) $d(x) = \cos^2(\sin x)$
- (e) $e(x) = \cos^2(x + \alpha)$, $\alpha \in \mathbb{R}$
- (f) $f(x) = \frac{1}{x^2 + a^2}$, $a \neq 0$
- (g) $g(x) = \frac{1}{x^3 + a^3}$, $a \neq 0$

8*. Soit $p : \mathbb{R} \rightarrow \mathbb{R}$ une fonction paire et $i : \mathbb{R} \rightarrow \mathbb{R}$ une impaire. Quelle est la parité des fonctions suivantes ?

- (a) $p + p$, $p + i$, $i + i$
- (b) $p \cdot p$, $p \cdot i$, $i \cdot i$
- (c) $p \circ p$, $p \circ i$, $i \circ p$, $i \circ i$

Note : $(f \cdot g)(x) = f(x)g(x)$ et $(f \circ g)(x) = f(g(x))$.

9*. Développer les fonctions suivantes, définies sur $[0, \pi]$, en série de cosinus.

- (a) $f(x) = \sin x$
- (b) $f(x) = x$
- (c) $f(x) = x + x^2$

10. (a) Développer la fonction $f(x) = 1$, définie sur $[0, \pi]$ en série de sinus.

(b) En supposant la convergence de la série calculée en (a) aux points où la fonction prolongée est continue, démontrer l'expression suivante :

$$\frac{\pi}{4} = 1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} + \dots$$

11*. Existe-t-il des fonctions qui soient à la fois paires et impaires ?

12*. Soit $f : [0, \frac{\pi}{2}] \rightarrow \mathbb{R}$ une fonction continue. Sauriez-vous proposer un prolongement $\tilde{f} : (-\pi, \pi] \rightarrow \mathbb{R}$ de f qui coïncide avec f sur $[0, \frac{\pi}{2}]$ et tel que sa série de Fourier ne contienne que

(a) les termes $\cos nx$ avec n impairs ?

(b) les termes $\sin nx$ avec n pairs ?

13. (a) Montrer que, pour toute fonction $f : \mathbb{R} \rightarrow \mathbb{R}$, il existe une fonction paire p et une impaire i telles que $f(x) = p(x) + i(x)$, $\forall x \in \mathbb{R}$.

(b) Si f est continue, est-ce que les fonctions p et i le sont également ?

14. Si p est une fonction paire, sa dérivée p' possède-t-elle une parité ? Laquelle ?

1.3 Séries de période L

Les fonctions étudiées à la section 1.1 étaient de période 2π . Le but de la présente section est d'étendre notre discussion aux fonctions de période quelconque. Soit $f : \mathbb{R} \rightarrow \mathbb{R}$ une fonction périodique de période L :

$$f(x + L) = f(x).$$

Alors la fonction $\phi : \mathbb{R} \rightarrow \mathbb{R}$ définie par $\phi(x) = f(Lx/2\pi)$ est de période 2π :

$$\phi(x + 2\pi) = f\left(\frac{L(x + 2\pi)}{2\pi}\right) = f\left(\frac{Lx}{2\pi} + L\right) = f\left(\frac{Lx}{2\pi}\right) = \phi(x)$$

où nous avons utilisé la périodicité de f pour établir la troisième égalité.

Nous connaissons les expressions des coefficients a_n et b_n pour le développement en série de ϕ :

$$\phi(x) = f\left(\frac{Lx}{2\pi}\right) = \frac{a_0}{2} + \sum_{m=1}^{\infty} (a_m \cos mx + b_m \sin mx).$$

Nous pouvons écrire ces expressions en termes de la fonction f comme suit :

$$\begin{aligned} a_n &= \frac{1}{\pi} \int_{-\pi}^{\pi} \phi(x) \cos nx \, dx, \quad n \geq 0 \\ &= \frac{1}{\pi} \int_{-L/2}^{L/2} f\left(\frac{Lx}{2\pi}\right) \cos\left(\frac{2\pi nt}{L}\right) \frac{2\pi}{L} dt, \quad \text{où } t = \frac{Lx}{2\pi} \\ &= \frac{2}{L} \int_{-L/2}^{L/2} f(t) \cos\left(\frac{2\pi nt}{L}\right) dt \end{aligned}$$

et similairement

$$b_n = \frac{2}{L} \int_{-L/2}^{L/2} f(t) \sin\left(\frac{2\pi nt}{L}\right) dt, \quad n \geq 1.$$

Les expressions ci-dessus donnent donc les coefficients a_n et b_n du développement en série de f :

$$f(t) = \frac{a_0}{2} + \sum_{m=1}^{\infty} \left(a_m \cos\left(\frac{2\pi mt}{L}\right) + b_m \sin\left(\frac{2\pi mt}{L}\right) \right)$$

qui est obtenu de celui de ϕ en remplaçant simplement x par $2\pi t/L$.

Exemple. Développons en série de cosinus la fonction $f : [0, 1] \rightarrow \mathbb{R}$ définie par

$$f(x) = \begin{cases} 1 - \frac{x}{h}, & 0 \leq x \leq h, \\ 0, & h \leq x \leq 1, \end{cases}$$

où h est un nombre tel que $0 < h < 1$.

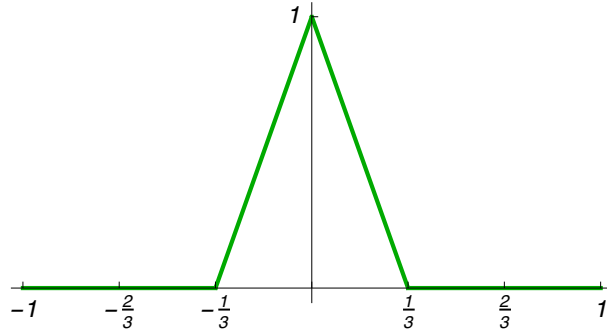


Figure 1.5 – La fonction g pour $h = \frac{1}{3}$.

Puisque nous désirons une série de cosinus, il faut prolonger f sur l'intervalle $[-1, 0)$ de façon à ce que la fonction prolongée soit paire. La fonction $g : [-1, 1] \rightarrow \mathbb{R}$ sur le graphe ci-dessus remplit cette condition. Puisque la fonction g est définie sur $[-1, 1]$, nous devons fixer la période $L = 2$. Le coefficient a_0 est l'aire sous la courbe et est donc égale à h . Les autres coefficients ($n \geq 1$) sont donnés par la formule ci-dessus avec $L = 2$:

$$a_n = \frac{2}{2} \int_{-1}^{+1} g(x) \cos\left(\frac{2\pi nx}{2}\right) dx.$$

Puisque la fonction est paire :

$$\begin{aligned} &= 2 \int_0^h \left(1 - \frac{x}{h}\right) \cos \pi n x \, dx \\ &= 2 \left(\frac{\sin \pi n x}{\pi n} \Big|_0^h - \frac{1}{h} \left(\frac{x \sin \pi n x}{\pi n} + \frac{1}{\pi n} \frac{\cos \pi n x}{\pi n} \right) \Big|_0^h \right) \end{aligned}$$

où nous avons intégré le second terme par parties. Ainsi

$$a_n = \frac{2}{h\pi^2 n^2} (1 - \cos \pi n h).$$

Comme il se doit, les coefficients dépendent du paramètre h décrivant l'ouverture de la « dent » dans le graphe de g .

$$\begin{aligned} f(t) &= \frac{a_0}{2} + \sum_{m=1}^{\infty} \left(a_m \cos \left(\frac{2\pi m t}{L} \right) + b_m \sin \left(\frac{2\pi m t}{L} \right) \right) \\ a_n &= \frac{2}{L} \int_{-L/2}^{L/2} f(t) \cos \left(\frac{2\pi n t}{L} \right) dt, \quad n \geq 0 \\ b_n &= \frac{2}{L} \int_{-L/2}^{L/2} f(t) \sin \left(\frac{2\pi n t}{L} \right) dt, \quad n \geq 1 \end{aligned}$$

série de Fourier
de période L

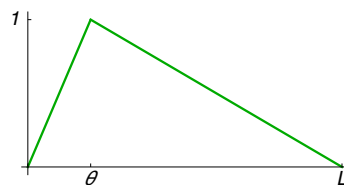
15*. Développer en série de Fourier la fonction en créneaux $f(x) = (-1)^{[x]}$ où $[x]$ désigne la partie entière de x : $[x] = n$, où $n \in \mathbb{Z}$ est l'entier tel que $n \leq x < n+1$. **Exercices**

16. Développer en série de Fourier la fonction donnée par $f(x) = x$ sur $(-1, 1]$ et de période 2.

17*. Développer en série de Fourier la fonction g de période 2 donnée sur $(-1, 1]$ par $g(x) = |x|$.

18*. (a) Obtenir les coefficients de la série de sinus de la fonction $f(x)$ donnée sur le graphique ci-dessous. (Le graphe est constitué de deux segments de droite, le maximum (= 1) étant situé en $x = \theta$, $0 < \theta < L$.)

(b) Soit $f, g : [0, L] \rightarrow \mathbb{R}$ deux fonctions dont les développements en série de sinus existent et dont les coefficients sont b_n^f et b_n^g . Si $f(x) = g(L-x)$, pour $x \in [0, L]$, trouver une relation entre b_n^f et b_n^g . Que pouvez-vous dire des coefficients calculés en (a) ?

Figure 1.6 – La corde pincée f

(c) Où une corde de guitare devrait-elle être pincée pour que l'amplitude de la première harmonique (c'est-à-dire le terme $n = 2$) soit maximale ? (On supposera que la déformation, lors du pincement, est de la forme donnée en (a)).

1.4 Séries complexes

Une autre forme des séries de Fourier est usuelle : c'est celle des séries complexes. Les fonctions de base, c'est-à-dire les fonctions $\cos nx$ et $\sin nx$, sont remplacées par les exponentielles complexes e^{inx} . Avant d'exprimer les coefficients pour ces séries, rappelons quelques définitions et propriétés élémentaires des nombres complexes.

nombres complexes

Un nombre $z = a + ib$ où a et b sont deux nombres réels et i est une des deux racines (fixée) de -1 ($i = \sqrt{-1}$) est dit un nombre complexe. Le nombre a est la partie réelle de z et b sa partie imaginaire. La somme et le produit de deux nombres complexes $z_1 = a_1 + ib_1$ et $z_2 = a_2 + ib_2$ sont donnés par

$$\begin{aligned} z_1 + z_2 &= (a_1 + a_2) + i(b_1 + b_2) \\ z_1 z_2 &= (a_1 a_2 - b_1 b_2) + i(a_1 b_2 + a_2 b_1). \end{aligned}$$

Le conjugué complexe du nombre z est noté z^* et est égal à $z^* = a - ib$. Sa valeur absolue est notée $|z|$ et est le nombre non négatif $\sqrt{a^2 + b^2}$.² Enfin, une relation remarquable relie les fonctions sinus et cosinus à la fonction exponentielle :

formule d'Euler

$$e^{i\theta} = \cos \theta + i \sin \theta.$$

Elle est parfois connue sous le nom de formule d'Euler. Si le nombre θ est un nombre réel, alors $|e^{i\theta}| = \sqrt{\cos^2 \theta + \sin^2 \theta} = 1$. L'ensemble des nombres complexes est dénoté par $\mathbb{C}$. Si ce rappel vous semble trop rapide, je vous somme d'aller faire les premiers exercices ci-dessous. Les nombres complexes joueront un rôle capital dans ce qui suit !

Utilisant la formule d'Euler, on peut exprimer les fonctions cosinus et sinus comme combinaison linéaire d'exponentielles :

$$\cos nx = \frac{1}{2}(e^{inx} + e^{-inx}) \quad \text{et} \quad \sin nx = \frac{1}{2i}(e^{inx} - e^{-inx}).$$

2. Un nombre non négatif est un nombre positif ou nul. Un nombre positif n'est pas nul.

Considérons une somme partielle de la série de Fourier d'une fonction f :

$$\frac{a_0}{2} + \sum_{n=1}^r (a_n \cos nx + b_n \sin nx)$$

et remplaçons les fonctions cosinus et sinus par les exponentielles :

$$\begin{aligned} \frac{a_0}{2} + \sum_{n=1}^r \frac{1}{2} (a_n(e^{inx} + e^{-inx}) - ib_n(e^{inx} - e^{-inx})) \\ = \frac{a_0}{2} + \sum_{n=1}^r \left(\frac{1}{2} e^{inx} (a_n - ib_n) + \frac{1}{2} e^{-inx} (a_n + ib_n) \right). \end{aligned}$$

Ainsi, avec la définition :

$$\begin{aligned} c_0 &= \frac{1}{2} a_0, \\ c_n &= \frac{1}{2} (a_n - ib_n), \quad n \geq 1, \\ c_{-n} &= \frac{1}{2} (a_n + ib_n), \quad n \geq 1, \end{aligned}$$

nous pouvons réécrire la somme partielle sous la forme compacte :

$$\sum_{n=-r}^r c_n e^{inx}.$$

(Pourquoi seules les sommes partielles $\sum_{m=1}^r$ ont-elles été considérées ici plutôt que toute la série de Fourier $\sum_{m=1}^{\infty}$? Parce qu'en général, il n'est pas possible d'intervertir l'ordre d'un nombre infini de termes dans une série sans en changer le résultat. Or, c'est précisément ce que nous avons fait ici.)

Quant aux coefficients c_n , ils peuvent être également exprimés comme intégrales simples de la fonction originale f . Par exemple, pour c_n , $n \geq 1$

$$\begin{aligned} c_n &= \frac{1}{2} (a_n - ib_n) = \frac{1}{2\pi} \left(\int_{-\pi}^{\pi} f(x) \cos nx \, dx - i \int_{-\pi}^{\pi} f(x) \sin nx \, dx \right) \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) (\cos nx - i \sin nx) \, dx \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) e^{-inx} \, dx. \end{aligned}$$

Un exercice facile vous convaincra que cette formule vaut également pour $n \leq 0$. Pour certaines fonctions, le calcul des coefficients c_n est plus aisé que celui des a_n et b_n .

Exemple. Développons la fonction $\cosh x$ en série d'exponentielles sur l'intervalle $[-\pi, \pi]$. Puisque

$$\cosh x = \frac{1}{2} (e^x + e^{-x}),$$

le calcul des c_n est particulièrement simple

$$\begin{aligned}
 c_n &= \frac{1}{2\pi} \int_{-\pi}^{\pi} \frac{1}{2} (e^x + e^{-x}) e^{-inx} dx \\
 &= \frac{1}{4\pi} \int_{-\pi}^{\pi} (e^{x(1-in)} + e^{x(-1-in)}) dx \\
 &= \frac{1}{4\pi} \left(\frac{e^{x(1-in)}}{1-in} + \frac{e^{x(-1-in)}}{-1-in} \right) \Big|_{-\pi}^{\pi} \\
 &= \frac{(-1)^n}{4\pi} \left(\frac{e^{\pi} - e^{-\pi}}{1-in} + \frac{e^{-\pi} - e^{\pi}}{-1-in} \right) \\
 &= \frac{(-1)^n \sinh \pi}{\pi(1+n^2)}.
 \end{aligned}$$

Donc

$$\cosh x = \sum_{n \in \mathbb{Z}} \frac{(-1)^n e^{inx} \sinh \pi}{\pi(1+n^2)}, \quad -\pi \leq x \leq \pi.$$

continuité sur le cercle

Nous allons nous préoccuper des problèmes de rigueur à partir du chapitre 2. Nous ferons alors un rappel des principales définitions nécessaires à cette étude fine. Cependant il est utile d'introduire immédiatement le concept de continuité sur le cercle. Soit $f : [-\pi, \pi) \rightarrow \mathbb{R}$ une fonction. Son prolongement périodique $\tilde{f} : \mathbb{R} \rightarrow \mathbb{R}$ est défini comme d'habitude par $\tilde{f}(x) = f(x + 2\pi n)$ où n est l'entier $\in \mathbb{Z}$ tel que $x + 2\pi n \in [-\pi, \pi)$. La fonction f est *continue sur le cercle* si $\tilde{f}$ est continue sur $\mathbb{R}$. La continuité de f sur le cercle est légèrement plus exigeante que la continuité de f sur $[-\pi, \pi)$: elle ajoute l'exigence que $\tilde{f}$ soit continue en π , ce qui se traduit sur la fonction f par les égalités suivantes entre les deux limites de f aux points $\pm\pi$:

$$f(-\pi) = \lim_{x \rightarrow -\pi^+} f(x) = \lim_{x \rightarrow \pi^-} f(x).$$

Cette définition est en fait un exemple de définition de la continuité d'une fonction sur une variété. A ce point, vous n'avez probablement rencontré que la définition de continuité d'une fonction en un point ou sur un intervalle de $\mathbb{R}$. Il est possible de définir la continuité de fonctions sur des espaces qui ressemblent à des espaces euclidiens (tels les $\mathbb{R}^n$) mais qui n'en sont pas. Un exemple simple est fourni par la fonction T qui donne la température en chaque point du globe. Son domaine n'est pas un sous-ensemble du plan mais bien la surface de la terre, c'est-à-dire une sphère. Dans notre exemple, la « bonne » surface est un cercle. Pour le définir correctement, considérons d'abord le cercle de rayon unité paramétré par l'angle $\theta \in [-\pi, \pi)$. Sur cet ensemble un intervalle ouvert peut-être de trois types : (i) ce peut être le cercle en entier, (ii) un intervalle ouvert (a, b) avec $-\pi \leq a < b \leq \pi$ qui est défini comme d'habitude comme étant l'ensemble des points du cercle correspondants aux valeurs de θ qui sont dans $\{\theta \mid a < \theta < b\}$, et (iii) un sous-ensemble noté (a, b) avec $-\pi \leq b < a < \pi$ qui est l'ensemble de points du cercle correspondants aux valeurs de θ dans $\{\theta \mid a < \theta < \pi \text{ ou } -\pi \leq \theta < b\}$. Notons que

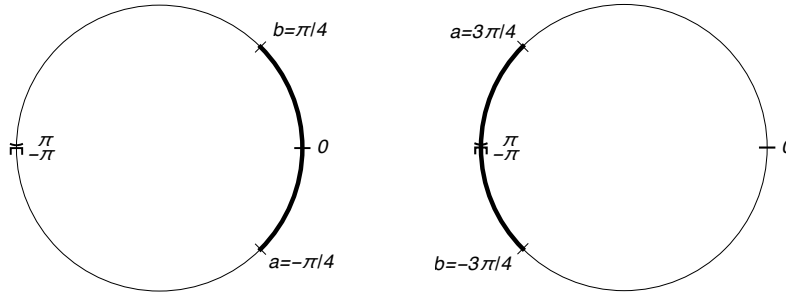


Figure 1.7 – À gauche l'intervalle $\{\theta \mid a = -\frac{\pi}{4} < \theta < b = \frac{\pi}{4}\}$ de type (ii) et à droite l'intervalle $\{\theta \mid a = \frac{3\pi}{4} < \theta < \pi \text{ ou } -\pi < \theta < b = -\frac{3\pi}{4}\}$ de type (iii).

si $b < a$, l'intervalle ouvert (a, b) contient le point correspondant à $-\pi$. De plus, tout intervalle ouvert contenant le point $-\pi$ contient un sous-intervalle $(\pi - \epsilon, \pi)$ avec $\epsilon > 0$. Le cercle $\mathbb{T}^1$ est le cercle paramétré par θ muni des d'intervalles ouverts le cercle $\mathbb{T}^1$ que nous venons de décrire.³ La continuité sur l'ensemble $\mathbb{T}^1$ que nous venons de définir est alors la continuité sur le cercle que nous avons définie ci-dessus. Nous noterons maintenant souvent les intégrales sur l'intervalle $[-\pi, \pi)$ comme suit

$$\int_{-\pi}^{\pi} f(x) dx = \int_{\mathbb{T}^1} f(x) dx.$$

$$f(x) = \sum_{n \in \mathbb{Z}} c_n e^{inx}$$

$$c_n = \frac{1}{2\pi} \int_{\mathbb{T}^1} f(x) e^{-inx} dx, \quad n \in \mathbb{Z}.$$

formules pour les séries d'exponentielles

19. Poser $z_1 = a_1 + ib_1$ et $z_2 = a_2 + ib_2$ où a_1, b_1, a_2 et b_2 sont des nombres réels. **Exercices**
Montrer que

- (a) $z_1 z_2^* - z_1^* z_2$ est un nombre imaginaire pur
- (b) et que $z_1/z_2 + z_1^*/z_2^*$ est un nombre réel si $z_2 \neq 0$.

Note : Rappelons qu'un nombre complexe z est un nombre réel si $z = z^*$ et un imaginaire pur si $z = -z^*$.

20*. Si $z_1 = r_1(\cos \theta_1 + i \sin \theta_1)$ et $z_2 = r_2(\cos \theta_2 + i \sin \theta_2)$, $r_1, r_2, \theta_1, \theta_2 \in \mathbb{R}$, montrer en utilisant des identités trigonométriques que :

³. Le choix de la lettre $\mathbb{T}^1$ pour le cercle vient de *tore*. Le cercle est un tore à une dimension et la surface d'un beigne est un tore à deux dimensions.

formule de de Moivre
Abraham de Moivre
(1667-1754)

(a) $1/z_1 = (\cos \theta_1 - i \sin \theta_1)/r_1$

(b) $z_1^* = r_1(\cos \theta_1 - i \sin \theta_1)$

(c) $z_1 z_2 = r_1 r_2 (\cos(\theta_1 + \theta_2) + i \sin(\theta_1 + \theta_2))$

(d) $z_1^n = r_1^n (\cos n\theta_1 + i \sin n\theta_1)$ qui est connue sous le nom de formule de de Moivre.

21*. Obtenir le développement en série d'exponentielles des fonctions suivantes définies sur $\mathbb{T}^1$.

(a) $f(x) = x$. (Vérifier dans ce cas le résultat avec celui du second exercice de la section 1.)

(b) $f(x) = \sin ax$, pour $a \in \mathbb{R}$

22*. Nous avons restreint notre discussion des séries de Fourier et des séries d'exponentielles à des fonctions à valeurs réelles. Cependant ce traitement s'étend sans changement aux fonctions qui prennent des valeurs complexes : $f : \mathbb{R} \rightarrow \mathbb{C}$. Dans ce cas les c_n sont tous indépendants. Certaines propriétés sur f se matérialisent par des relations sur les c_n .

(a) Montrer que, si f est à valeur réelle, alors $c_n = c_{-n}^*$.

(b) Quelles propriétés vérifient les c_n si f est une fonction paire ?

(c) Et si f est une fonction impaire ?

(d) Et si f est une fonction de période $\frac{2\pi}{n}$, $n \in \mathbb{N}$, qui a été développée comme si sa période était 2π ?

23*. (a) Soit f, g et h les fonctions définies sur l'intervalle $[-\pi, \pi]$ par $f(x) = 1$, $g(x) = x$ et $h(x) = x^2$ respectivement et étendues sur l'axe réel par périodicité. Lesquelles de ces trois fonctions sont continues sur le cercle ?

(b) Lesquelles sont continues sur le cercle et ont une dérivée continue sur le cercle ?

24*. (a) Montrer : Soit $f : \mathbb{T}^1 \rightarrow \mathbb{R}$ une fonction qui soit $(n-1)$ fois continûment différentiable sur le cercle et telle que $f^{(n-1)}$ soit différentiable avec dérivée continue (sur le cercle) sauf en un nombre fini de points $x_1, x_2, \dots, x_l$. Si $|f^{(n)}(x)| \leq M$ partout en $x \notin \{x_1, x_2, \dots, x_l\}$, alors les coefficients du développement en série complexe vérifient : $|c_r| \leq M/|r|^n$, pour $r > 0$.

(b) Obtenir le développement en exponentielles de $f : \mathbb{T}^1 \rightarrow \mathbb{R}$ donné par $f(x) = x(x^2 - \pi^2)$. Vérifier que les c_r satisfont (a).

polynômes de Bernoulli

25*. Le but de cet exercice est de construire des fonctions dont le m -ième coefficient de Fourier est simplement une puissance (négative) de m . Le résultat

intervient dans plusieurs domaines des mathématiques et particulièrement en combinatoire. On définit une famille de polynômes B_n , $n \geq 0$, par les deux relations :

$$B_n(x+1) - B_n(x) = nx^{n-1} \quad (\text{i})$$

$$\frac{dB_n}{dx}(x) = nB_{n-1}(x), \quad \text{si } n \geq 1 \quad (\text{ii})$$

et les deux conditions

$$B_0(x) = 1, \quad (\text{iii})$$

$$\text{le polynôme } B_n \text{ est de degré } n \text{ et} \quad (\text{iv})$$

$$\int_0^1 B_n(x) dx = 0, \quad \text{si } n \geq 1. \quad (\text{v})$$

Les polynômes B_n ainsi construits sont appelés les polynômes de Bernoulli.

(a) Construire $B_1(x)$ et $B_2(x)$. Note : Poser $B_1(x) = ax+b$, $B_2(x) = cx^2+dx+e$ et utiliser les relations ci-dessous pour fixer les a, b, c, d, e . Vous devrez construire partiellement $B_3(x)$.

(b) Identifier le cercle à l'intervalle $[0, 1)$ où les extrémités sont recollées et noter par le même caractère B_n la restriction des polynômes à cet intervalle : $B_n : \mathbb{T}^1 \rightarrow \mathbb{R}$ pour tout $n \geq 0$. Montrer que les $(n-2)$ premières dérivées de $B_n(x)$ sont continues sur le cercle. Note : Cette question consiste donc à montrer que $B_n^{(i)}(0) = B_n^{(i)}(1)$ pour $i = 0, 1, \dots, n-2$.

(c) Se convaincre, par intégration par parties, que les intégrales

$$\int_0^1 x^n \cos 2\pi mx \, dx \quad \text{et} \quad \int_0^1 x^n \sin 2\pi mx \, dx$$

sont de la forme

$$\sum_{i=0}^n \frac{d_i}{m^i}, \quad c_i \in \mathbb{R}.$$

(d) Montrer que les $B_n(x)$ ont la parité $(-1)^n$ par rapport au point $x = \frac{1}{2}$, c'est-à-dire, montrer que $B_n(\frac{1}{2} + \epsilon) = (-1)^n B_n(\frac{1}{2} - \epsilon)$. Note : par induction.

(e) En utilisant l'exercice précédent, montrer que les séries de Fourier de B_n sont

$$B_n(x) = c_n \sum_{m \in \mathbb{N}} \frac{\sin 2\pi mx}{m^n}, \quad n \text{ impair},$$

et

$$B_n(x) = c_n \sum_{m \in \mathbb{N}} \frac{\cos 2\pi mx}{m^n}, \quad n \text{ pair}.$$

(f) Montrer que la constante c_n est donnée par

$$c_n = \begin{cases} \frac{2(-1)^{\frac{n+1}{2}} n!}{(2\pi)^n}, & n \text{ impair}, \\ \frac{2(-1)^{\frac{n}{2}+1} n!}{(2\pi)^n}, & n \text{ pair}. \end{cases}$$

Note : Pour cela, commencer par montrer que le coefficient de x^n dans B_n est égal à 1.

(g) Montrer que la fonction génératrice des polynômes de Bernoulli est :

$$\frac{te^{xt}}{e^t - 1} = \sum_{n=0}^{\infty} B_n(x) \frac{t^n}{n!}.$$

Note : Nous introduirons plus loin, à la section 3.4, le concept de fonction génératrice. Si vous n'êtes pas familiers avec ce concept, ignorez (pour l'instant) cet exercice.

1.5 La corde vibrante

Les séries de Fourier ont été introduites dans le but de résoudre l'équation de la chaleur. Leur rôle peut être aisément étendu aux équations aux dérivées partielles linéaires dont font partie, outre l'équation de la chaleur, l'équation des ondes (qui décrit la corde vibrante, le tambour ou les ondes électromagnétiques), l'équation du fléchissement des poutres, etc. A la section 3.1, nous reverrons les définitions relatives aux équations différentielles ordinaires et aux dérivées partielles. Pour la section présente, aucune connaissance particulière n'est supposée sur ce sujet.

Nous n'étudierons ici qu'une seule application : la solution de l'équation de la corde vibrante qui décrit entre autres la dynamique des cordes du violon, de la guitare et du piano. D'origine physique, ce problème révèle la puissance des séries de Fourier. Avant de passer à ce problème, nous allons cependant résoudre un problème plus simple, celui de la corde plombée. Les techniques utilisées nous indiqueront la voie pour la corde vibrante et certaines propriétés remarquables seront redécouvertes dans les sections suivantes (3.2 et 3.3).

Dans les deux cas, notre discussion suivra les étapes suivantes :

1. description de l'équation d'évolution,
2. séparabilité et solution générale,
3. conditions initiales.

1.5.1 La corde plombée

1. *L'équation de la corde plombée* — Une corde plombée est une corde tendue de masse nulle sur laquelle sont fixés à distance égale n poids de masse égale m .

La position d'équilibre correspond à la corde rectiligne. (Voir la figure 1.8.) Nous restreindrons le mouvement de ces masses à un plan. Pour obtenir l'équation du mouvement de ce système, nous ferons deux hypothèses : premièrement, les composantes horizontales de la tension des cordes à gauche et à droite de chacune des masses s'annulent (ce qui implique que le mouvement longitudinal de ces masses est nul) et, deuxièmement, le déplacement transversal (c'est-à-dire vertical) est petit par rapport à la distance l entre les masses. Cette dernière hypothèse est souvent appelée l'hypothèse des petites oscillations. Le déplacement vertical de la i -ième masse par rapport à sa position d'équilibre sera noté x_i .

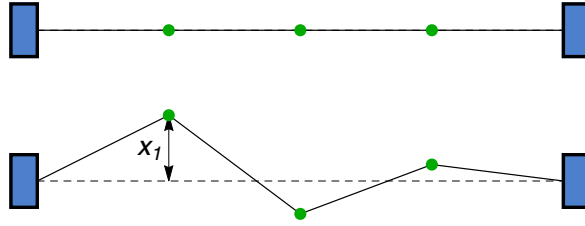


Figure 1.8 – La corde plombée $n = 3$ à sa position d'équilibre (en haut) et dans en vibration (en bas)

Le point de départ est l'équation du mouvement de Newton $F = ma$ que nous devons écrire pour chacune des masses. L'accélération est évidemment la dérivée seconde d^2x_i/dt^2 . La seule force sur la masse m_i , $1 \leq i \leq n$, est due à la tension dans la corde. Puisque le déplacement n'est que selon l'axe vertical, seule la composante verticale de la tension T entre en jeu. Ainsi la force F_i sur m_i ($1 < i < n$) sera la somme des composantes verticales de la tension dans les segments de la corde à gauche et à droite de la masse. La composante verticale de la tension dans le segment gauche est $T \sin \theta_{i-1,i}$ où $\theta_{i-1,i}$ est l'angle que la corde entre m_{i-1} et m_i fait par rapport à l'horizontale. (Nous avons utilisé ici le fait que la tension dans la corde est constante et égale à T .) Si les déplacements x_{i-1} et x_i sont petits, alors le sinus est approximativement égal à $-(x_i - x_{i-1})/l$. Une expression similaire est obtenue pour le segment à droite de m_i . Ainsi l'équation de Newton $ma_i = F_i$ pour la i -ième masse ($1 < i < n$) est :

$$\begin{aligned} m \frac{d^2x_i}{dt^2} &= -T \left(\frac{x_i - x_{i-1}}{l} + \frac{x_i - x_{i+1}}{l} \right) \\ &= -\frac{T}{l} (-x_{i-1} + 2x_i - x_{i+1}). \end{aligned}$$

Puisque les oscillations sont petites (seconde hypothèse), le rapport $\frac{T}{ml}$ est presque constant ; nous le poserons égal à 1. (Y perd-t-on en généralité ?) Alors l'équation ci-dessus devient simplement

$$\frac{d^2x_i}{dt^2} = -(-x_{i-1} + 2x_i - x_{i+1}).$$

Les deux poids extrêmes vérifient les équations :

$$\begin{aligned}\frac{d^2 x_1}{dt^2} &= -(2x_1 - x_2) \\ \frac{d^2 x_n}{dt^2} &= -(-x_{n-1} + 2x_n).\end{aligned}$$

(Pourquoi ?) Ces équations sont donc les équations du mouvement de la corde plombée.

Ce système d'équations différentielles ordinaires peut être mis sous la forme matricielle

$$\frac{d^2 \vec{x}}{dt^2} = -A\vec{x} \quad (*)$$

où

$$\vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ \vdots \\ x_{n-1} \\ x_n \end{pmatrix} \in \mathbb{R}^n \quad \text{et} \quad A = \begin{pmatrix} 2 & -1 & 0 & \cdots & 0 & 0 \\ -1 & 2 & -1 & \cdots & 0 & 0 \\ 0 & -1 & 2 & \cdots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \cdots & 2 & -1 \\ 0 & 0 & 0 & \cdots & -1 & 2 \end{pmatrix} \in \mathbb{R}^{n \times n}.$$

Pour le physicien, ce système décrit l'évolution au cours du temps des n poids. Pour le mathématicien, il s'agit d'un système d'équations différentielles (à cause des $\frac{d^2 x_i}{dt^2}$ qui y interviennent) qui décrit la dépendance des fonctions x_i en fonction de leur argument t . Résoudre ce système, c'est donner les n fonctions x_i explicitement. C'est ce que nous faisons à la prochaine étape.

2. Une hypothèse simplificatrice et la solution générale — Comment résoudre ce système d'équations différentielles ordinaires ? La solution générale de l'équation différentielle

$$\frac{d^2 y}{dt^2} = -ay$$

où $y = y(t)$ est une fonction réelle du temps t et a une constante positive est bien connue :

$$y = \alpha \cos \sqrt{a}t + \beta \sin \sqrt{a}t$$

où α et β sont des constantes réelles. (Que cette fonction soit une solution peut être vérifiée directement. Que cette fonction soit la solution *générale* vient du fait que l'équation différentielle étant du second ordre, sa solution générale dépend de deux constantes réelles, ici α et β . Nous en reparlerons à la section 3.1.) Nous

allons chercher une solution du système (*) de forme similaire. Soit

$$\vec{x} = \begin{pmatrix} x_{10} \\ x_{20} \\ x_{30} \\ \vdots \\ x_{n-1,0} \\ x_{n0} \end{pmatrix} \cos \omega t = \vec{x}_0 \cos \omega t$$

où ω est une constante à déterminer et $\vec{x}_0$ un vecteur constant $\in \mathbb{R}^n$ également à déterminer. Nous aurions pu prendre tout aussi bien $\sin \omega t$ plutôt que $\cos \omega t$; nous en aurons d'ailleurs besoin pour écrire la solution générale. Cette forme particulière de la solution recherchée $\vec{x}$ implique que :

toutes les variables x_i auront la même dépendance en t ,
c'est-à-dire elles oscilleront avec la même fréquence.

Puisque ω et $\vec{x}_0$ sont des constantes, la dérivée seconde de $\vec{x}$ se calcule aisément :

$$\frac{d^2 \vec{x}}{dt^2} = \frac{d^2}{dt^2} (\vec{x}_0 \cos \omega t) = -\vec{x}_0 \omega^2 \cos \omega t = -\omega^2 \vec{x}.$$

Le système différentiel (*) devient alors

$$\begin{aligned} -\omega^2 \vec{x} &= -A\vec{x} \\ 0 &= (A - \omega^2 I)\vec{x} \end{aligned}$$

où I est la matrice identité. Puisque cette égalité doit être satisfaite pour tout t , nous devons avoir :

$$0 = (A - \omega^2 I)\vec{x}_0.$$

Mais cette équation est un problème matriciel. Le vecteur $\vec{x}_0$ sera non-nul que si $(A - \omega^2 I)$ est singulière. Ainsi :

la solution du système différentiel est réduit au problème
de trouver les valeurs propres ω^2 de la matrice A .

Avant même de commencer le travail, nous pouvons observer que la matrice A est symétrique. De cette observation suivent les propriétés suivantes :

la matrice A est diagonalisable, ses valeurs propres sont réelles et
il existe une base orthonormale constituée de ses vecteurs propres.

Ceci est un résultat d'algèbre linéaire. Mais ceci termine la recherche de la solution générale du système (*) ! Le système (*) contient n équations du second ordre. La solution générale fera donc intervenir $2n$ constantes. Or la matrice A possède n vecteurs propres linéairement indépendants et chacun fournit une solution

$$\vec{x}_0^i \cos \omega_i t$$

où ω_i est la racine carrée de la i -ième valeur propre de A ($1 \leq i \leq n$) et $\vec{x}_0^i$ le vecteur propre correspondant. Mais, comme nous l'avons dit plutôt, nous aurions pu également considérer les solutions

$$\vec{x}_0^i \sin \omega_i t$$

qui forment, avec les précédentes, $2n$ solutions linéairement indépendantes. La solution générale du système (*) est donc une combinaison linéaire de ces $2n$ solutions. Chacune de ces $2n$ solutions est appelée un *mode propre* de la corde plombée. Les constantes ω_i , $1 \leq i \leq n$, sont appelées *fréquences propres* ou *naturelles* et leur ensemble est appelé le *spectre* de la corde plombée.

mode propre
fréquence propre
spectre

3. Conditions initiales — Ce résultat d'existence peut « vous laisser sur votre faim »... Voici donc un exemple particulier où nous décrirons une solution particulière pour la corde plombée avec trois masses ($n = 3$). Pour obtenir une solution particulière, il faut spécifier complètement les conditions initiales. Puisque l'équation de Newton est du second ordre, il faut donner les positions et vitesses initiales des n masses. Dans le cas $n = 3$, il faudra donc donner les six constantes :

$$\begin{aligned} x_1(0), \quad x_2(0), \quad x_3(0), \\ \frac{dx_1}{dt}(0), \quad \frac{dx_2}{dt}(0), \quad \frac{dx_3}{dt}(0). \end{aligned}$$

Commençons par trouver les modes propres. La matrice à diagonaliser est

$$A = \begin{pmatrix} 2 & -1 & 0 \\ -1 & 2 & -1 \\ 0 & -1 & 2 \end{pmatrix}.$$

Le polynôme caractéristique est $\det(A - xI) = -4 + 10x - 6x^2 + x^3$ où, il faut se rappeler, x a été mis pour ω^2 . Les trois valeurs propres sont donc : $2 - \sqrt{2}$, 2 , $2 + \sqrt{2}$. Les vecteurs propres associées sont, dans l'ordre :

$$\vec{x}_0^1 = \begin{pmatrix} 1 \\ \sqrt{2} \\ 1 \end{pmatrix}, \quad \vec{x}_0^2 = \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix} \quad \text{et} \quad \vec{x}_0^3 = \begin{pmatrix} 1 \\ -\sqrt{2} \\ 1 \end{pmatrix}.$$

Nous venons donc de trouver six solutions linéairement indépendantes. Deux modes propres correspondent à la fréquence la plus basse :

$$\begin{pmatrix} 1 \\ \sqrt{2} \\ 1 \end{pmatrix} \cos \sqrt{2 - \sqrt{2}} t \quad \text{et} \quad \begin{pmatrix} 1 \\ \sqrt{2} \\ 1 \end{pmatrix} \sin \sqrt{2 - \sqrt{2}} t.$$

Les quatre autres modes peuvent être écrits similairement. La solution générale pour $n = 3$ est donc

$$\begin{aligned}\vec{x} = & \vec{x}_0^1(\alpha_1 \cos \omega_1 t + \beta_1 \sin \omega_1 t) \\ & + \vec{x}_0^2(\alpha_2 \cos \omega_2 t + \beta_2 \sin \omega_2 t) \\ & + \vec{x}_0^3(\alpha_3 \cos \omega_3 t + \beta_3 \sin \omega_3 t)\end{aligned}$$

qui dépend, comme il se doit, de six constantes arbitraires : $\alpha_1, \alpha_2, \alpha_3$ et $\beta_1, \beta_2, \beta_3$.

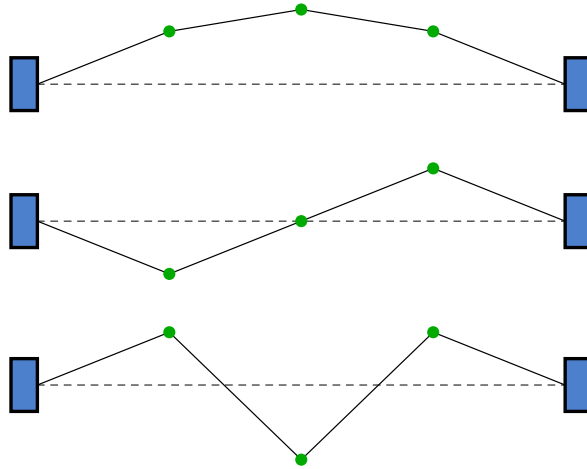


Figure 1.9 – Les trois configurations propres de la corde plombée $n = 3$

Ces modes sont représentés à la figure 1.9. Les modes propres avec la fréquence propre la plus faible $\omega_1 = \sqrt{2} - \sqrt{2}$ correspond à une configuration où les trois masses forment un seul ventre. Les modes correspondant à la seconde fréquence $\omega_2 = \sqrt{2}$ montrent deux ventres et les modes pour la plus grande fréquence $\omega_3 = \sqrt{2} + \sqrt{2}$ en ont trois.

Soit les conditions initiales suivantes :

$$\begin{aligned}x_1(0) &= 1, & x_2(0) &= 0, & x_3(0) &= 1 \\ \frac{dx_1}{dt}(0) &= 0, & \frac{dx_2}{dt}(0) &= 0, & \frac{dx_3}{dt}(0) &= 0.\end{aligned}$$

Pour trouver la solution particulière correspondant à ces conditions initiales, il faut donc fixer les constantes indéterminées $\alpha_1, \alpha_2, \alpha_3$ et $\beta_1, \beta_2, \beta_3$ de la solution générale ci-dessus de façon à ce que

$$\vec{x}(0) = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} \quad \text{et} \quad \frac{d\vec{x}}{dt} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}.$$

En prenant la dérivée de la solution générale et en l'évaluant en $t = 0$, nous obtenons :

$$\begin{aligned} 0 &= \frac{d\vec{x}}{dt}(0) \\ &= \vec{x}_0^1 \omega_1 \beta_1 + \vec{x}_0^2 \omega_2 \beta_2 + \vec{x}_0^3 \omega_3 \beta_3. \end{aligned}$$

Puisque les trois vecteurs propres sont linéairement indépendants, il faut donc que $\beta_1 = \beta_2 = \beta_3 = 0$. Posons maintenant $t = 0$ dans la solution

$$\begin{aligned} \vec{x}(0) &= \vec{x}_0^1 \alpha_1 + \vec{x}_0^2 \alpha_2 + \vec{x}_0^3 \alpha_3 \\ &= \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}. \end{aligned}$$

Un rapide coup d'oeil montre que la solution est $\alpha_1 = \alpha_3 = \frac{1}{2}$ et $\alpha_2 = 0$. La solution désirée est donc

$$\vec{x} = \frac{1}{2} \vec{x}_0^1 \cos \omega_1 t + \frac{1}{2} \vec{x}_0^3 \cos \omega_3 t.$$

Passons maintenant au problème continu.

1.5.2 La corde vibrante

1. *L'équation de la corde vibrante* — Considérons une corde fixée entre 0 et L ne vibrant que dans un plan de l'espace et notons $u(x, t)$ l'écart (vertical) d'un élément de corde Δx centré au point x et au temps t . La densité linéaire de la corde sera notée ρ . Nous ferons les deux mêmes hypothèses que précédemment. La première stipule que la corde demeure de densité constante (c'est-à-dire qu'il n'y a pas de vibrations longitudinales) et donc que la tension T dans la corde est partout la même. La seconde est celle des petites oscillations : $u(x, t) \ll L$.

Comme précédemment le point de départ est l'équation de Newton : $F = ma$. Ici cependant, la corde est un milieu continu et nous nous concentrerons sur un élément de longueur Δx qui jouera le rôle d'un des poids. La masse de cet élément est $\rho \Delta x$. L'accélération a de l'élément est la dérivée seconde de u par rapport à t évaluée au point x où est situé l'élément Δx . Les forces (verticales) sur cet élément sont les parties verticales de la tension dans la corde aux deux extrémités de l'élément. Ainsi l'équation du mouvement $ma = F$ se lit :

$$\rho \Delta x \frac{\partial^2}{\partial t^2} u(x, t) = \text{partie verticale de } (T_2) - \text{partie verticale de } (T_1).$$

Les parties verticales sont obtenues à l'aide des sinus des angles θ_1 et θ_2 définis sur la figure.

$$\begin{aligned} &= T(\sin \theta_2 - \sin \theta_1) \\ &\cong T(\tan \theta_2 - \tan \theta_1). \end{aligned}$$

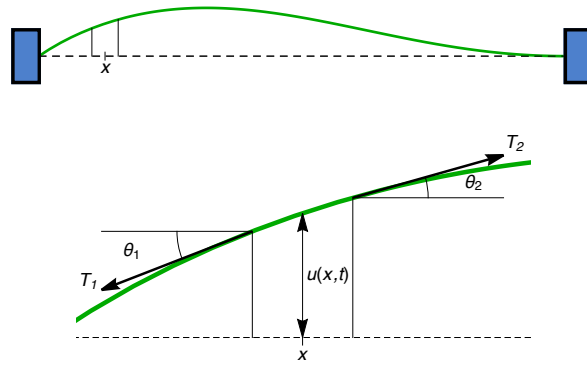


Figure 1.10 – La corde vibrante (en haut) et un agrandissement d'un élément Δx avec les forces qui agissent sur lui (en bas).

Pour cette dernière étape, nous avons utilisé le fait que les fonctions sin et tan sont les mêmes pour des arguments suffisamment petits. Plus précisément, le premier terme du développement en série de Taylor autour de 0 de sin et de tan coïncident. C'est ici que nous utilisons l'hypothèse des petites oscillations. Puisque la tangente de θ_1 est précisément la pente de la droite tangente au graphe de u en $x - \frac{\Delta x}{2}$, on peut remplacer $\tan \theta_1$ par $\frac{\partial u}{\partial x}(x - \frac{\Delta x}{2}, t)$ et similairement pour $\tan \theta_2$:

$$\cong T \left(\frac{\partial u}{\partial x}(x + \frac{\Delta x}{2}, t) - \frac{\partial u}{\partial x}(x - \frac{\Delta x}{2}, t) \right)$$

et alors

$$\frac{\partial^2}{\partial t^2} u(x, t) \cong \frac{T}{\rho} \frac{\left(\frac{\partial u}{\partial x}(x + \frac{\Delta x}{2}, t) - \frac{\partial u}{\partial x}(x - \frac{\Delta x}{2}, t) \right)}{\Delta x}.$$

Pour Δx suffisamment petit, le membre de droite devient T/ρ fois la dérivée seconde de u par rapport à x . L'équation de la corde vibrante que nous étudierons est donc

$$\frac{\partial^2}{\partial t^2} u(x, t) = a^2 \frac{\partial^2}{\partial x^2} u(x, t). \quad (**)$$

La constante *positive* T/ρ a été remplacée par la constante a^2 . Cette équation est un exemple classique d'équations différentielles aux dérivées partielles. Elle est *linéaire*, c'est-à-dire si u_1 et u_2 sont des solutions, alors $c u_1$ et $u_1 + u_2$ (où c est une constante) le sont aussi.

2. Séparabilité et solution générale — Nous avons vu que, pour la corde plombée, la forme de la solution ($\vec{x} = \vec{x}_0 \cos \omega t$) donnait à toutes les masses la même

dépendance en t , à savoir $\cos \omega t$. L'analogie de cette hypothèse pour le cas de la corde continue est de chercher une solution de la forme :

$$u(x, t) = X(x)T(t)$$

où les fonctions X et T ne dépendent que d'une variable. Il est clair qu'en général, la solution de l'équation de la corde vibrante n'aura pas cette forme. Cependant, comme pour la corde plombée, nous pouvons espérer obtenir la solution générale en faisant des combinaisons linéaires de solutions de la forme $u(x, t) = X(x)T(t)$. L'équation devient alors

$$\frac{\partial^2 XT}{\partial t^2} = a^2 \frac{\partial^2 XT}{\partial x^2}$$

et puisque X ne dépend pas de t ni T de x :

$$X \frac{d^2 T}{dt^2} = a^2 T \frac{d^2 X}{dx^2}$$

et donc

$$\frac{1}{T} \frac{d^2 T}{dt^2} = a^2 \frac{1}{X} \frac{d^2 X}{dx^2}.$$

Nous avons remplacé les signes de dérivée partielle (∂) par des signes de dérivée totale (d) puisque les fonctions sont maintenant d'une seule variable. Utilisant la notation X'' et T'' pour les dérivées secondes, l'équation à résoudre est donc simplement :

$$\frac{1}{a^2} \frac{T''}{T} = \frac{X''}{X}.$$

Remarquons que le membre de gauche ne dépend que de t alors que le membre de droite ne dépend que de x . Ceci n'est possible que si les deux membres sont en fait égaux à la même constante. (Beaucoup sont déroutés par cet argument. Voici une autre façon de vous convaincre de ce résultat. Essayer de trouver deux fonctions f et g telles que l'équation $f(t) = g(x)$ soit satisfaite pour toute valeur de t et de x . Vous arriverez rapidement à la conclusion que f et g sont des fonctions constantes et égales.) Soit $-\lambda^2$ cette constante. Alors la solution de l'équation se résume à résoudre les deux équations différentielles ordinaires :

$$X'' = -\lambda^2 X \quad \text{et} \quad T'' = -a^2 \lambda^2 T.$$

séparabilité

L'hypothèse sur la forme de la solution cherchée, c'est-à-dire $u(x, t) = X(x)T(t)$, mène donc à la séparation de l'équation aux dérivées partielles en deux équations ordinaires. Le truc que nous avons utilisé pour *séparer* l'équation de la corde vibrante ne fonctionne pas toujours. Il dépend de l'équation différentielle considérée, du système de coordonnées utilisé et des conditions aux limites. Lorsqu'un problème donné permet la séparation, il est dit *séparable*. Cette propriété remarquable

joue un rôle central dans l'étude des équations différentielles aux dérivées partielles. Nous y reviendrons à la section 3.2.

La première de ces équations a pour solution générale :

$$X(x) = c_1 \cos \lambda x + c_2 \sin \lambda x,$$

où c_1 et c_2 sont deux constantes arbitraires. Puisqu'on exige de la corde que ses extrémités soient fixes, il faut donc que

$$u(0, t) = u(L, t) = 0,$$

c'est-à-dire

$$X(0) = X(L) = 0.$$

La solution ci-dessus donne en $x = 0$:

$$X(0) = 0 = c_1 + 0$$

et donc :

$$c_1 = 0.$$

La condition en $x = L$ donne :

$$X(L) = 0 = c_2 \sin \lambda L$$

qui a pour solution :

$$\lambda_n = \frac{n\pi}{L}, n \in \mathbb{N}.$$

Nous avons étiqueté les différentes solutions de λ par n . Nous ferons de même pour les deux fonctions X et T ainsi que pour leur produit u ; alors $u_n(x, t) = X_n(x)T_n(t)$ est la solution qui correspond à la constante de séparation λ_n . Notons que ce sont les conditions aux limites $u(0, t) = u(L, t) = 0$ qui ont fixé les valeurs possibles λ_n .

L'équation pour $T_n(t)$ est alors un jeu d'enfant à résoudre. Elle se lit maintenant :

$$T_n'' = -\frac{a^2 n^2 \pi^2}{L^2} T_n$$

et sa solution générale est :

$$T_n(t) = A_n \cos \frac{an\pi t}{L} + B_n \sin \frac{an\pi t}{L}$$

où A_n et B_n sont deux constantes arbitraires. La constante de séparation λ_n peut donc être interprétée physiquement comme étant la fréquence du mode u_n étudié. Ainsi, par la remarque à la fin du paragraphe précédent :

les conditions aux limites fixent les fréquences naturelles
du système physique étudié.

L'ensemble des fréquences naturelles d'un système physique est appelé le *spectre* du système. Les mathématiciens parlent également du *spectre d'un opérateur linéaire*; dans le cas présent, l'opérateur linéaire est

$$\frac{d^2}{dx^2}$$

qui agit sur l'espace des fonctions possédant des dérivées partielles secondes et qui s'annulent en $x = 0$ et L . Cette définition est plus abstraite mais fort utile.

Nous venons donc de trouver une famille de solutions étiquetées par $n \in \mathbb{N}$:

$$u_n(x, t) = \left(A_n \cos \frac{an\pi t}{L} + B_n \sin \frac{an\pi t}{L} \right) \sin \frac{n\pi x}{L}.$$

(La constante c_2 dans la solution $X_n(x)$ a été absorbée dans les deux constantes A_n et B_n parce que le produit de deux constantes arbitraires demeure une constante arbitraire. Comme il deviendra clair ci-dessous, nous n'avons perdu aucune généralité en faisant cette redéfinition des A_n et B_n .) Rappelons-nous que nous avons obtenu la solution générale de la corde plombée en faisant une superposition linéaire des solutions particulières. Si nous utilisons l'analogie de cet argument dans le cas présent, nous pouvons affirmer que la solution générale de l'équation de la corde vibrante peut être écrite comme

$$u(x, t) = \sum_{n=1}^{\infty} \left(A_n \cos \frac{an\pi t}{L} + B_n \sin \frac{an\pi t}{L} \right) \sin \frac{n\pi x}{L}.$$

Cette série fait maintenant intervenir un nombre infini de modes propres paramétrés par n puisque le spectre du problème contient un nombre infini de fréquences. A nouveau, nous devons faire face au problème de convergence de cette série. (Voir sections 2.3 et 2.5.)

3. Conditions initiales — Nous allons maintenant donner un exemple de solution particulière en fixant des conditions initiales. Nous allons supposer qu'à l'instant $t = 0$, on donne à la corde un certain profil représenté par le graphe de la fonction $f(x)$. On relâche la corde sans lui donner d'élan. Ces deux hypothèses peuvent être réécrites mathématiquement :

$$u(x, 0) = f(x), \quad \text{profil en } t = 0$$

et

$$\frac{\partial u}{\partial t}(x, 0) = 0, \quad \text{la vitesse de chacun des points est nulle en } t = 0.$$

Utilisant la solution générale ci-dessus évaluée en $t = 0$, nous obtenons :

$$u(x, 0) = f(x) = \sum_{n=1}^{\infty} A_n \sin \frac{n\pi x}{L}$$

et donc les A_n sont précisément les coefficients de Fourier de la fonction f qui donne le profil initial de la corde. En supposant que nous puissions dériver la solution générale terme à terme, la condition « vitesse initiale nulle » mène à

$$\begin{aligned}\frac{\partial u}{\partial t}(x, 0) &= 0 \\ &= \sum_{n=1}^{\infty} \left(-A_n \frac{an\pi}{L} \sin \frac{an\pi t}{L} + B_n \frac{an\pi}{L} \cos \frac{an\pi t}{L} \right) \sin \frac{n\pi x}{L} \Big|_{t=0} \\ &= \sum_{n=1}^{\infty} \left(B_n \frac{an\pi}{L} \right) \sin \frac{n\pi x}{L}.\end{aligned}$$

Cette dernière équation affirme que les coefficients $C_n = B_n \frac{an\pi}{L}$ sont les coefficients du développement en série de Fourier de la fonction $g(x) = 0$ pour tout x . Ces coefficients doivent donc être tous nuls et $B_n = 0$. Ainsi la solution correspondant aux conditions initiales ci-dessus est

$$u(x, t) = \sum_{n=1}^{\infty} A_n \sin \frac{n\pi x}{L} \cos \frac{an\pi t}{L}$$

où les A_n sont les coefficients de Fourier de f :

$$A_n = \frac{2}{L} \int_0^L f(x) \sin \frac{n\pi x}{L} dx.$$

Le parallèle avec la corde plombée nous a aidé à poser l'hypothèse sur la forme de la solution cherchée : $u(x, t) = X(x)T(t)$. Cependant d'autres propriétés remarquables sont apparues lors de la solution de la corde plombée. Premièrement, le problème différentiel s'est vu donner une formulation purement algébrique, la diagonalisation de la matrice A . Deuxièmement, la matrice A obtenue était symétrique, nous assurant un spectre réel et des modes propres orthogonaux. Qu'en est-il ici ? Que veut dire « être symétrique » pour un opérateur différentiel ? Et qu'est-ce qui nous assure que les fréquences propres de la corde vibrante soient réelles ? Ces questions sont importantes et nous y répondrons à la section 3.3.

26*. Pour obtenir le spectre de la corde vibrante, nous avons d'abord obtenu la limite continue de la corde plombée et calculé son spectre par séparation. Il est également possible d'obtenir la limite du spectre de la corde plombée quand le nombre n des poids tend vers l'infini. **Exercices**

(a) Vérifier que les valeurs propres de la matrice A , $n \times n$, intervenant dans le système (*) décrivant la corde plombée sont

$$\lambda_i = 2 \left(1 - \cos \frac{i\pi}{n+1} \right), \quad i = 1, 2, \dots, n$$

et que le vecteur propre $\vec{x}_i$ associé à la valeur propre λ_i a pour composantes

$$x_{i,j} = (\vec{x}_i)_j = \sin \frac{ij\pi}{n+1}.$$

Il s'ensuit que, dans le i -ième mode propre, les n poids se situent précisément sur le graphe de $\sin(ix\pi/(n+1))$ si la longueur totale de la corde est $n+1$.

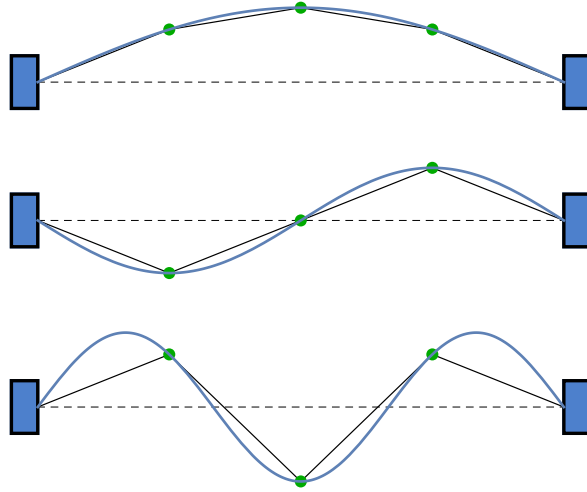


Figure 1.11 – Les trois modes propres de la corde plombée $n = 3$ trouvés précédemment avec, surimposées, les courbes $\sin(ix\pi/(n+1))$ pour $i = 1, 2$ et 3 .

(b) Dans la notation de (a), il est possible de prendre la limite vers la corde vibrante en faisant $n \rightarrow \infty$. Montrer que l'on retrouve bien ses fréquences naturelles

$$\frac{i\pi}{L}, \quad i \in \mathbb{N}.$$

(Se rappeler que n est une mesure de la longueur de la corde L et prendre l'approximation que i est beaucoup plus petit que n .)

(c) Quoiqu'il soit difficile d'imaginer un dispositif physique qui réalise un tel système, supposons maintenant que le système de n masses ne soit plus attaché à deux points fixes mais que la masse m_1 ait comme voisine à sa gauche la masse m_n . Dans cette configuration la corde plombée forme un collier dont les perles sont les masses. Quelle est la matrice B qui représente ce nouveau système physique. Est-elle symétrique ?

(d) Les valeurs propres de B sont-elles réelles ? Comment interpréter la fréquence nulle ? Pouvez-vous obtenir les valeurs propres de B d'une façon aussi complète qu'en (a) ?

27. Dans le traitement de la corde vibrante, nous avons choisi $-\lambda^2$ comme constante de séparation, sous-entendant donc que cette constante est négative ou nulle. Montrer que, si cette constante est choisie strictement positive, l'équation différentielle pour X n'a pas de solution non-nulle qui satisfasse les conditions aux limites.

28. La limite continue du système physique décrit à l'exercice 26 (b) correspond à une corde vibrante dont les extrémités sont attachées à des tiges sur lesquelles elles sont libres de glisser. Alors la corde arrive toujours perpendiculairement à ces deux tiges qui la tendent. En d'autres termes :

$$\frac{\partial u}{\partial x}(0, t) = \frac{\partial u}{\partial x}(L, t) = 0.$$

Obtenir le spectre de ce nouveau problème physique, écrire la solution générale et la solution particulière correspondant aux conditions initiales : $u(x, 0) = 0$ et $\frac{\partial u}{\partial t}(x, 0) = f(x)$.

29*. L'équation décrivant la température $u(x, t)$ le long d'une barre de longueur L et dont la surface est isolée thermiquement est :

$$\frac{1}{a^2} \frac{\partial u}{\partial t} = \frac{\partial^2 u}{\partial x^2},$$

où a^2 est une constante (positive), x est la position le long de la barre ($0 \leq x \leq L$) et t est le temps.

(a) Trouver une famille de solutions de la forme $u_n(x, t) = X_n(x)T_n(t)$ où X_n et T_n ne dépendent que d'une variable, si les deux extrémités de la barre sont maintenues à la température nulle :

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

(Se rappeler que la température d'une telle barre ne peut pas augmenter indéfiniment avec le temps.)

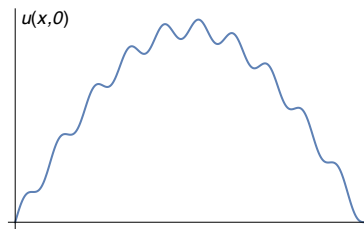


Figure 1.12 – La fonction température u au temps $t = 0$

(b) Dire comment les solutions obtenues en (a) peuvent donner la solution de l'équation de la chaleur si, initialement, la température le long de la barre est

$$u(x, 0) = f(x) \quad \text{et} \quad f(0) = f(L) = 0.$$

(c) Supposons que f ait qualitativement la forme ci-contre. Tracer qualitativement $u(x, t_1)$ pour $t_1 = \frac{L^2}{a^2 \pi^2} \ln 2$. Justifier.

Remarque : L'équation de la chaleur est un des principaux acteurs dans l'histoire rocambolesque du premier câble transatlantique. A lire dans Körner, pp. 332-337.

Chapitre 2

Convergence des séries de Fourier

À la fin de la section 1.1, nous avons soulevé les questions fondamentales reliant la fonction f à sa série de Fourier. Rappelons ces questions :

- Est-ce que les sommes partielles $s_m = a_0/2 + \sum_{n=1}^m (a_n \cos nx + b_n \sin nx)$ tendent vers $f(x)$ en tout point $x \in [-\pi, \pi)$? Sous quelles conditions ?
- Peut-on reconstruire f en tout point $x \in [-\pi, \pi)$ même si les sommes partielles ne convergent pas ?
- Peut-on mesurer la « distance » entre les sommes partielles s_m et la fonction originale f ? Cette distance tend-elle vers 0 lorsque m tend vers l'infini ?

Chacune de ces questions se verra consacrer une section du présent chapitre. Les réponses sont parfois surprenantes. Les sections 2.1 et 2.4 rappellent les outils d'analyse et d'algèbre linéaire dont nous aurons besoin.

La présentation des théorèmes fondamentaux de ce chapitre (sections 2.2, 2.3 et 2.5) suit de près celle de Körner.

2.1 Rappel d'analyse

L'étude de la convergence ponctuelle et en moyenne des séries de Fourier reposent sur les concepts d'analyse suivants :

1. la continuité,
2. la convergence d'une série et d'une suite,
3. l'intégrale de Riemann,
4. la convergence uniforme.

Dans cette section, nous rappelons rapidement ces concepts. Des exemples sont donnés pour rafraîchir ou développer l'intuition. Je suis la présentation de Spivak, *Calculus*, 2nd ed., Publish or Perish (1980).

continuité

2.1.1 La continuité

Définition 1. (a) Une fonction $f : (a, b) \rightarrow \mathbb{R}$ est continue en $y \in (a, b)$ si, pour tout $\epsilon > 0$, il existe $\delta > 0$ tel que, si $|y - x| < \delta$ alors $|f(y) - f(x)| < \epsilon$.
 (b) Une fonction f est continue sur un intervalle ouvert si elle est continue en tout point de cet ouvert.

Le concept de continuité peut être mis en mots de la façon suivante : une fonction f est continue en y si la valeur $f(x)$ est arbitrairement proche de la valeur $f(y)$ pour tout x suffisamment proche de y . Rappelons que $|a - l| < \epsilon$ signifie que a appartient au segment $(l - \epsilon, l + \epsilon)$. (Exercice !)

Les fonctions continues ont les propriétés suivantes.

1. Si f et g sont continues en a , alors $f + g$ et $f \cdot g$ le sont également.
2. Si g est continue en a et f l'est en $g(a)$, alors $f \circ g$ l'est en a .

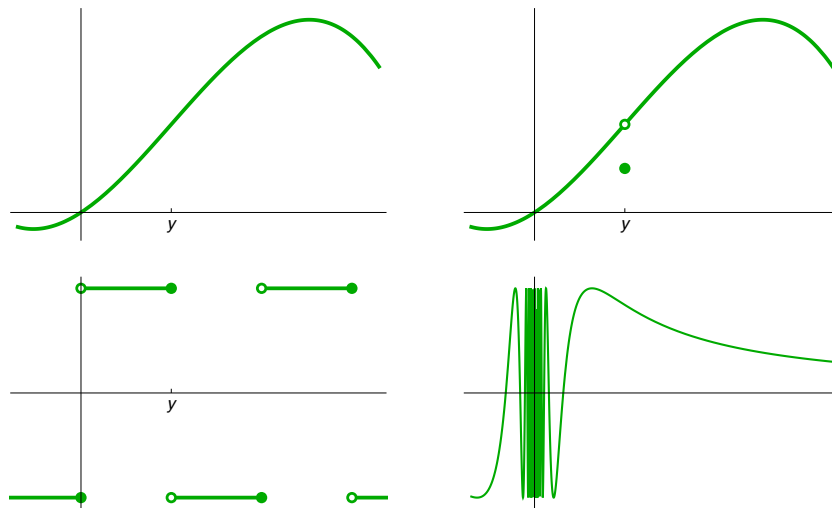


Figure 2.1 – Quatre graphes : seul le premier (en haut à gauche) représente une fonction continue sur son domaine.

La figure 2.1 présentent les graphes de quatre fonctions. Les deux fonctions du haut ne diffèrent qu'au point y . Pour la première, toutes les valeurs de $f(x)$ pour x suffisamment proche de y sont arbitrairement proche de la valeur de $f(y)$. Cette fonction est donc continue en y . Le même argument peut être répété aux autres points de l'intervalle. Cette fonction est donc continue sur tout l'intervalle représenté. Sa voisine coïncide avec la première à l'exception du point y . Elle n'est pas continue en ce point car il existe des points x proches de y qui prennent des valeurs fort loin de la valeur $f(y)$ donnée par le point. Cette fonction n'est donc pas continue en y ; elle l'est cependant partout ailleurs.

Le troisième graphe est celui de la fonction créneau. Les points (pleins) indiquent les valeurs de la fonction aux points de saut. Remarquons qu'en chacun de ces points de saut, tout point x à gauche et suffisamment proche d'un saut possède la valeur de la fonction en ce saut. Cependant les points immédiatement à droite ne le sont pas. La définition de continuité impose une fenêtre autour de y contenant des points à gauche et à droite. Ainsi la fonction créneau n'est continue en aucun de ses sauts. En dehors de ces derniers, elle est continue.

Ces premiers exemples sont assez évidents. Il est difficile de comprendre sur ceux-ci pourquoi une définition aussi abstraite est nécessaire. Le prochain exemple est plus délicat. Le quatrième graphe est celui de la fonction $f : \mathbb{R} \rightarrow \mathbb{R}$ donnée par $\sin \frac{1}{x}$ si $x \neq 0$ et par $f(0) = 0$ en $x = 0$. Pour tout $x \neq 0$, la fonction f se comporte bien et son graphe est bien lisse. Autour de $x = 0$, la fonction f oscille « sauvagement ». La fonction f n'est pas continue en ce point comme l'argument suivant le montre. Choisissons $\epsilon < 1$ et supposons qu'il existe un $\delta > 0$ tel que, si $|x - 0| < \delta$, alors $|f(x) - f(0)| = |f(x)| < \epsilon$. Mais, par la propriété archimédienne des nombres réels, il existe un $n \in \mathbb{N}$ tel que

$$0 < \frac{1}{2\pi(n+1)} < \frac{1}{2\pi n} < \delta$$

et il existe donc un x tel que $|x| < \delta$ et $f(x) = 1 > \epsilon$. (Le nombre $x = \frac{1}{2\pi(n+1/4)}$ satisfait par exemple ces conditions.) C'est donc une contradiction et f n'est pas continue en $x = 0$. En fait, une preuve très semblable montre que, même si f prenait une autre valeur α en $x = 0$, elle n'y serait tout de même pas continue. Ainsi f est continue sur $\mathbb{R} \setminus \{0\}$, mais pas en $x = 0$.

2.1.2 La convergence d'une suite et d'une série

Une suite est un ensemble de nombres réels étiquetés par les entiers naturels : $\{a_1, a_2, \dots, a_n, \dots\}$. Une suite contient donc un nombre infini de nombres.

Définition 2. Une suite converge vers l si, pour tout $\epsilon > 0$, il existe $N \in \mathbb{N}$ tel que, si $n > N$, alors $|a_n - l| < \epsilon$. On écrit $\{a_n\} \rightarrow l$, $a_n \rightarrow l$ ou encore $\lim_{n \rightarrow \infty} a_n = l$. convergence d'une suite

En mots courants cette définition affirme qu'une suite converge vers l si, quelque soit $\epsilon > 0$, tous les éléments de la suite sont à une distance de l plus petite que ϵ à l'exception, peut-être, des premiers N éléments. (Le nombre N doit être fini.)

Exemple 1. La suite $\{1, \frac{1}{2}, \frac{1}{3}, \dots, \frac{1}{n}, \dots\} \rightarrow 0$. Que cette suite converge vers 0 devrait être intuitivement clair. En voici la preuve. Soit $\epsilon > 0$. On choisira N tel que $N > \frac{1}{\epsilon}$, $N \in \mathbb{N}$. Alors, si $n > N$, on a $\frac{1}{N} > \frac{1}{n}$ et donc

$$\left| \frac{1}{n} - l \right| = \left| \frac{1}{n} - 0 \right| = \frac{1}{n} < \frac{1}{N} < \epsilon.$$

Exemple 2. La suite $\{-1, +1, -1, +1, \dots, (-1)^n, \dots\}$ ne converge pas. A nouveau, cet affirmation devrait être assez claire. La preuve procède comme suit. Si $\epsilon = \frac{1}{2} > 0$ et la limite était l , alors quelque soit N , les éléments pairs de la suite avec $n > N$ seraient tels que $|a_n - l| = |1 - l|$ alors que les impairs seraient tels que $|a_n - l| = |-1 - l|$. Il faut donc que

$$|1 - l| \text{ et } |-1 - l| < \frac{1}{2}.$$

Mais ceci est impossible car ceci signifie que l est à une distance plus petite que $\frac{1}{2}$ à la fois de 1 et de -1 .

Soit $\{a_1, a_2, \dots, a_n, \dots\}$ une suite. Considérons les sommes partielles

$$s_1 = a_1,$$

$$s_2 = a_1 + a_2,$$

$$s_3 = a_1 + a_2 + a_3,$$

$$\dots,$$

$$s_n = \sum_{i=1}^n a_i,$$

$$\dots$$

L'ensemble $\{s_1, s_2, s_3, \dots, s_n, \dots\}$ est lui-même une suite.

convergence
d'une série

Définition 3. La suite $\{a_1, a_2, \dots, a_n, \dots\}$ est dite sommable si la suite $\{s_1, s_2, s_3, \dots, s_n, \dots\}$ converge. Si la suite $\{s_1, s_2, s_3, \dots, s_n, \dots\}$ converge vers s , on dit alors que la série $\sum_{i=1}^{\infty} a_i$ converge vers s et on écrit

$$\sum_{i=1}^{\infty} a_i = s.$$

Les suites sommables ont les propriétés suivantes.

1. Si $\{a_n\}$ et $\{b_n\}$ sont sommables, alors $\{a_n + b_n\}$ l'est aussi et $\sum(a_n + b_n) = \sum a_n + \sum b_n$.
2. Si $\{a_n\}$ est sommable, alors $\{ca_n\}$ l'est aussi pour tout nombre c . De plus $\sum ca_n = c \sum a_n$.
3. Si $\{a_n\}$ est sommable, alors $\lim_{n \rightarrow \infty} a_n = 0$.

suite géométrique

Exemple 3. Soit la suite $\{1, r, r^2, r^3, \dots, r^{n-1}, \dots\}$ pour $|r| < 1$. Cette suite est-elle sommable? Pour y répondre, notons d'abord qu'il existe une expression simple pour la somme partielle $s_n = \sum_{i=0}^{n-1} r^i$. En effet :

$$s_n = 1 + r + r^2 + \dots + r^{n-1}$$

et

$$rs_n = r + r^2 + r^3 \cdots + r^n.$$

Donc $rs_n - s_n = r^n - 1$ et

$$s_n = \frac{1 - r^n}{1 - r}.$$

Alors $\sum_{i=0}^{\infty} r^i$ existera si $\lim_{n \rightarrow \infty} \frac{1 - r^n}{1 - r}$ existe. Mais

$$\lim_{n \rightarrow \infty} \frac{1 - r^n}{1 - r} = \frac{1}{1 - r} - \frac{1}{1 - r} \lim_{n \rightarrow \infty} r^n = \frac{1}{1 - r}$$

puisque $|r| < 1$. Alors la suite $\{1, r, r^2, r^3, \dots, r^{n-1}, \dots\}$ converge et

$$\sum_{n=1}^{\infty} r^{n-1} = \sum_{n=0}^{\infty} r^n = \frac{1}{1 - r}.$$

Exemple 4. Soit la suite $\{1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \frac{1}{5}, \dots\}$. Cette suite n'est pas sommable. Pour le voir, il suffit de remarquer que les termes de la somme peuvent être regroupés de la façon suivante :

$$\begin{aligned} & 1 + \underbrace{\frac{1}{2}}_{\geq \frac{1}{2}} \\ & + \underbrace{\frac{1}{3} + \frac{1}{4}}_{\geq \frac{1}{2}} \\ & + \underbrace{\frac{1}{5} + \frac{1}{6} + \frac{1}{7} + \frac{1}{8}}_{\geq \frac{1}{2}} \\ & + \underbrace{\frac{1}{9} + \frac{1}{10} + \frac{1}{11} + \frac{1}{12} + \frac{1}{13} + \frac{1}{14} + \frac{1}{15} + \frac{1}{16}}_{\geq \frac{1}{2}} + \dots \end{aligned}$$

Remarquons que la réciproque de la troisième propriété ci-dessus n'est donc pas vraie : une suite $\{a_n\}$ peut ne pas être sommable même si $\lim_{n \rightarrow \infty} a_n = 0$.

2.1.3 L'intégrale de Riemann

L'intégration n'est pas un concept difficile à comprendre ; on l'introduit habituellement comme l'aire sous le graphe. C'est pourtant un concept difficile à définir rigoureusement et il existe en fait plusieurs définitions non-équivalentes. L'intégrale de Riemann est habituellement celle définie dans les premiers cours d'analyse.

Définition 4. Soit $a < b$. Une partition de l'intervalle $[a, b]$ est une collection finie de points distincts $\in [a, b]$, un étant a et un autre étant b . (On prendra pour acquis que ces points sont ordonnés $a = t_0 < t_1 < \dots < t_{n-1} < t_n = b$.)

Définition 5. Soit f bornée sur $[a, b]$ et $P = \{t_0, t_1, \dots, t_n\}$ une partition de $[a, b]$. Soit

$$m_i = \inf \{f(x) : t_{i-1} \leq x \leq t_i\}$$

$$M_i = \sup \{f(x) : t_{i-1} \leq x \leq t_i\}.$$

La somme inférieure de f pour P , notée $I(f, P)$, est définie par :

$$I(f, P) = \sum_{i=1}^n m_i(t_i - t_{i-1}).$$

La somme supérieure de f pour P , notée $S(f, P)$, est définie par :

$$S(f, P) = \sum_{i=1}^n M_i(t_i - t_{i-1}).$$

intégrale de Riemann

Définition 6. Une fonction f bornée sur $[a, b]$ est intégrable (au sens de Riemann) sur $[a, b]$ si

$$\sup_P \{I(f, P)\} = \inf_P \{S(f, P)\}.$$

Dans ce cas, on note ce nombre commun par $\int_a^b f$ ou $\int_a^b f(x)dx$.

L'intégration au sens de Riemann est celle que nous utiliserons par la suite. Certaines fonctions ne sont pas intégrables en ce sens. Par exemple la fonction

$$f(x) = \begin{cases} 0, & x \text{ irrationnel,} \\ 1, & x \text{ rationnel.} \end{cases}$$

ne l'est pas sur l'intervalle $[a, b]$. Il est aisé de s'en convaincre en effet puisque, quelque soit la partition P , les sommes inférieure et supérieure sont $L(f, P) = 0$ et $U(f, P) = (b - a)$. Alors

$$\sup_P \{L(f, P)\} = 0 \neq (b - a) = \inf_P \{U(f, P)\}$$

et l'intégrale de f sur $[a, b]$ n'existe pas.

Toutes les propriétés usuelles de l'intégration vues dans les cours de calcul demeurent vraies pour les fonctions qui sont intégrables.

1. Si f est intégrable sur $[a, b]$ et $c \in [a, b]$, alors f est intégrable sur $[a, c]$ et sur $[c, b]$ et $\int_a^b f = \int_a^c f + \int_c^b f$.
2. Si f est continue sur $[a, b]$, alors f est intégrable.

3. Si f et g sont intégrables sur $[a, b]$, alors la somme $f + g$ est aussi intégrable et $\int_a^b (f + g) = \int_a^b f + \int_a^b g$.
4. Si f est intégrable sur $[a, b]$, cf où $c \in \mathbb{R}$ l'est aussi et $\int_a^b cf = c \int_a^b f$.
5. Si f est intégrable sur $[a, b]$, alors $F(x) = \int_a^x f$ est une fonction continue sur $[a, b]$.

2.1.4 Convergence uniforme

A la prochaine section, nous commencerons l'étude de la convergence des séries de Fourier. Pour ces séries, chacun des termes est une fonction. Par exemple, le développement de $f(x) = x^2$, $-\pi \leq x \leq \pi$, est

$$\frac{\pi^2}{3} + 4 \sum_{n=1}^{\infty} \frac{(-1)^n}{n^2} \cos nx$$

et donc $a_0 = \frac{\pi^2}{3}$, $a_1 = -4 \cos x$, $a_2 = \cos 2x$, $a_3 = -\frac{4}{9} \cos 3x$, etc. Ainsi chacun des a_i est une fonction et établir la convergence de la série pour chaque valeur de x est un nouveau problème. (On remarquera que l'indice de la suite commence ici à 0 et non à 1. Nous prendrons cette liberté quand elle sera utile.)

Le grand problème des fonctions définies par des suites $\{f_n\}$ est que, même si les f_n ont les meilleures propriétés (continuité, différentiabilité, etc.), la fonction limite f , même si elle existe, peut vraiment être « sauvage ». Voici un exemple. (Voir la figure ci-dessus.) Soit la suite

$$f_n(x) = \begin{cases} x^n, & 0 \leq x \leq 1, \\ 1, & 1 \leq x. \end{cases}$$

On montre aisément que $f_n \rightarrow f$ avec

$$f(x) = \begin{cases} 0, & 0 \leq x < 1, \\ 1, & 1 \leq x. \end{cases}$$

Les fonctions f_n sont continues sur tout le demi-axe $\{x \mid x \geq 0\}$. Cependant la fonction limite f n'est pas continue en $x = 1$. Donc la limite de fonctions continues peut ne pas être continue.

Un certain type de convergence (plus exigeant que la convergence en chaque point) est la convergence uniforme définie comme suit.

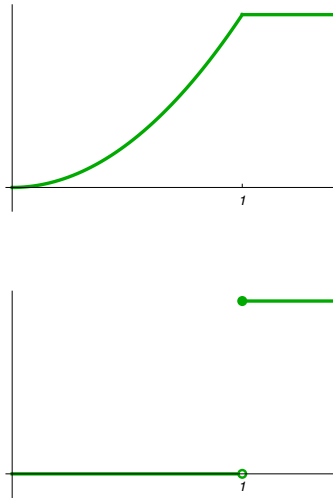


Figure 2.2 – La fonction f_2 (a) et la limite f (b)

convergence uniforme

Définition 7. Une suite $\{f_n\}$ de fonctions définies sur le domaine A converge uniformément vers une fonction f sur A si, pour tout $\epsilon > 0$, il existe $N \in \mathbb{N}$ tel que, pour tout $x \in A$, si $n > N$ alors $|f(x) - f_n(x)| < \epsilon$.

La seule différence entre « convergence de $\{f_n\}$ » et « convergence uniforme de $\{f_n\}$ » repose dans l'ordre des quantificateurs : « $\exists N \in \mathbb{N}$ tel que $\forall x \in A$ » pour la continuité uniforme plutôt que « $\forall x \in A, \exists N \in \mathbb{N}$ » pour la continuité. Cette nouvelle exigence demande donc qu'il soit possible de choisir un $N \in \mathbb{N}$ commun à tous les $x \in A$.

La convergence uniforme est extrêmement utile. En voici quelques propriétés.

1. Soit $\{f_n\}$ une suite de fonctions intégrables sur l'intervalle $[a, b]$ et qui converge uniformément sur $[a, b]$ vers une fonction f qui est elle-même intégrable. Alors $\int_a^b f = \lim_{n \rightarrow \infty} \int_a^b f_n$.
2. Si les f_n sont continues sur $[a, b]$ et convergent uniformément sur $[a, b]$ vers f , alors f est continue.

Exercices

1*. Les fonctions suivantes sont-elles continues en $x = 0$?

(a) $f(x) = \frac{x+1}{x-1}$

(b) $f(x) = \frac{\tan x}{x}$ lorsque $x \neq 0$ et $f(0) = 1$.

(c) $f(x) = x \sin \frac{1}{x}$ si $x \neq 0$ et $f(0) = 0$.

2. Prouver les propriétés des fonctions continues citées dans le texte. (La seconde est un peu plus difficile.)

3*. Décider si les suites suivantes convergent ou non et le montrer.

(a) $a_n = \frac{n-1}{n+1}$

(b) $a_n = nc^n$ où $|c| < 1$.

(c) $a_n = \frac{c^n - d^n}{c^n + d^n}$. Dans ce cas préciser la limite en fonction de c et d .

4*. Décider si les séries suivantes convergent ou non et le montrer.

(a) $\sum_{n=0}^{\infty} \frac{1}{n!}$

(b) $\sum_{n=2}^{\infty} \frac{1}{\ln n}$

5. Démontrer les propriétés sur les séries citées dans le texte.

6. (a) Montrer que la fonction constante $f(x) = c$ est intégrable au sens de Riemann.

(b) Montrer que la fonction $f(x) = x$ est intégrable au sens de Riemann et obtenir $\int_0^a f$. (Plus difficile.)

7*. Déterminer la fonction f vers laquelle converge les suites $\{f_n\}$ suivantes et décider si les suites convergent uniformément.

(a) $f_n(x) = \sum_{i=1}^n x^i$ sur l'intervalle $(-1, 1)$

(b) Même fonction sur l'intervalle $[-\frac{1}{2}, \frac{1}{2}]$

(c) $f_n(x) = e^{-nx^2}$ sur l'intervalle $[-1, 1]$

2.2 Le théorème de Fejér

Après plus d'un chapitre sur les séries de Fourier, il est temps d'avouer que la série de Fourier ne converge pas tout le temps vers la fonction originale f même si f est continue. En fait :

Théorème 1 (DuBois-Reymond (1896)). *Il existe une fonction continue $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ telle que ses sommes partielles vérifient*

$$\limsup_{m \rightarrow \infty} |s_m(0)| = \infty.$$

En d'autres mots, il existe une fonction *continue* dont la série de Fourier évaluée en $x = 0$ diverge ! La preuve de ce théorème mettait fin aux espoirs de plusieurs grands mathématiciens, Riemann, Weierstraß et Dedekind par exemple, qui avaient cru que la continuité d'une fonction serait une condition suffisante pour assurer la convergence de sa série de Fourier. Cette absence de convergence (en un point de $\mathbb{T}^1$) n'est que la pointe de l'iceberg cependant. En effet, il y a quelques années Kahane et Katznelson (1966) ont montré que, donné un sous-ensemble $S \subset \mathbb{T}^1$ de mesure nulle¹, il est possible de construire une fonction continue $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ telle que $\limsup_{m \rightarrow \infty} |s_m(s)| = \infty$ pour tout $s \in S$. Un exemple d'ensemble de mesure nulle est le sous-ensemble des nombres rationnels sur l'intervalle $[-\pi, \pi)$. Ainsi il existe une fonction continue $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ dont la série de Fourier diverge en tous les nombres rationnels... Ceci semble décourageant. Ce ne l'est pas complètement car la plupart des fonctions que nous avons rencontrées converge en presque tous les points de $\mathbb{T}^1$. La construction d'une fonction continue dont la série diverge en un point est faite dans Körner (chapitre 18). Nous ne la référons pas. Plutôt nous nous concentrerons tout d'abord sur notre seconde question :

1. Le concept de mesure d'un ensemble est habituellement défini dans un cours sur l'intégrale de Lebesgue ou sur la théorie de la mesure.

— Peut-on reconstruire f en tout point $x \in D$ même si les sommes partielles ne convergent pas ?

La réponse contenue dans le théorème de Fejér est simplement : oui, si la fonction est continue. A première vue, ceci semble en contradiction avec le théorème de DuBois-Reymond. Mais ce ne l'est pas comme nous le verrons à l'instant. Le théorème de Fejér ainsi que sa preuve sont remarquablement beaux et satisfaisants. Il jouera aussi un rôle crucial dans la preuve de la convergence en moyenne des fonctions continues (section 2.5).

La preuve de Fejér est basée sur l'observation que, même si les sommes partielles

$$s_m = \sum_{n=-m}^m c_n e^{inx}$$

ne convergent pas, les moyennes des sommes partielles

$$\sigma_m = \frac{1}{m+1} \sum_{n=0}^m s_n$$

peuvent, elles, converger. Que représentent ces σ_m ? Si les fonctions s_m approchent bien f quand $m \rightarrow \infty$, alors elles devraient être toutes à peu près égales à f pour m suffisamment grand. Donc la moyenne devrait être aussi à peu près f . En fait les nombres σ_m se comportent mieux que les nombres s_m !

Lemme 2. (i) Si $s_n \rightarrow s$, alors

$$\sigma_n = \frac{1}{n+1} \sum_{j=0}^n s_j \rightarrow s.$$

(ii) Il existe des suites $\{s_n\}$ qui n'ont pas de limite, mais pour lesquelles $\{\sigma_m\} \rightarrow l$ existe.

Démonstration. (i) Soit $\epsilon > 0$. Puisque $s_n \rightarrow s$, alors il existe $N(\epsilon)$ tel que $|s_n - s| \leq \frac{\epsilon}{2}$ si $n > N(\epsilon)$. Soit $A = \sum_{j=0}^{N(\epsilon)} |s_j - s|$; choisir $M(\epsilon) \geq N(\epsilon)$ tel que $M(\epsilon) \geq \frac{2A}{\epsilon}$. Alors, si $n > M(\epsilon) \geq N(\epsilon)$:

$$\begin{aligned} \left| \frac{1}{n+1} \sum_{j=0}^n s_j - s \right| &= \frac{1}{n+1} \left| \sum_{j=0}^n (s_j - s) \right| \leq \frac{1}{n+1} \sum_{j=0}^n |s_j - s| \\ &= \frac{1}{n+1} \left(\sum_{j=0}^{N(\epsilon)} |s_j - s| + \sum_{j=N(\epsilon)+1}^n |s_j - s| \right) \\ &\leq \frac{1}{n+1} \left(A + (n - N(\epsilon)) \frac{\epsilon}{2} \right) \\ &< \frac{1}{n+1} \left(\frac{\epsilon}{2} (n+1) + (n+1) \frac{\epsilon}{2} \right) = \epsilon, \end{aligned}$$

où la dernière étape vient du fait que $n > M \geq \frac{2A}{\epsilon}$ et donc $\frac{\epsilon n}{2} > A$.

(ii) La suite $s_n = (-1)^n$, nous l'avons vu, ne converge pas. Cependant

$$\left| \frac{1}{n+1} \sum_{j=0}^n s_j \right| = \frac{1}{n+1} \underbrace{\left| \sum_{j=0}^n s_j \right|}_{=1 \text{ ou } 0} \leq \frac{1}{n+1} \rightarrow 0.$$

Donc les $\sigma_n \rightarrow 0$. □

Ce lemme justifie la définition suivante.

Définition 8. Une suite $\{a_n, n = 0, 1, 2, \dots\}$ est dite convergente au sens de convergence
Cesàro si la suite des moyennes $\{\sigma_n = \frac{1}{n+1} \sum_{j=0}^n a_j\}$ converge au sens usuel. au sens de Cesàro

Une suite $\{a_n, n = 0, 1, 2, \dots\}$ est dite sommable au sens de Cesàro si la sommabilité
suite des moyennes des sommes partielles $\{\sigma_n = \frac{1}{n+1} \sum_{j=0}^n s_j\}$ où $s_n = \sum_{i=0}^n a_i$ au sens de Cesàro
converge au sens usuel.

Nous allons étudier la convergence des s_n au sens de Cesàro et donc construire Ernesto Cesàro
explicitement les σ_n pour les séries complexes. Soit $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ une fonction (1859-1906)
Riemann-intégrable. Alors ses coefficients c_r existent. Notons par $s_j(f, t)$ la j -
ième somme partielle évaluée en t , c'est-à-dire $s_j(f, t) = \sum_{r=-j}^j c_r e^{irt}$, et par
 $\sigma_n(f, t)$ la moyenne de ses n premières sommes partielles évaluée en t , c'est-à-
dire $\frac{1}{n+1} \sum_{j=0}^n s_j(f, t)$. Nous pouvons réécrire σ_n comme suit :

$$\begin{aligned} \sigma_n(f, t) &= \frac{1}{n+1} \sum_{j=0}^n s_j(f, t) = \frac{1}{n+1} \sum_{j=0}^n \sum_{r=-j}^j c_r e^{irt} \\ &= \frac{1}{n+1} \sum_{r=-n}^n (n+1 - |r|) c_r e^{irt}. \end{aligned}$$

(La dernière égalité est-elle évidente ? Pour vous en convaincre, comptez le nombre
de fois que le terme $r = 0$ apparaît dans la double somme $\sum_{j=0}^n \sum_{r=-j}^j$, puis le
nombre de fois que $r = 1$ apparaît, etc.) On peut réécrire ces sommes de Cesàro
sous une forme plus élégante :

$$\begin{aligned} \sigma_n(f, t) &= \sum_{r=-n}^n \frac{n+1 - |r|}{n+1} c_r e^{irt} \\ &= \sum_{r=-n}^n \frac{n+1 - |r|}{n+1} \frac{1}{2\pi} \left(\int_{\mathbb{T}^1} f(x) e^{-irx} dx \right) e^{irt} \\ &= \sum_{r=-n}^n \frac{n+1 - |r|}{n+1} \frac{1}{2\pi} \int_{\mathbb{T}^1} f(x) e^{ir(t-x)} dx \\ &= \frac{1}{2\pi} \int_{\mathbb{T}^1} f(x) K_n(t-x) dx \end{aligned}$$

où nous avons introduit le noyau de Fejér

noyau de Fejér

$$K_n(x) = \sum_{r=-n}^n \frac{n+1-|r|}{n+1} e^{irx}.$$

Faisant la substitution $y = t - x$, nous pouvons réécrire :

$$\begin{aligned} \sigma_n(f, t) &= \frac{1}{2\pi} \int_{t-\pi}^{t+\pi} f(t-y) K_n(y) (-dy) \\ &= \frac{1}{2\pi} \int_{t-\pi}^{t+\pi} f(t-y) K_n(y) dy \\ &= \frac{1}{2\pi} \int_{\mathbb{T}^1} f(t-y) K_n(y) dy. \end{aligned}$$

Le théorème de Fejér reposera presque entièrement sur la famille de fonctions K_n . La figure 2.3 présente les fonctions de Fejér pour $n = 1, 2, 5$ et 10 . Nous présentons tout d'abord deux lemmes qui en décrivent les propriétés et nous reviendrons après à une discussion plus précise de leur graphe.

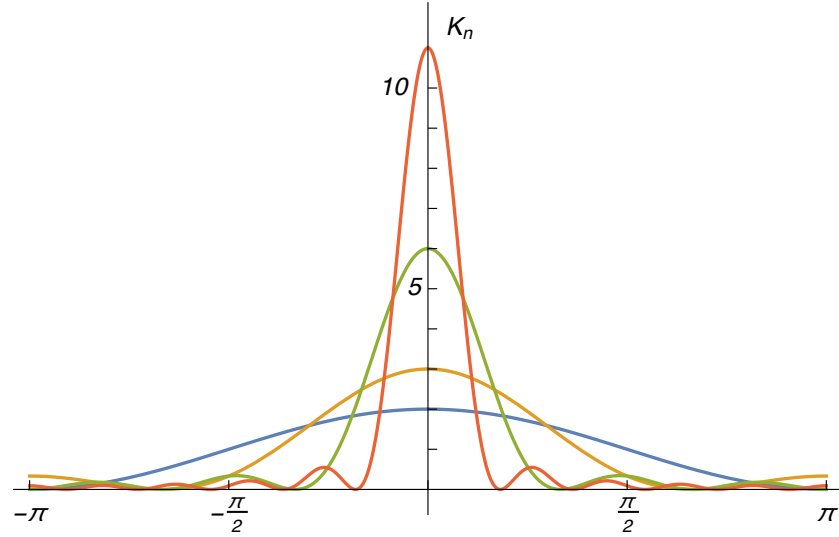


Figure 2.3 – Les fonctions de Fejér K_n pour $n = 1, 2, 5, 10$.

Lemme 3. La fonction K_n a la forme :

$$K_n(x) = \frac{1}{n+1} \left(\frac{\sin \frac{(n+1)x}{2}}{\sin \frac{x}{2}} \right)^2, \quad x \neq 0,$$

$$K_n(0) = n+1.$$

Démonstration. Si $x = 0$, alors

$$\begin{aligned} K_n(0) &= \sum_{r=-n}^n \frac{n+1-|r|}{n+1} 1 = \sum_{r=-n}^n 1 - \frac{1}{n+1} \sum_{r=-n}^n |r| \\ &= 2n+1 - \frac{2}{n+1} \sum_{r=1}^n r = 2n+1 - \frac{2}{n+1} \frac{(n+1)n}{2} \\ &= 2n+1 - n = n+1. \end{aligned}$$

Si $x \neq 0$, alors :

$$\sum_{r=-n}^n (n+1-|r|)e^{irx} = \left(\sum_{k=0}^n \exp i(k - \frac{n}{2})x \right)^2.$$

Voici un tour de passe-passe presque miraculeux ! Pourtant cette égalité n'est pas trop difficile à comprendre. Pour vous en convaincre, écrivez d'abord l'égalité pour $n = 1$. Le membre de gauche contiendra trois termes pour lesquels le facteur $(n+1-|r|)$ sera respectivement 1, 2 et 1. La somme du membre de droite contiendra deux termes et, après le carré, elle contiendra des exponentielles d'arguments $-ix, 0$ et $+ix$. Dans ce carré, le second terme peut être obtenu de deux façons différentes alors que le premier et le troisième ne peuvent être obtenus que d'une seule façon. Leur facteur est donc 1, 2 et 1 comme au membre de gauche, ce qui prouve l'égalité pour $n = 1$. Pour faire le cas n quelconque, il suffit de se demander de combien de façons les exponentielles avec argument irx avec $-n \leq r \leq n$ peuvent être produites par le carré du membre de droite. (Rappelez-vous, dans le cas $n = 1$, le cas $k = 0$ était produit de deux façons différentes.) Il est aisé de se convaincre que l'exponentielle d'argument irx est obtenue de $(n-1-|r|)$, ce qui démontre l'égalité. Ainsi la somme ci-dessus devient

$$\begin{aligned} &= \left(\exp(-\frac{i}{2}nx) \sum_{k=0}^n \exp ikx \right)^2 \\ &= \left(\exp(-\frac{i}{2}nx) \frac{1 - \exp ix(n+1)}{1 - \exp ix} \right)^2 \end{aligned}$$

où nous avons utilisé le fait que $\sum_{k=0}^n r^k = (1 - r^{n+1})/(1 - r)$ (voir paragraphe 2.1.2). Donc

$$\begin{aligned} &= \left(\frac{\exp(-\frac{i}{2}nx) - \exp(\frac{i}{2}nx + ix)}{\exp \frac{i}{2}x (\exp -\frac{i}{2}x - \exp \frac{i}{2}x)} \right)^2 \\ &= \left(\frac{\exp -\frac{i}{2}(n+1)x - \exp \frac{i}{2}(n+1)x}{\exp -\frac{i}{2}x - \exp \frac{i}{2}x} \right)^2 \\ &= \left(\frac{\sin \frac{(n+1)x}{2}}{\sin \frac{x}{2}} \right)^2 \end{aligned}$$

qui est le résultat désiré. $\square$

Le prochain lemme ajoute à ces premières propriétés.

Lemme 4. (i) $K_n(x) \geq 0$ sur $\mathbb{T}^1$.
(ii) $K_n(x) \rightarrow 0$ uniformément en dehors de $[-\delta, \delta]$, $\forall \delta > 0$.
(iii) $\frac{1}{2\pi} \int_{\mathbb{T}^1} K_n(x) dx = 1$.

Avant d'entreprendre la preuve de ces propriétés, il est utile de donner intuitivement leur utilité. Qu'essayons-nous de montrer ? Que la suite des $\sigma_n(f, t)$ définie par

$$\sigma_n(f, t) = \frac{1}{2\pi} \int_{\mathbb{T}^1} f(t-y) K_n(y) dy$$

converge vers f pour tout t . La seule chose qui change dans la définition de σ_n lorsque $n \rightarrow \infty$ est le noyau de Fejér K_n . Ainsi les propriétés de convergence de la série de Fourier de f dépend de façon importante des propriétés du noyau K_n . Quand on y regarde de plus près, on remarque que, pour que l'intégrale

$$\frac{1}{2\pi} \int_{\mathbb{T}^1} f(t-y) K_n(y) dy \rightarrow f(t),$$

il faut que K_n soit très très grand proche de $y = 0$, mais presque 0 partout ailleurs. Le lemme ci-dessus donnait la valeur en $x = 0$: $K_n(0) = n+1$ qui effectivement devient énorme quand $n \rightarrow \infty$ et le lemme présent (la partie (ii)) montre justement que K_n , lorsque n est assez grand, est « pratiquement » nul en dehors d'un tout petit intervalle centré en $x = 0$, malgré que son intégrale sur le cercle demeure égale à 2π (partie (iii)). Les graphiques de la figure 2.3 confirme cette intuition. Donc ce lemme est crucial.

Démonstration. (i) K_n est le carré d'un nombre réel par le lemme précédent.
(ii) Puisque $\sin \frac{x}{2}$ est monotone croissante sur l'intervalle $[-\pi, \pi]$, $\sin \frac{\delta}{2} < |\sin \frac{x}{2}|$ pour tout x tel que $\pi \geq |x| > \delta$. Ainsi

$$K_n(x) \leq \frac{1}{n+1} \left(\frac{1}{\sin \frac{x}{2}} \right)^2 \leq \frac{1}{n+1} \left(\frac{1}{\sin \frac{\delta}{2}} \right)^2 \rightarrow 0$$

lorsque $n \rightarrow \infty$ et ce, pour tout $\delta < |x| \leq \pi$. (La convergence est uniforme puisque la rapidité de convergence ne dépend que du δ choisi et non du point x .)

(iii) Calcul direct :

$$\frac{1}{2\pi} \int_{\mathbb{T}^1} K_n(x) dx = \frac{1}{2\pi} \sum_{r=-n}^n \frac{n+1-|r|}{n+1} \underbrace{\int_{\mathbb{T}^1} e^{-irx} dx}_{=2\pi\delta_{r,0}} = \frac{n+1}{n+1} = 1.$$

$\square$

Théorème 5 (Fejér (1904)). (i) Soit $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ une fonction intégrable. Si f est continue en t , alors Lipót Fejér
(1880-1959)

$$\sigma_n(f, t) \rightarrow f(t), \quad \text{lorsque } n \rightarrow \infty.$$

(ii) Si $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ est continue, alors $\sigma_n(f, t) \rightarrow f(t)$ uniformément sur $\mathbb{T}^1$.

Démonstration. (i) Puisque f est intégrable, elle est bornée sur $\mathbb{T}^1$ et donc il existe M tel que

$$|f(x)| \leq M, \quad \text{pour } x \in \mathbb{T}^1. \quad (*)$$

Puisque f est continue en t , pour tout $\epsilon > 0$, il existe $\delta > 0$ tel que

$$\text{si } |t - x| < \delta, \quad \text{alors } |f(t) - f(x)| \leq \frac{\epsilon}{2}. \quad (**)$$

Notons que ce $\delta = \delta(\epsilon, t)$ dépend en général de ϵ et de t . Par le lemme précédent, il existe donc maintenant $N = N(\epsilon, t)$ tel que :

$$\text{si } n \geq N(\epsilon, t) \quad \text{et} \quad x \notin [-\delta, \delta] \quad \text{alors} \quad K_n(x) \leq \frac{\epsilon}{4M}. \quad (***)$$

Alors

$$\begin{aligned} |\sigma_n(f, t) - f(t)| &= \left| \frac{1}{2\pi} \int_{\mathbb{T}^1} f(t-x) K_n(x) dx - \underbrace{\frac{1}{2\pi} \int_{\mathbb{T}^1} K_n(x) f(t) dx}_{\text{par le lemme 4 (iii)}} \right| \\ &= \left| \frac{1}{2\pi} \int_{\mathbb{T}^1} (f(t-x) - f(t)) K_n(x) dx \right| \\ &\leq \left| \frac{1}{2\pi} \int_{x \in [-\delta, \delta]} (f(t-x) - f(t)) K_n(x) dx \right| \\ &\quad + \left| \frac{1}{2\pi} \int_{x \notin [-\delta, \delta]} (f(t-x) - f(t)) K_n(x) dx \right|. \end{aligned}$$

Utilisant (**) pour le premier terme et (*) pour le second :

$$\leq \frac{\epsilon}{2} \frac{1}{2\pi} \int_{x \in [-\delta, \delta]} |K_n(x)| dx + \frac{2M}{2\pi} \int_{x \notin [-\delta, \delta]} |K_n(x)| dx.$$

Nous avons remplacé $|f(t-x) - f(t)| \leq |f(t-x)| + |f(t)| \leq 2M$. Le lemme 4 (i) permet d'enlever les valeurs absolues de la première intégrale alors que (***) permet de réécrire :

$$\begin{aligned} &\leq \frac{\epsilon}{2} \frac{1}{2\pi} \int_{\mathbb{T}^1} K_n(x) dx + \frac{2M}{2\pi} \int_{x \notin [-\delta, \delta]} \frac{\epsilon}{4M} dx \\ &\leq \frac{\epsilon}{2} + \frac{\epsilon}{2} \\ &\leq \epsilon \end{aligned}$$

pour tout $n \geq N(\epsilon, t)$.

(ii) Si f est continue, alors f est uniformément continue sur $\mathbb{T}^1$ et le choix du δ peut être fait indépendamment du t sur $\mathbb{T}^1$. $\square$

Théorème 6. Si $f, g : \mathbb{T}^1 \rightarrow \mathbb{C}$ sont continues et tous leurs coefficients de Fourier coïncident deux à deux $c_r^f = c_r^g, \forall r \in \mathbb{Z}$, alors $f = g$.

Démonstration. En effet

$$0 = \sigma_n(f, t) - \sigma_n(g, t) \rightarrow f(t) - g(t),$$

lorsque $n \rightarrow \infty$. Ainsi $f(t) - g(t) = 0$ pour tout t et $f = g$. $\square$

Avant d'énoncer le prochain théorème, rappelons la définition de « sup ».

sup A

Définition 9. Soit $A \subset \mathbb{R}$ un sous-ensemble des réels. Le nombre réel x est une borne supérieure de A si $x \geq y, \forall y \in A$. L'ensemble A possède une plus petite borne supérieure, s'il existe une borne supérieure de A , notée $\sup A$, telle que, si b est une borne supérieure de A , alors $\sup A \leq b$.

La plus petite borne supérieure est également appelée la *borne supérieure* dans la littérature mathématique française. Le mot employé en anglais, *supremum*, est également devenu d'usage courant en français.

Théorème 7. Si $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ est continue et $\epsilon > 0$, alors il existe un polynôme trigonométrique P tel que

$$\sup_{t \in \mathbb{T}^1} |P(t) - f(t)| \leq \epsilon.$$

Démonstration. Un polynôme trigonométrique Q est, par définition, de la forme $Q(t) = \sum_{r=-n}^n a_r e^{irt}$ et σ_n est précisément de cette forme. Or, pour un $\epsilon > 0$ donné, le théorème 5 (ii) assure l'existence d'un n tel que $|\sigma_n(f, t) - f(t)| \leq \epsilon$. $\square$

Remarques : (1) Le théorème précédent dit que σ_n converge en tous les points ou que f peut être approximée en tout $t \in \mathbb{T}^1$ par σ_n . Il ne dit pas cependant que les σ_n convergent vers f .

(2) Le théorème de Fejér répond à notre seconde question. Si f est intégrable et continue en t , alors la valeur de f en t peut être restituée à partir des c_r .

Exercices

8*. Il est possible de réécrire les sommes partielles

$$s_n = \sum_{k=-n}^n c_k e^{ikx}$$

noyau de Dirichlet

sous la forme

$$s_n = \frac{1}{2\pi} \int_{\mathbb{T}^1} f(x+t) d_n(t) dt$$

où $d_n(t)$ est appelé le noyau de Dirichlet.

(a) Montrer que le noyau de Dirichlet est

$$d_n(t) = \frac{\sin(n + \frac{1}{2})t}{\sin \frac{t}{2}}.$$

(b) Au lemme 4, nous avons démontré trois propriétés importantes du noyau de Fejér. Lesquelles de ces propriétés demeurent vraies pour le noyau de Dirichlet ?

9. (a) Soit la série $\{s_n\}$ avec $s_0 = \frac{1}{2}$ et $s_n = \frac{1}{2} + \sum_{j=1}^n \cos jx$. En réécrivant $\cos jx$ sous forme de sommes d'exponentielles, montrer que cette suite n'est pas sommable quelque soit la valeur de $x \in [-\pi, \pi)$.

(b) Montrer cependant que, quelque soit $x \neq 0$ dans l'intervalle $[-\pi, \pi)$, cette suite est sommable au sens de Cesàro et tend vers 0.

(c) Montrer que

$$2 \sum_{j=1}^{\infty} \sin jx = \cot \frac{x}{2}$$

au sens de Cesàro pour $x \in [-\pi, \pi) \setminus \{0\}$.

10. Montrer que la série harmonique n'est pas sommable au sens de Cesàro.

11. Soit $\{D_n, n \in \mathbb{N}\}$ un ensemble de fonctions $D_n : [0, 1] \times [0, 1] \rightarrow \mathbb{R}$ symétriques, c'est-à-dire telles que $D_n(t, x) = D_n(x, t)$. On dit que $\{D_n\}$ est une famille de Dirac si :

(i) $D_n(t, x) \geq 0$, pour tout $t, x \in [0, 1]$.

(ii) Pour tout $\delta > 0$, $D_n(t, x) \xrightarrow[n \rightarrow \infty]{} 0$ uniformément sauf pour t, x tels que $|t - x| < \delta$.

(iii) $\int_0^1 D_n(t, x) dx = 1$.

(a) Montrer que les noyaux de Fejér forment une famille de Dirac si on pose

$$D_n(t, x) = K_n(2\pi(t - x)).$$

(b) Montrer que, si f est intégrable et $\{D_n\}$ est une famille de Dirac, alors

$$\int_0^1 f(x) D_n(t, x) dx \xrightarrow[n \rightarrow \infty]{} f(t)$$

en les points t où f est continue.

(c) Si f est continue sur tout l'intervalle $[0, 1]$, alors

$$\int_0^1 f(x) D_n(t, x) dx \rightarrow f(t)$$

uniformément.

12. Les ondelettes. Nous verrons dans les chapitres qui suivent que les fonctions trigonométriques ne sont pas les seules à permettre des développements en série. Dans le dernier quart de siècle, de nouvelles familles de fonctions ont retenu l'attention à cause de certaines de leurs propriétés. Ces familles de fonctions, appelées ondelettes (*wavelets*), jouent un rôle de plus en plus important dans l'analyse du signal et certains standards de transmission les utilisant ont été proposés pour la télévision haute définition. Un des exemples les plus simples de ces familles est la famille des ondelettes de Haar. Les éléments $f_{ij} : [0, 1] \rightarrow \mathbb{R}$ de cette famille sont indicés par deux indices : le premier i est $\in \mathbb{N} \cup \{0\}$ et le second j est $\in \{0, 1, 2, \dots, 2^{i-1} - 1\}$, sauf pour $i = 0$ auquel cas $j = 0$. Les premiers éléments de cette famille sont définis comme suit :

$$\begin{aligned} f_{00} &= 1 \\ f_{10} &= \begin{cases} 1, & 0 \leq x < \frac{1}{2} \\ -1, & \frac{1}{2} \leq x \leq 1 \end{cases} \\ f_{20} &= \begin{cases} \sqrt{2}, & 0 \leq x < \frac{1}{4} \\ -\sqrt{2}, & \frac{1}{4} \leq x < \frac{1}{2} \\ 0, & \frac{1}{2} \leq x \leq 1 \end{cases} \\ f_{21} &= \begin{cases} 0, & 0 \leq x < \frac{1}{2} \\ \sqrt{2}, & \frac{1}{2} \leq x < \frac{3}{4} \\ -\sqrt{2}, & \frac{3}{4} \leq x \leq 1 \end{cases} \end{aligned}$$

Les autres f_{ij} sont plus facilement définis en mots qu'à l'aide des relations explicites. Il y a 2^{i-1} fonctions dont le premier indice est i . Chacune est non-nulle que dans le j -ième intervalle de largeur $\frac{1}{2^{i-1}}$. (Rappelons que l'indice j balaie l'ensemble $\{0, 1, 2, \dots, 2^{i-1} - 1\}$ et que le 0-ième intervalle correspond donc à l'intervalle $[0, \frac{1}{2^{i-1}}]$.) Ainsi f_{46} est non-nulle dans l'intervalle $[\frac{3}{4}, \frac{7}{8}]$ de largeur $\frac{1}{2^{4-1}} = \frac{1}{8}$. Sur l'intervalle où f_{ij} est non-nulle, elle est égale à $2^{(i-1)/2}$ sur la première moitié et à $-2^{(i-1)/2}$ sur la seconde. Aux points de discontinuité, la fonction prend la valeur prise par les points immédiatement à sa droite (sauf pour le point $x = 1$ où f prend la valeur de f en les points immédiatement à sa gauche). Ainsi

$$f_{46} = \begin{cases} 0, & 0 \leq x < \frac{3}{4} \\ 2^{\frac{3}{2}}, & \frac{3}{4} \leq x < \frac{13}{16} \\ -2^{\frac{3}{2}}, & \frac{13}{16} \leq x < \frac{7}{8} \\ 0, & \frac{7}{8} \leq x \leq 1. \end{cases}$$

Les premières fonctions f_{ij} sont tracées à la figure 2.4. Si on suit les étapes de la construction du noyau de Fejér mais pour les fonctions f_{ij} plutôt que pour les

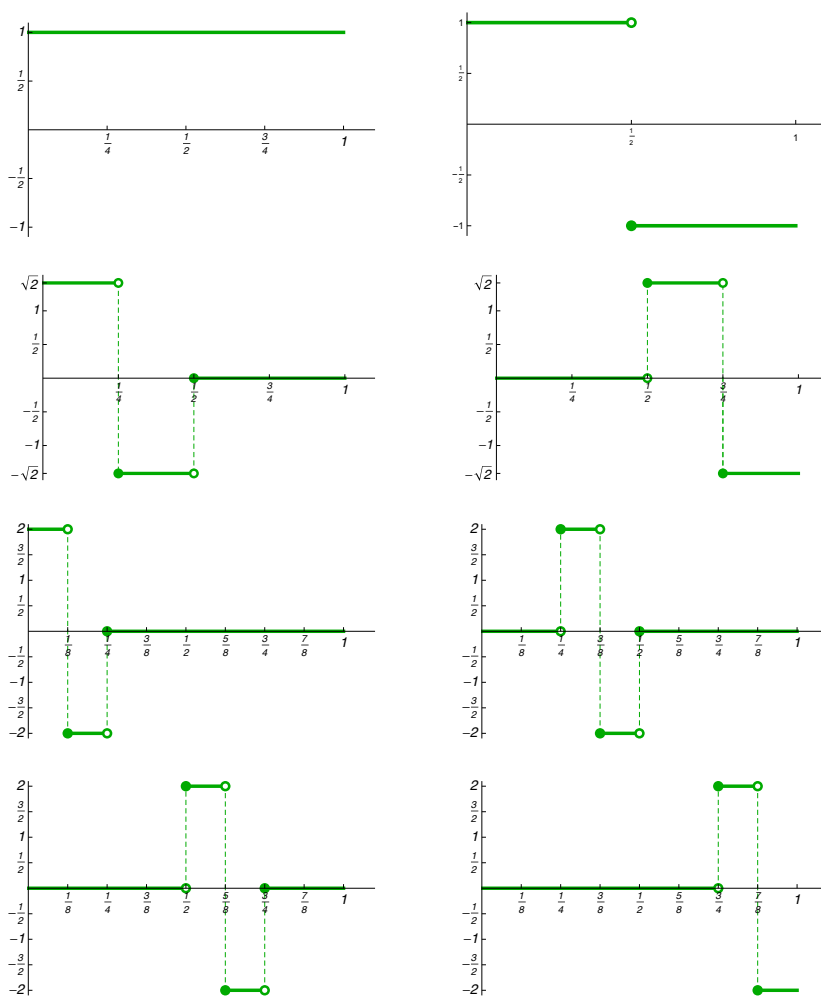


Figure 2.4 – Les premières fonctions f_{ij} : $f_{00}, f_{10}, f_{20}, f_{21}, f_{30}, f_{31}, f_{32}, f_{33}$. Les échelles verticales varient d'un graphe à l'autre.

fonctions trigonométriques, on est amené à définir la fonction

$$H_n(t, x) = 1 + \sum_{i=1}^n \frac{n+1-i}{n+1} \sum_{j=0}^{2^{i-1}-1} f_{ij}(t) f_{ij}(x), \quad n \geq 1.$$

Montrer que $\{H_n, n \in \mathbb{N}\}$ est une famille de Dirac.

Note : J'ai vérifié les propriétés de l'exercice précédent dans l'ordre (iii), (ii) et (i) et vous aurez peut-être besoin de l'identité :

$$\sum_{i=1}^{n-1} i2^i = (n-2)2^n + 2.$$

2.3 Convergence ponctuelle

Le théorème de Fejér assure la convergence au sens de Cesàro des séries de Fourier des fonctions continues. Il existe cependant des critères (plus exigeants que la simple continuité) qui impliquent la convergence, au sens usuel, des séries de Fourier. Le but de la présente section est d'en expliciter quelques-uns et donc de répondre à notre première question :

— Est-ce que les sommes partielles $s_m = a_0/2 + \sum_{n=1}^m (a_n \cos nx + b_n \sin nx)$ tendent vers $f(x)$ en tout point $x \in [-\pi, \pi]$? Sous quelles conditions ?

Nous allons nous concentrer sur les séries complexes ; les séries trigonométriques feront l'objet d'un exercice. La question est donc : quand les séries partielles

$$s_n = \sum_{r=-n}^n c_r e^{irt}$$

convergent-elles vers $f(t)$?

Rappelons tout d'abord la définition de suite de Cauchy ; un lemme préparatoire suivra.

suite de Cauchy

Définition 10. Une suite $\{a_n\}$ est dite de Cauchy si, pour tout $\epsilon > 0$, il existe N tel que

$$\text{si } m, n > N, \quad \text{alors } |a_n - a_m| < \epsilon.$$

L'utilité des suites de Cauchy repose sur le théorème suivant que nous rappelons sans preuve.

Théorème 8. Une suite $\{a_n\}$ converge si et seulement si elle est une suite de Cauchy.

Montrons maintenant le lemme suivant.

Lemme 9. Si $\sum_{r=-n}^n |a_r|$ converge lorsque $n \rightarrow \infty$, alors $\sum_{r=-n}^n a_r e^{irt}$ converge uniformément sur $\mathbb{T}^1$ lorsque $n \rightarrow \infty$ vers une fonction $g : \mathbb{T}^1 \rightarrow \mathbb{C}$ continue et dont les c_r sont précisément les $a_r, r \in \mathbb{Z}$.

Démonstration. Puisque $\sum_{r=-n}^n |a_r|$ converge, pour tout $\epsilon > 0$, il existe N tel que

$$\text{si } m \geq n \geq N \quad \text{alors} \quad \sum_{n < |r| \leq m} |a_r| < \epsilon. \quad (\text{Cauchy}).$$

Donc, si $m \geq n \geq N$

$$\left| \sum_{n < |r| \leq m} a_r e^{irt} \right| \leq \sum_{n < |r| \leq m} |a_r e^{irt}| = \sum_{n < |r| \leq m} |a_r| < \epsilon$$

et $\{s_n = \sum_{r=-n}^n a_r e^{irt}\}$ est donc aussi une suite de Cauchy qui converge uniformément pour tout $t \in \mathbb{T}^1$ car N ne dépend pas de t mais seulement de ϵ . Puisque les s_n sont continues, la limite g est aussi continue (propriété des suites uniformément convergentes, voir section 2.1).

Si $s_n = \sum_{r=-n}^n a_r e^{irt} \rightarrow g$ uniformément, alors $s_n e^{-ikt} \rightarrow g e^{-ikt}$ uniformément et si $n \geq |k|$:

$$\begin{aligned} a_k &= \sum_{r=-n}^n a_r \frac{1}{2\pi} \underbrace{\int_{\mathbb{T}^1} e^{i(r-k)t} dt}_{2\pi\delta_{r,k}} = \frac{1}{2\pi} \int_{\mathbb{T}^1} \sum_{r=-n}^n a_r e^{irt} e^{-ikt} dt \\ &\xrightarrow{n \rightarrow \infty} \frac{1}{2\pi} \int_{\mathbb{T}^1} g(t) e^{-ikt} dt = c_k^g. \end{aligned}$$

Donc $c_k^g = a_k$ pour tout $k \in \mathbb{Z}$. □

Nous sommes maintenant prêts pour notre premier théorème de convergence ponctuelle pour les suites $\{s_n\}$.

Théorème 10. Soit $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ une fonction continue et c_n ses coefficients de Fourier. Si $\sum_{r=-n}^n |c_r|$ converge, alors $s_n(f, t) \rightarrow f(t)$ uniformément sur $\mathbb{T}^1$ lorsque $n \rightarrow \infty$.

Démonstration. Puisque $\sum_{r=-n}^n |c_r|$ converge, les $s_n = \sum_{r=-n}^n c_r e^{irt}$ convergent uniformément vers une certaine fonction $g : \mathbb{T}^1 \rightarrow \mathbb{C}$ continue dont les coefficients de Fourier c_r^g sont précisément les c_r . Rien ne nous dit a priori que cette fonction g soit la fonction originale f . Cependant, par le théorème 6², nous pouvons conclure que $f = g$. □

Remarque : Le théorème 10 est

- *bon* parce que sa convergence est uniforme et le critère est simple ;
- *mauvais* parce que le critère dépend des c_r et pas seulement de f . En d'autres mots, il faut calculer les c_r pour l'utiliser.

Le critère qui suit (théorème 12) est en ce sens plus puissant.

2. Il est important de souligner que l'utilisation du théorème 6, et donc du théorème de Fejér, est capitale. Les théorèmes de la présente section reposent donc tous sur le théorème de Fejér.

Lemme 11. Soit $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ une fonction continûment différentiable $(n-1)$ fois et telle que $f^{(n-1)}$ est différentiable avec dérivée continue sauf peut-être en un nombre fini de points $x_1, x_2, \dots, x_l$. Si $|f^{(n)}(t)| \leq M$ (sauf en $t = x_1, x_2, \dots, x_l$), alors

$$|c_r| \leq \frac{M}{r^n}, \quad \forall r \in \mathbb{Z} \setminus \{0\}.$$

Démonstration. Voir l'exercice 24 à la section 1.4. $\square$

Théorème 12. Si $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ est dérivable continûment deux fois alors $s_n \rightarrow f$ uniformément.

Démonstration. Si $f^{(2)}$ est continue et donc bornée, alors, par le lemme précédent, il existe M tel que $\sum_{r=-n}^n |c_r| < M \sum_{r=-n}^n r^{-2}$ pour n assez grand. Le théorème 10 s'applique donc et $s_n \rightarrow f$ uniformément. (Nous avons utilisé ici la sommabilité de la suite $\{\frac{1}{n^2}\}$. Peut-être est-il bon de rappeler une preuve de cette propriété élémentaire. Les sommes partielles $s_n = \sum_{i=1}^n \frac{1}{i^2}$ forment une suite de Cauchy puisque, pour $m > n$

$$s_m - s_n = \sum_{i=n+1}^m \frac{1}{i^2} < \sum_{i=n+1}^m \frac{1}{(i-1)i} = \sum_{i=n+1}^m \left(\frac{1}{i-1} - \frac{1}{i} \right) = \frac{1}{n} - \frac{1}{m} < \frac{1}{n}$$

qui peut être choisi arbitrairement petit pour n suffisamment grand.) $\square$

Terminons par un résultat classique que nous énoncerons ici sans preuve. Celle-ci fera l'objet d'un exercice à la section 2.5. C'est ce résultat qu'on retrouve le plus souvent cité dans les livres élémentaires.

continuité
par morceaux

Définition 11. (a) Une fonction $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ est continue par morceaux (i) si elle est continue partout sur $\mathbb{T}^1$ sauf en un nombre fini de points $\{x_1, x_2, \dots, x_N\}$ et (ii)

$$\lim_{x \rightarrow x_i^+} f(x) \quad \text{et} \quad \lim_{x \rightarrow x_i^-} f(x)$$

existent (et sont finies) pour $1 \leq i \leq N$. On écrira alors $f(x_i^+)$ et $f(x_i^-)$ pour ces valeurs.

fonction régulière

(b) Une fonction $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ est régulière si f et sa dérivée sont continues par morceaux.

théorème de
Dirichlet
Johann Peter Dirichlet
(1805-1859)

Théorème 13 (Dirichlet (1829)). Si $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ est régulière, alors

$$s_n(f, x) \xrightarrow{n \rightarrow \infty} \frac{1}{2} \left(\lim_{t \rightarrow x^+} f(t) + \lim_{t \rightarrow x^-} f(t) \right).$$

Ce théorème affaiblit les conditions du théorème 12 et s'applique donc à une fonction avec un nombre fini de sauts telle la fonction à créneaux.

Exemple 2. La série de Fourier de la fonction $f : \mathbb{T}^1 \rightarrow \mathbb{R}$ définie par $f(x) = x$, $x \in [-\pi, \pi)$ a été obtenue en exercice à la section 1.1 :

$$x \stackrel{?}{=} -2 \sum_{n=1}^{\infty} \frac{(-1)^n}{n} \sin nx.$$

Comme nous allons le voir, le théorème 12 ne peut être appliqué à cette fonction. Celle-ci satisfait cependant les hypothèses du théorème de Dirichlet. Vérifions-le. Tout d'abord, la fonction f est continue sur l'intervalle $(-\pi, \pi)$. Elle n'est cependant pas continue sur le cercle $\mathbb{T}^1$! En effet, si on la prolonge périodiquement sur $\mathbb{R}$, son prolongement a un saut en tous les $(2k+1)\pi$, $k \in \mathbb{Z}$. Ce saut se produit une fois sur le cercle et f est donc continue partout sur $\mathbb{T}^1$ à l'exception d'un seul point. Elle vérifie donc la condition (i) de la définition ci-dessus, mais non les hypothèses du théorème 12. Sa dérivée sur l'intervalle $(-\pi, \pi)$ est simplement 1. En $-\pi$, sa dérivée ne peut pas exister puisque la fonction (prolongée) n'y est même pas continue. Donc f' est continue sur le cercle sauf en un nombre fini (= 1) de points. Au point $x_1 = -\pi$, la limite à droite est

$$\lim_{x \rightarrow x_1^+} f(x) = \lim_{x \rightarrow -\pi^+} x = -\pi.$$

Pour calculer la limite à gauche, il faut utiliser le prolongement de f et

$$\lim_{x \rightarrow x_1^-} f(x) = \lim_{x \rightarrow +\pi^-} x = +\pi.$$

Les limites à droite et à gauche existent et f est continue par morceaux. Les limites à gauche et à droite de f' au point $x = -\pi$ sont simplement 1. Donc f' est également continue par morceaux. Le théorème de Dirichlet assure donc que la série de Fourier ci-dessus est égale à $f(x) = x$ pour tout $x \in (-\pi, \pi)$ et à

$$\frac{1}{2} \left(\lim_{x \rightarrow -\pi^+} x + \lim_{x \rightarrow -\pi^-} (x + 2\pi) \right) = 0$$

lorsqu'elle est évaluée en $-\pi$. (Remarquons que cette dernière observation est évidente puisque chacun des termes de la série de Fourier s'annule en $x = -\pi$.)

13. Plusieurs fonctions ont été développées depuis le début du cours, en exemples et en exercices. Identifiez les domaines de convergence de ces séries. Quelles sont celles dont la convergence est uniforme ? (L'exemple de la section 1.2 est intéressant. Rappelons qu'on y calcule le développement d'une fonction $f : [0, \pi] \rightarrow \mathbb{R}$ en série de cosinus puis en série de sinus. Notez que les deux séries ne convergent pas aussi rapidement l'une que l'autre.)

Exercices

14*. Utilisez les différentes séries calculées pour démontrer les identités suivantes. Citez le théorème qui vous permet d'affirmer l'égalité de la série et de la fonction originale au point désiré.

(a)

$$\frac{\pi}{4} = 1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} + \dots$$

(b)

$$\frac{\pi}{2} = 1 - 2 \sum_{j=1}^{\infty} \frac{(-1)^j}{4j^2 - 1}$$

(c)

$$\frac{\pi}{\sinh \pi} = \sum_{n \in \mathbb{Z}} \frac{(-1)^n}{n^2 + 1}$$

(d)

$$\frac{\pi}{\tanh \pi} = \sum_{n \in \mathbb{Z}} \frac{1}{n^2 + 1}$$

(e) Lesquelles des sommes suivantes savez-vous faire :

$$\sum \frac{1}{n^i} \quad \text{et} \quad \sum \frac{(-1)^n}{n^i}$$

Ici i est un entier ≥ 2 .

(f)

$$\ln 2 = \sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n}$$

Attention : pour cette série, est-il possible de démontrer la convergence de la série à l'aide des théorèmes démontrés ?

15*. (Körner.) Justifiez les étapes importantes à l'aide des théorèmes appropriés.

(a) Montrer que

$$\frac{3x^2 - \pi^2}{12} = \sum_{n=1}^{\infty} \frac{(-1)^n}{n^2} \cos nx, \quad -\pi \leq x \leq \pi.$$

(b) Trouver une forme polynomiale pour la série

$$\sum_{n=1}^{\infty} \frac{(-1)^n}{n^3} \sin nx, \quad -\pi \leq x \leq \pi.$$

(c) Soit $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ définie par

$$f(x, y) = \sum_{n=1}^{\infty} \frac{(-1)^n}{n^2} \sin nx \sin ny.$$

Montrer que f est continue et que l'ensemble $\Omega = \{(x, y) \in \mathbb{R}^2 | f(x, y) = 0\}$ est constitué de droites s'intersectant à angle droit.

(d) Evaluer $f(\frac{3\pi}{4}, -\frac{5\pi}{4})$.

16. (a) Donner un critère sur les coefficients a_n et b_n des séries trigonométriques qui permette de réécrire le théorème 10 pour les séries partielles

$$S_n = \frac{a_0}{2} + \sum_{m=0}^n (a_m \cos mx + b_m \sin mx), \quad n \geq 1.$$

Montrer que le critère proposé est suffisant.

(b) Est-ce que le théorème 12 demeure vrai pour les sommes partielles S_n ci-dessus? Le montrer ou trouver un contre-exemple.

2.4 Rappel d'algèbre linéaire

Les preuves de convergence, à ce point, ont reposé sur des arguments d'analyse seulement. Pourtant, si nous nous rappelons l'analogie entre la corde plombée et la corde vibrante (voir section 1.5), nous aimerions interpréter les fonctions $\sin \frac{\pi x}{L}$ comme les modes naturels de la corde vibrante, ou encore, en poussant l'analogie plus loin, comme les vecteurs propres de l'opérateur $\frac{d^2}{dx^2}$ sur l'ensemble des fonctions vérifiant les conditions aux limites $f(0) = f(L) = 0$. Ce vocabulaire est clairement celui de l'algèbre linéaire. Cette section rappelle donc les concepts élémentaires qui seront cruciaux pour rendre rigoureuse l'analogie précédente.

Le premier concept important est celui d'espace vectoriel. Les ingrédients sont V , un ensemble dont les éléments sont appelés vecteurs, et deux opérations : l'addition « + » qui permet d'associer à toute paire ordonnée de deux éléments de V un autre élément appelé leur somme et la multiplication « · » qui associe à une paire $c \in \mathbb{F}$ et $v \in V$ un vecteur de V noté simplement cv . Dans notre cas $\mathbb{F}$, le corps de définition de V , sera l'ensemble des nombres réels $\mathbb{R}$ ou des nombres complexes $\mathbb{C}$.

Définition 12. L'ensemble V muni des deux opérations + et · est un espace vectoriel si les propriétés suivantes sont satisfaites :

- (i) + est commutative : $v_1 + v_2 = v_2 + v_1$, $v_1, v_2 \in V$,
- (ii) + est associative : $(v_1 + v_2) + v_3 = v_1 + (v_2 + v_3)$, $v_1, v_2, v_3 \in V$,
- (iii) il existe un unique élément neutre noté 0 : $0 + v = v + 0 = v$, $v \in V$,
- (iv) il existe, pour tout $v \in V$, un unique inverse noté $-v$: $v + (-v) = (-v) + v = 0$,

- (v) si $1 \in \mathbb{F}$ est l'unité dans $\mathbb{F}$, alors $1 \cdot v = v, v \in V$,
 (vi) pour tout $c_1, c_2 \in \mathbb{F}$ et $v \in V : (c_1 c_2)v = c_1(c_2 v)$,
 (vii) pour tout $c \in \mathbb{F}$ et $v_1, v_2 \in V : c(v_1 + v_2) = cv_1 + cv_2$,
 (viii) pour tout $c_1, c_2 \in \mathbb{F}$ et $v \in V : (c_1 + c_2)v = c_1 v + c_2 v$.

Cette définition est fort longue et vous avez sûrement vérifié qu'elle était satisfaite pour les espaces vectoriels les plus usuels. La plupart du temps, plusieurs des propriétés peuvent être vérifiées mentalement.

application linéaire

isomorphisme

Les définitions suivantes complètent la définition d'espace vectoriel. On dit qu'une fonction (ou application) $f : V \rightarrow W$ entre deux espaces vectoriels est linéaire si $f(cv + w) = cf(v) + f(w)$ pour tout $c \in \mathbb{F}$ et $v, w \in V$. Un isomorphisme f entre deux espaces vectoriels V et W est une application linéaire $f : V \rightarrow W$ qui possède un inverse $f^{-1} : W \rightarrow V$.

Un sous-ensemble $X = \{v_1, v_2, \dots, v_N\}$ d'un espace vectoriel V est linéairement indépendant si une combinaison linéaire d'éléments de V

$$a_1 v_1 + a_2 v_2 + \dots + a_N v_N$$

qui est nulle ne peut avoir que des coefficients nuls :

$$a_1 = a_2 = \dots = a_N = 0.$$

base

Une base d'un espace vectoriel est un sous-ensemble linéairement indépendant $X = \{e_1, e_2, \dots, e_n\}$ de V qui engendre V , c'est-à-dire tel que, pour tout $v \in V$, il existe n coefficients $a_1, a_2, \dots, a_n \in \mathbb{F}$ tels que

$$v = a_1 e_1 + a_2 e_2 + \dots + a_n e_n.$$

dimension

Le nombre d'éléments dans une base d'un espace vectoriel est appelé la dimension de V .

Voici quelques exemples.

Exemple 1. L'ensemble $\mathbb{R}^n = \{(x_1, x_2, \dots, x_n) | x_1, x_2, \dots, x_n \in \mathbb{R}\}$ des n -tuplets forme un espace vectoriel si on le munit des deux opérations suivantes : l'addition des deux éléments $(x_1, x_2, \dots, x_n)$ et $(y_1, y_2, \dots, y_n)$ est $(x_1 + y_1, x_2 + y_2, \dots, x_n + y_n)$ et la multiplication par un réel $r \in \mathbb{F} = \mathbb{R}$ est donnée par $(rx_1, rx_2, \dots, rx_n)$. Il est aisé de vérifier les propriétés ci-dessus. Par exemple, le neutre est le n -tuplet $(0, 0, \dots, 0)$ et l'inverse de $(x_1, x_2, \dots, x_n)$ est $(-x_1, -x_2, \dots, -x_n)$. La propriété (vii) découle de la distributivité dans les réels :

$$\begin{aligned} r((x_1, x_2, \dots, x_n) + (y_1, y_2, \dots, y_n)) &= r(x_1 + y_1, x_2 + y_2, \dots, x_n + y_n) \\ &= (rx_1 + ry_1, rx_2 + ry_2, \dots, rx_n + ry_n) \\ &= r(x_1, x_2, \dots, x_n) + r(y_1, y_2, \dots, y_n) \end{aligned}$$

$\mathbb{R}^n$ est un espace vectoriel de dimension n .

Voici maintenant un exemple plus proche de l'usage que nous voulons faire des concepts d'algèbre linéaire.

Exemple 2. Soit P_n l'ensemble des polynômes de degré n (et inférieur). L'addition de deux polynômes et la multiplication par un réel suivent les définitions usuelles : la somme de deux fonctions évaluée en un point est la somme de l'évaluation des deux fonctions en ce point, etc. Ainsi la somme des deux polynômes $p(x) = \sum_{i=0}^n p_i x^i$ et $q(x) = \sum_{i=0}^n q_i x^i$ est le polynôme $(p + q)(x) = \sum_{i=0}^n (p_i + q_i) x^i$. A nouveau, la vérification des propriétés d'espace vectoriel est directe. Le neutre 0 est le polynôme nul, c'est-à-dire le polynôme qui associe à tout x la valeur zéro. (Il serait peut-être utile de le noter $0(x)$.)

P_n est un espace vectoriel de dimension $(n + 1)$ puisqu'à l'aide d'une combinaison linéaire des $(n + 1)$ polynômes $e_0 = 1$, $e_1 = x$, $e_2 = x^2, \dots, e_n = x^n$, il est possible d'écrire tous les éléments de P_n . De plus, ces $(n + 1)$ polynômes sont linéairement indépendants. (Comment s'en convaincre ?) Si c'est la première fois que vous voyez cet exemple, vous aurez peut-être la tendance à confondre le rôle de l'indice dans l'exemple 1. et celui de la variable x dans l'exemple 2.. Ces deux rôles sont différents. Le fait que x puisse prendre plusieurs valeurs tel un indice ne doit pas vous induire en erreur. Il est facile de comprendre à l'aide de la base donnée ci-dessus que c'est l'*exposant* de la variable x qui joue le rôle de l'étiquette et non la variable x elle-même.

Exemple 3. Voici maintenant un exemple qui sort des cours élémentaires d'algèbre linéaire. Soit $\mathcal{C}^0(\mathbb{R})$ l'ensemble des fonctions continues définies sur l'axe réel. Cet ensemble peut être muni des deux opérations d'addition $+$ (addition usuelle de fonctions) et de multiplication $\cdot$ (multiplication par une constante). Rappelons que la somme de deux fonctions continues est une fonction continue et que, si r est un nombre réel, alors la fonction rf est continue si f l'est. Il est aisé de montrer que les propriétés d'espace vectoriel sont vérifiées. Par exemple, l'associativité suit de :

$$\begin{aligned} (f + (g + h))(x) &= f(x) + (g + h)(x) \\ &= f(x) + g(x) + h(x) \\ &= (f + g)(x) + h(x) \\ &= ((f + g) + h)(x). \end{aligned}$$

Ainsi $\mathcal{C}^0(\mathbb{R})$ est un espace vectoriel. Ses éléments (ses vecteurs) sont des fonctions.

Une façon naturelle d'introduire de nouveaux espaces vectoriels est par l'introduction de contraintes sur les éléments d'un ensemble déjà identifié comme étant un espace vectoriel. Il n'est pas alors nécessaire de vérifier toutes les propriétés d'espace vectoriel.

Théorème 14. *Un sous-ensemble W d'un espace vectoriel V est un espace vectoriel si l'addition de vecteurs et la multiplication par un scalaire sont fermées pour l'ensemble W . On dit de W qu'il est un sous-espace vectoriel de V .*

Exemple 4. Soit $W \subset \mathbb{R}^n$ l'ensemble des n -uplets de $\mathbb{R}^n$ dont la somme des composantes est nulle :

$$W = \{(x_1, x_2, \dots, x_n) \in \mathbb{R}^n \mid x_1 + x_2 + \dots + x_n = 0\}.$$

W est un sous-espace vectoriel de $\mathbb{R}^n$. En effet, si $x = (x_1, x_2, \dots, x_n), y = (y_1, y_2, \dots, y_n) \in W$ et $r \in \mathbb{R}$, alors :

$$x + y = (x_1, x_2, \dots, x_n) + (y_1, y_2, \dots, y_n) = (x_1 + y_1, x_2 + y_2, \dots, x_n + y_n)$$

avec

$$\sum_{i=1}^n (x_i + y_i) = \sum_{i=1}^n x_i + \sum_{i=1}^n y_i = 0$$

et

$$rx = (rx_1, rx_2, \dots, rx_n)$$

avec

$$\sum_{i=1}^n (rx_i) = r \sum_{i=1}^n x_i = 0.$$

Donc $x + y$ et rx sont dans le sous-ensemble W et W est un sous-espace vectoriel de V .

Exemple 5. Soit $Q_n \subset P_n$ le sous-ensemble des polynômes de degré n qui s'annulent en $x = 0$:

$$Q_n = \{p \in P_n \mid p(0) = 0\}.$$

Q_n est un sous-espace vectoriel. Si $p, q \in Q_n$ et $r \in \mathbb{R}$, alors :

$$\begin{array}{lll} (p+q)(x) = p(x) + q(x) & \text{et} & (p+q)(0) = p(0) + q(0) = 0 \\ (rp)(x) = rp(x) & \text{et} & (rp)(0) = rp(0) = 0 \end{array}$$

et $p + q$ et rp sont dans Q_n .

Exemple 6. Soit $\mathcal{C}^1(\mathbb{R}) \subset \mathcal{C}^0(\mathbb{R})$ le sous-ensemble des fonctions dérivables continûment sur les réels. $\mathcal{C}^1(\mathbb{R})$ est un sous-espace vectoriel de $\mathcal{C}^0(\mathbb{R})$. La vérification repose sur les propriétés des fonctions dont la dérivée est continue. Si $f, g \in \mathcal{C}^1(\mathbb{R})$, alors la dérivée de la fonction $f + g$ est la somme des dérivées $f' + g'$. Puisque f' et g' sont continues, $f' + g'$ est continue et donc $(f + g) \in \mathcal{C}^1(\mathbb{R})$. De même, si $r \in \mathbb{R}$, alors $(rf)' = rf'$ et puisque f' est continue, $(rf)'$ l'est et $rf \in \mathcal{C}^1(\mathbb{R})$.

Exemple 7. Avant de quitter ce premier concept, il est utile de donner un exemple d'un sous-ensemble W d'un espace vectoriel V qui *n'est pas* un sous-espace vectoriel. Soit W le sous-ensemble des vecteurs de $\mathbb{R}^n$ dont la première composante est 1 :

$$W = \{(x_1, x_2, \dots, x_n) \in \mathbb{R}^n \mid x_1 = 1\}.$$

Ce sous-ensemble W n'est pas un sous-espace vectoriel. En effet, la somme de $x = (1, x_2, x_3, \dots, x_n)$ et $y = (1, y_2, y_3, \dots, y_n)$, deux éléments de W , n'est pas un élément de W puisque la première composante de $x + y$ est 2. L'addition n'est donc pas fermée à l'intérieur de W et W n'est pas un espace vectoriel.

Le second concept fondamental de l'algèbre linéaire dont nous aurons besoin par la suite est celui de produit scalaire ou produit intérieur. (L'appellation « produit intérieur » est calquée sur l'appellation anglaise « inner product ».)

Définition 13. Soit V un espace vectoriel et $(\cdot, \cdot) : V \times V \rightarrow \mathbb{F}$ une fonction produit scalaire qui associe à une paire de vecteurs de V un élément du corps $\mathbb{F}$. La fonction $(\cdot, \cdot)$ est un produit scalaire si elle satisfait les propriétés suivantes pour tout vecteur $u, v, w \in V$ et $r \in \mathbb{F}$:

- (i) $(u + v, w) = (u, w) + (v, w)$,
- (ii) $(cv, w) = c(v, w)$,
- (iii) $(v, w) = (w, v)^*$,
- (iv) $(v, v) \geq 0$, avec l'égalité seulement pour $v = 0$.

L'astérisque dans la propriété (iii) dénote la conjugaison complexe. Il suit de (i) et de (iii) que, si $\mathbb{F} = \mathbb{R}$, alors le produit scalaire est linéaire sur ses deux composantes. Si $\mathbb{F} = \mathbb{C}$, alors $(\cdot, \cdot)$ est linéaire sur sa première composante et anti-linéaire sur sa seconde : $(v, cw) = c^*(v, w)$.

Exemple 8. La fonction suivante $(\cdot, \cdot)$ est un produit scalaire pour $\mathbb{R}^n$. Si $x = (x_1, x_2, \dots, x_n)$ et $y = (y_1, y_2, \dots, y_n)$ sont des éléments de $\mathbb{R}^n$, alors

$$\begin{aligned}(x, y) &= \sum_{i=1}^n x_i y_i \\ &= y^T x\end{aligned}$$

où T dénote la transposition matricielle. (Il est usuel de définir $(x, y) = x^T y$. Le choix inhabituel $y^T x$ est celui qui se généralisera naturellement aux espaces vectoriels $\mathbb{C}^n$. Voir ci-dessous.) La seule propriété du produit intérieur qui n'est pas immédiate est la quatrième. Mais puisque

$$(x, x) = \sum_{i=1}^n x_i^2,$$

(x, x) est un nombre non-négatif qui sera nul que si tous les carrés x_i^2 sont nuls, c'est-à-dire si toutes les composantes du vecteur x sont nulles et donc si $x = 0$. Pour obtenir un produit scalaire sur $\mathbb{C}^n$, il faut remplacer la définition précédente par

$$\begin{aligned}(x, y) &= \sum_{i=1}^n x_i y_i^* \\ &= y^\dagger x\end{aligned}$$

où † dénote la conjugaison complexe suivie de la transposition matricielle. Ce choix assure que la linéarité (ii) sur le premier vecteur x est satisfaite.

Nous pourrions donner plusieurs autres exemples de produit scalaire. En fait, la première étape à la prochaine section sera justement de définir un produit scalaire sur l'ensemble des fonctions continues sur le cercle $\mathbb{T}^1$. Cet exemple est non-trivial et nous reporterons donc cette discussion à cette section.

Nous terminerons cette section en rappelant une preuve de l'inégalité de Cauchy-Schwarz-Buniakowski aussi appelée l'inégalité du triangle.

inégalité du triangle

Lemme 15 (Cauchy, Schwarz, Buniakowski). *Soit V un espace vectoriel de corps $\mathbb{F}$ muni d'un produit scalaire $(\cdot, \cdot)$. Si $v, w \in V$, alors $|(v, w)|^2 \leq (v, v)(w, w)$ avec l'égalité si et seulement si $cv + dw = 0$ pour certains $c, d \in \mathbb{F}$ dont au moins un est non-nul.*

Démonstration. Si $v = w = 0$, il n'y a rien à montrer. Supposons donc que $v \neq 0$ et donc que $(v, v) \neq 0$. Alors, quelque soient c et $d \in \mathbb{F}$:

$$\begin{aligned} 0 &\leq (cv + dw, cv + dw) \\ &= c(v, cv + dw) + d(w, cv + dw) \\ &= cc^*(v, v) + cd^*(v, w) + dc^*(w, v) + dd^*(w, w) \\ &= \left(c(v, v)^{\frac{1}{2}} + d \frac{(w, v)^{\frac{1}{2}}}{(v, v)^{\frac{1}{2}}} \right) \left(c(v, v)^{\frac{1}{2}} + d \frac{(w, v)^{\frac{1}{2}}}{(v, v)^{\frac{1}{2}}} \right)^* \\ &\quad + dd^* \left((w, w) - \frac{|(w, v)|^2}{(v, v)} \right) \\ &= \left| c(v, v)^{\frac{1}{2}} + d \frac{(w, v)^{\frac{1}{2}}}{(v, v)^{\frac{1}{2}}} \right|^2 + |d|^2 \left((w, w) - \frac{|(w, v)|^2}{(v, v)} \right). \end{aligned}$$

En particulier, pour $d = 1$ et $c = -\frac{(w, v)}{(v, v)}$, on a

$$0 \leq |0|^2 + \underbrace{|d|^2}_{=1} \left((w, w) - \frac{|(w, v)|^2}{(v, v)} \right)$$

c'est-à-dire

$$|(w, v)|^2 \leq (v, v)(w, w)$$

avec l'égalité seulement si $0 = (cv + dw, cv + dw)$, et donc, seulement si $cv + dw = 0$ avec, par exemple à nouveau $d = 1$ et $c = -\frac{(w, v)}{(v, v)}$.

Réciproquement, si $cv + dw = 0$ pour certains $c \neq 0$ et $d \in \mathbb{F}$, alors $v = -\frac{d}{c}w$ et

$$|(v, w)|^2 = \left| \frac{d}{c} \right|^2 (w, w)^2 = \left(-\frac{d}{c} \right) \left(-\frac{d}{c} \right)^* (w, w)(w, w) = (v, v)(w, w).$$

□

17. Dire lesquels des ensembles suivants sont des espaces vectoriels.

Exercices

(a) $V = \mathbb{R}^{n \times n}$ est l'ensemble des matrices réelles $n \times n$, $+$ l'addition matricielle et $\cdot$ la multiplication de tous les éléments de matrice par une constante.

(b) V est l'ensemble des fonctions sur $\mathbb{R}^2$, $+$ est l'addition usuelle de fonctions et $\cdot$ la multiplication par une constante.

(c) V est l'ensemble des continues sur $\mathbb{R}$, $\cdot$ est le produit par une constante et $+$ est la composition de fonctions : $(f + g)(x) = (f \circ g)(x)$.

(d) F_N est l'ensemble des combinaisons linéaires (avec coefficients complexes) des fonctions $e_j(x) = e^{ijx}$, $-N \leq j \leq N$, $+$ est l'addition de fonctions et $\cdot$ est la multiplication par une constante.

(e) $V = \{(s_1, s_2, s_3, \dots) \mid \sum_{i=1}^{\infty} s_i \text{ converge}\}$ est l'ensemble des séries qui convergent, $+$ est l'addition des suites terme à terme et $\cdot$ est le produit de tous les termes de la série par une constante.

18. Dire lesquels de ces sous-ensembles sont des sous-espaces vectoriels.

(a) $Q = \{p \in P_n \mid p(-1) = 0\}$.

(b) $Q = \{p \in P_n \mid \int_{-1}^1 p(x) dx = 0\}$.

(c) $W_a = \{(x_1, x_2, \dots, x_n) \in \mathbb{R}^n \mid x^T a = 0\}$ où $a \in \mathbb{R}^n$ est un n -tuplet fixé.

(d) $A = \{M \in \mathbb{R}^{n \times n} \mid M^T = -M\}$.

(e) $W = \{(x_1, x_2, \dots, x_n) \in \mathbb{R}^n \mid x_n \geq 0\}$.

(f) $\mathcal{C}^0(\mathbb{T}^1) = \{f \in \mathcal{C}^0(\mathbb{R}) \mid f(x + 2\pi) = f(x)\}$.

(g) $S = \{f \in \mathcal{C}^0(\mathbb{R}) \mid f(-x) = f(x)\}$.

(h) $\ker A = \{x \in \mathbb{R}^n \mid Ax = 0\}$ où A est une matrice réelle $n \times n$ fixée.

(i) $W = \{f \in \mathcal{C}^1(\mathbb{R}) \mid f'(0) = 0\}$.

19*. Soit $S \subset \mathbb{R}^{n \times n}$, l'espace vectoriel des matrices symétriques réelles. Montrer que $(\cdot, \cdot) : S \times S \rightarrow \mathbb{R}$ défini pour $A, B \in S$ par

$$(A, B) = \text{Tr } A^T B$$

est un produit scalaire.

20. Soit V un espace vectoriel muni d'un produit scalaire $(\cdot, \cdot)$. Soit $\{e_1, e_2, \dots, e_n\}$ une base de V qui est orthogonale par rapport à $(\cdot, \cdot)$, c'est-à-dire

orthogonalité
d'une base

$$(e_i, e_j) = d_i \delta_{ij}$$

où les d_i , $i = 1, \dots, n$, sont n nombres positifs. Montrer que si $v = \sum_{i=1}^n v_i e_i$, alors $v_i = (v, e_i)/d_i$. Si tous les d_i sont égaux à 1, la base est dite orthonormale.

orthonormalité
d'une base

21*. Soit $\mathcal{P} = \mathcal{P}([-1, 1])$ l'ensemble des polynômes à coefficients réels sur l'intervalle $[-1, 1]$.

(a) Montrer que $\mathcal{P}$ est un espace vectoriel. Quel est le polynôme qui joue le rôle de l'origine ?

(b) Montrer que $(\cdot, \cdot)$ défini par

$$(p, q) = \int_{-1}^1 p(x)q(x) dx$$

est un produit scalaire sur $\mathcal{P}$.

(c) Montrer que les polynômes $p_0 = 1, p_1 = x, p_2 = x^2, p_3 = x^3$ sont linéairement indépendants dans $\mathcal{P}$.

(d) Utiliser le processus de Gram-Schmidt pour trouver quatre polynômes P_0, P_1, P_2 et P_3 orthogonaux engendrant le sous-espace $\{p \in \mathcal{P} \mid p(x) = a + bx + cx^2 + cx^3\}$.

Note : Les premiers polynômes que vous venez de calculer sont, à une constante près, appelés polynômes de Legendre. Nous les étudierons de plus près à la section 3.4.

22*. Soit F_N l'espace vectoriel défini à l'exercice 1.(d).

(a) Soit $(\cdot, \cdot) : F_N \times F_N \rightarrow \mathbb{C}$ la fonction définie par

$$(f, g) = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x)g(x)^* dx, \quad f, g \in F_N.$$

Montrer que $(\cdot, \cdot)$ est un produit scalaire pour F_N .

(b) Montrer que $\{e_j, -N \leq j \leq N\}$ est un ensemble de vecteurs orthonormaux.

(c) Montrer que $\{e_j, -N \leq j \leq N\}$ est une base de F_N .

(d) Quelle est la dimension de F_N ?

(e) En conclure, par l'exercice précédent, que si $f \in F_N$, alors elle peut être écrite sous la forme $f = \sum_{j=-N}^N c_j e_j$ où $c_j = (f, e_j)$.

double série
de Fourier

23*. Le développement en série de Fourier de fonctions à deux variables n'est guère plus compliqué que celui en une variable. Cet exercice essaiera de vous en convaincre. (Cet exercice n'a aucune prétention à la rigueur. Il permet d'indiquer que les concepts de l'algèbre linéaire que nous venons de rappeler suggèrent souvent des généralisations aisées aux idées introduites au premier chapitre.)

(a) Soit $a_{m,n}$, $m, n \geq 0$, $b_{m,n}$, $c_{n,m}$, $m \geq 0, n \geq 1$ et $d_{m,n}$, $m, n \geq 1$, les fonctions trigonométriques en deux variables suivantes :

$$a_{m,n}(x, y) = \cos mx \cos ny,$$

$$b_{m,n}(x, y) = \cos mx \sin ny,$$

$$c_{m,n}(x, y) = \sin mx \cos ny,$$

$$d_{m,n}(x, y) = \sin mx \sin ny.$$

Soit T_N l'ensemble des combinaisons linéaires des $a_{m,n}$, $b_{m,n}$, $c_{m,n}$ et $d_{m,n}$ avec $m, n \leq N$. Montrer que T_N est un espace vectoriel.

(b) Soit $(\cdot, \cdot) : T_N \times T_N \rightarrow \mathbb{R}$ la règle suivante

$$(f, g) = \frac{1}{\pi^2} \int_{-\pi}^{\pi} \int_{-\pi}^{\pi} f(x, y)g(x, y) dx dy, \quad f, g \in T_N.$$

Montrer que les fonctions $a_{m,n}$, $b_{m,n}$, $c_{m,n}$ et $d_{m,n}$ sont orthogonales, c'est-à-dire, si x et y tiennent lieu d'un des symboles a, b, c ou d , alors

$$(x_{k,l}, y_{m,n}) = \epsilon_{x,k,l} \delta_{xy} \delta_{km} \delta_{ln}.$$

Déterminer les $\epsilon_{x,k,l}$.

(c) Montrer que l'ensemble $a_{m,n}$, $b_{m,n}$, $c_{m,n}$ et $d_{m,n}$ avec $m, n \leq N$ forme une base de T_N .

(d) Soit $f(x, y)$ une fonction périodique et de période 2π en les deux variables x et y : $f(x + 2\pi, y) = f(x, y + 2\pi) = f(x, y)$. En supposant qu'il soit possible d'écrire :

$$f(x, y) = \sum_{m,n} (A_{m,n} \cos mx \cos ny + B_{m,n} \cos mx \sin ny + C_{m,n} \sin mx \cos ny + D_{m,n} \sin mx \sin ny),$$

dites comment on peut calculer les $A_{m,n}$, $B_{m,n}$, $C_{m,n}$ et $D_{m,n}$.

(e) Développer $f(x, y) = x + y$ en double série de Fourier.

(f) A l'aide d'un langage de manipulations symboliques, comparer les graphes de f et des premières sommes partielles de la série obtenue en (e). (Même si ces programmes de manipulations symboliques produisent de remarquables graphiques en trois dimensions, pour faire une telle comparaison il est utile de se restreindre à des valeurs particulières de x ou de y .)

Note : Les idées de l'exercice ci-dessus reviendront par la suite. Faites-le !

24*. Une membrane est fixée au cadre d'un tambour carré de côtés de longueur π . Les petites oscillations transverses $z = z(x, y, t)$ sont décrites par l'équation

$$a^2 \left(\frac{\partial^2 z}{\partial x^2} + \frac{\partial^2 z}{\partial y^2} \right) = \frac{\partial^2 z}{\partial t^2}$$

où a est une constante.

(a) Obtenir les solutions z de la forme $z(x, y, t) = X(x)Y(y)T(t)$.

(b) Quel est le spectre du tambour carré ?

(c) On déforme la membrane en $t = 0$ selon la forme $z(x, y, 0) = f(x, y)$ et on la relâche sans vitesse. Décrire son évolution à l'aide des solutions obtenues en (a). Note : Pour faire ce problème, vous aurez besoin des résultats de l'exercice précédent.

(d) Supposons que la membrane carrée soit maintenue fixe le long d'une de ses diagonales : $z(x, x, t) = 0$. Quelles sont les combinaisons linéaires des solutions de (a) qui peuvent être utilisées pour résoudre ce problème ? Quel est le spectre dans ce cas ? (Considérer les combinaisons de solutions de même fréquence.)

2.5 Convergence en moyenne

Dans tout ce qui précède, nous nous sommes intéressés à la convergence ponctuelle. Même si celle-ci a son importance, elle n'est sûrement pas la seule propriété importante. Le théorème 13 ci-dessus indique que la fonction créneau que nous avons étudiée à plusieurs reprises possède une série qui converge partout sauf aux sauts où la série tend vers 0. Donc la valeur de la série en certains points peut différer de la valeur originale.

Le but de cette section est d'introduire une autre mesure de la convergence qui nous permettra de s'assurer que ces écarts sont négligeables. La première étape dans la mesure de cet écart est la définition d'une « distance » entre fonctions. Définissons tout d'abord un produit scalaire.

Définition 14. Soit $\mathcal{C}(\mathbb{T}^1)$ l'espace vectoriel des fonctions continues sur le cercle et à valeurs complexes. Soit $(\cdot, \cdot) : \mathcal{C}(\mathbb{T}^1) \times \mathcal{C}(\mathbb{T}^1) \rightarrow \mathbb{C}$ la fonction

$$(f, g) = \frac{1}{2\pi} \int_{\mathbb{T}^1} f(t)g(t)^* dt, \quad f, g \in \mathcal{C}(\mathbb{T}^1).$$

Lemme 16. La fonction $(\cdot, \cdot) : \mathcal{C}(\mathbb{T}^1) \times \mathcal{C}(\mathbb{T}^1) \rightarrow \mathbb{C}$ introduite dans la définition précédente est un produit scalaire. En d'autres mots, si $f, g, h \in \mathcal{C}(\mathbb{T}^1)$ et $\lambda, \mu \in \mathbb{C}$, alors

(i) et (ii) $(\lambda f + \mu g, h) = \lambda(f, h) + \mu(g, h)$

(iii) $(f, g) = (g, f)^*$

(iv) $(f, f) \geq 0$ et $(f, f) = 0$ si et seulement si $f = 0$.

Démonstration. Calculs directs pour (i), (ii) et (iii). Toutes les intégrations sont sur le cercle.

(i) et (ii)

$$\begin{aligned}
(\lambda f + \mu g, h) &= \frac{1}{2\pi} \int (\lambda f + \mu g) h^* dt \\
&= \frac{1}{2\pi} \left(\lambda \int f h^* dt + \mu \int g h^* dt \right) \\
&= \lambda(f, h) + \mu(g, h)
\end{aligned}$$

(iii)

$$\begin{aligned}
(g, f)^* &= \left(\frac{1}{2\pi} \int g(t) f(t)^* dt \right)^* \\
&= \frac{1}{2\pi} \int (g(t) f(t)^*)^* dt \\
&= \frac{1}{2\pi} \int f(t) g(t)^* dt = (f, g).
\end{aligned}$$

Pour (iv), notons tout d'abord que

$$(f, f) = \frac{1}{2\pi} \int f f^* dt = \frac{1}{2\pi} \int \underbrace{|f|^2}_{\geq 0 \text{ et réel}} dt \geq 0$$

Maintenant, si g est une fonction continue positive sur un intervalle $[a, b]$, alors $\int_a^b g(t) dt > 0$. Alors, si on prend $g(t) = |f(t)|^2 \geq 0$, on constate que, si $f \neq 0$, alors $(f, f) \neq 0$. $\square$

Puisque $(\cdot, \cdot)$ est un produit scalaire, le lemme 15 de Cauchy, Schwarz et Buniakowski donne le corollaire immédiat :

Corollaire 17. Si $f, g \in \mathcal{C}(\mathbb{T}^1)$, alors $|(f, g)|^2 \leq (f, f)(g, g)$ avec l'égalité si et seulement si $\lambda f + \mu g = 0$ pour certains $\lambda, \mu \in \mathbb{C}$ dont au moins un est non-nul.

Définition 15. Si $f \in \mathcal{C}(\mathbb{T}^1)$, la norme de la fonction f , notée $\|f\|_2$, est définie norme d'une fonction par

$$\|f\|_2 = (f, f)^{\frac{1}{2}}.$$

Lemme 18. Si $f, g \in \mathcal{C}(\mathbb{T}^1)$ et $\lambda \in \mathbb{C}$, alors

- (i) $\|\lambda f\|_2 = |\lambda| \|f\|_2$
- (ii) $\|f\|_2 \geq 0$ avec l'égalité si et seulement si $f = 0$
- (iii) (inégalité du triangle) $\|f\|_2 + \|g\|_2 \geq \|f + g\|_2$

Démonstration. (i) Calculer $\|\lambda f\|_2^2$.

(ii) Lemme 16 (iv).

(iii) Calcul direct :

$$\begin{aligned}
 \|f + g\|_2^2 &= (f + g, f + g) = (f, f) + (g, f) + (f, g) + (g, g) \\
 &= \|f\|_2^2 + 2 \operatorname{Re}(f, g) + \|g\|_2^2 \\
 &\leq \|f\|_2^2 + 2|(f, g)| + \|g\|_2^2 \\
 &\leq \|f\|_2^2 + 2\|f\|_2\|g\|_2 + \|g\|_2^2, \quad \text{par le corollaire 17} \\
 &= (\|f\|_2 + \|g\|_2)^2.
 \end{aligned}$$

□

Pour raccourcir les notations, nous écrivons

$$e_n(t) = e^{int}, \quad t \in \mathbb{T}^1.$$

Ceci nous permettra aussi de référer à la fonction e_n sans argument comme on le fait souvent pour la fonction f , comme par exemple dans l'expression (f, g) .

Lemme 19. $(e_n, e_m) = \delta_{n,m}$

Démonstration. Calcul direct. □

Ce lemme a comme conséquence immédiate que les coefficients de Fourier c_j de $f \in \mathcal{C}(\mathbb{T}^1)$

$$f(t) = \sum_{j \in \mathbb{Z}} c_j e_j(t)$$

sont simplement

$$c_j = (f, e_j).$$

interprétation
géométrique

Avec ce produit scalaire entre fonctions continues et cette norme, l'analogie avec l'algèbre linéaire devient attrayante. L'espace des fonctions continues sur le cercle forme un espace vectoriel sur lequel est définie une longueur, notée $\|\cdot\|_2$. Cette longueur permet à son tour d'introduire une distance entre fonctions (voir la définition 16 ci-dessous). Il est donc possible de reformuler la troisième des questions que nous avons soulevées à la première section sous la forme : est-ce que la distance entre f et les sommes partielles $\sum_{r=-n}^n c_r e_r$ de sa série de Fourier approche 0 lorsque $n \rightarrow \infty$? Cette reformulation en termes géométriques peut même être raffinée en les deux questions suivantes.

Question A : Quelles valeurs de $\lambda_{-n}, \dots, \lambda_n \in \mathbb{C}$ (s'il y en a) minimisent la longueur

$$\left\| f - \sum_{r=-n}^n \lambda_r e_r \right\|_2 ?$$

Question B : (Même question en d'autres mots...) Soit $E_n = \{\sum_{j=-n}^n \lambda_j e_j \mid \lambda_j \in \mathbb{C}\}$ l'espace vectoriel complexe engendré par les $\{e_j, -n \leq j \leq n\}$. Existe-t-il un $g \in E_n$ qui minimise la longueur $\|f - g\|_2$?

Un premier élément de réponse peut être trouvé dans l'intuition géométrique que nous avons des espaces vectoriels de dimension finie.

Intuition : Il existe un unique point $g_0 \in E_n$ tel que $\|f - g\|_2 > \|f - g_0\|_2$ pour tout $g \in E_n$, avec $g \neq g_0$. Ce g_0 est déterminé par la condition que $f - g_0$ soit perpendiculaire à E_n .

Ceci correspond à la situation usuelle pour les espaces de dimension finie. La « meilleure » approximation d'un vecteur $\vec{v} \in \mathbb{R}^3$ qui se trouve dans un plan $\mathcal{P}$ est la projection $P\vec{v}$ de $\vec{v}$ sur ce plan. Alors $\vec{v} - P\vec{v}$ est un vecteur perpendiculaire au plan $\mathcal{P}$. Remarquablement, cette intuition développée pour les espaces vectoriels de dimension finie s'avère juste dans le cas présent.

Théorème 20. Si $g_0 = \sum_{j=-n}^n (f, e_j) e_j$ et $g = \sum_{j=-n}^n \lambda_j e_j$ alors

$$\|f\|_2^2 \geq \|g_0\|_2^2 = \sum_{j=-n}^n |(f, e_j)|^2$$

et

$$\|f - g\|_2^2 \geq \|f - g_0\|_2^2 = \sqrt{\|f\|_2^2 - \|g_0\|_2^2}$$

avec l'égalité si et seulement si $\lambda_j = (f, e_j), \forall j, |j| \leq n$.

Démonstration. Transformons $\|f - g\|_2^2$ (toutes les sommes sans bornes sont sur j et vont de $-n$ à n) :

$$\begin{aligned} \|f - g\|_2^2 &= \left(f - \sum \lambda_j e_j, f - \sum \lambda_j e_j \right) \\ &= (f, f) - \sum \lambda_j (e_j, f) - \sum \lambda_j^* (f, e_j) + \sum_{j,k=-n}^n \lambda_j \lambda_k^* \underbrace{(e_j, e_k)}_{=\delta_{j,k}} \\ &= (f, f) - \left(\sum \lambda_j (f, e_j)^* + \sum \lambda_j^* (f, e_j) \right) + \sum \lambda_j \lambda_j^* \\ &= (f, f) - \sum (f, e_j)(f, e_j)^* + \sum (\lambda_j - (f, e_j))(\lambda_j - (f, e_j))^* \\ &= (f, f) - \sum |(f, e_j)|^2 + \sum |\lambda_j - (f, e_j)|^2 \\ &\geq (f, f) - \sum |(f, e_j)|^2 = \|f\|_2^2 - \sum |(f, e_j)|^2 \end{aligned}$$

où le signe $\geq$ se transforme en $=$ si et seulement si $\lambda_j = (f, e_j), -n \leq j \leq n$. Dans le cas de l'égalité, on a donc

$$\|f - g_0\|_2^2 = \|f\|_2^2 - \|g_0\|_2^2 \geq 0$$

où l'égalité peut être interprétée comme une généralisation du théorème de Pythagore, les fonctions $f - g_0$ et g_0 ayant un produit scalaire nul, $(f - g_0, g_0) = 0$, et étant donc « perpendiculaires ». $\square$

inégalité de Bessel

Corollaire 21 (Inégalité de Bessel). Si $f \in \mathcal{C}(\mathbb{T}^1)$, alors

$$\|f\|_2^2 \geq \sum_{j \in \mathbb{Z}} |c_j|^2.$$

Ainsi le carré de la longueur du « vecteur » f est au moins égal à la somme des carrés de la norme de ses « composantes » c_j . Dans le cas d'un espace vectoriel de dimension finie, la longueur d'un vecteur est précisément la somme des carrés de la norme de ses composantes. Le théorème 23 ci-dessous énoncera précisément ce résultat pour l'espace $\mathcal{C}(\mathbb{T}^1)$. Ainsi le présent corollaire n'est qu'une étape vers cet énoncé plus fort.

Démonstration. Rappelons que $c_j = \frac{1}{2\pi} \int_{\mathbb{T}^1} f(t) e^{-ijt} dt = (f, e_j)$. Par le théorème 20, on a :

$$\|f\|_2^2 \geq \sum_{j=-n}^n |c_j|^2, \quad \forall n$$

et puisque tous les $|c_j|^2 \geq 0$, la série converge et sa limite demeure plus petite ou devient égale à $\|f\|_2^2$:

$$\|f\|_2^2 \geq \sum_{j \in \mathbb{Z}} |c_j|^2.$$

□

distance moyenne
entre f et g

Définition 16. La distance moyenne entre f et $g \in \mathcal{C}(\mathbb{T}^1)$ est définie par :

$$\|f - g\|_2 = \sqrt{\frac{1}{2\pi} \int_{\mathbb{T}^1} |f - g|^2 dt}.$$

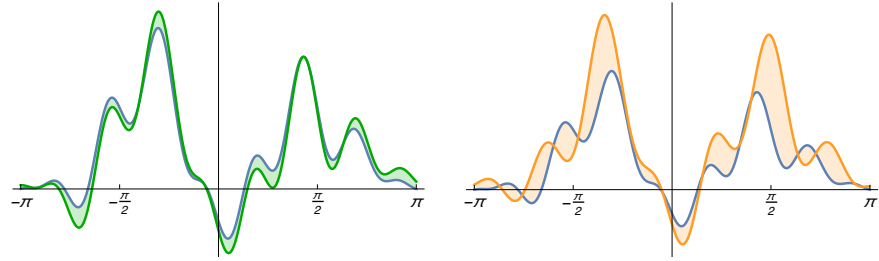


Figure 2.5 – Les fonctions à gauche sont plus voisines l'une de l'autre que ne l'est la paire à droite.

On appelle cette distance « en moyenne » (*mean square distance*) car elle fait la moyenne des carrés des écarts entre $f(t)$ et $g(t)$ pour tout $t \in \mathbb{T}^1$. L'analogie avec la distance entre $\vec{x}$ et $\vec{y} \in \mathbb{R}^n$ est donc

$$\sum_{i=1}^n (x_i - y_i)^2 \longleftrightarrow \frac{1}{2\pi} \int_{\mathbb{T}^1} |f - g|^2 dt.$$

À la figure 2.5, deux paires de fonctions sont tracées. Les deux fonctions à gauche sont plus voisines que ne le sont les deux fonctions à droite car la distance moyenne mesure l'écart moyen, c'est-à-dire une quantité reliée à l'aire entre les deux courbes. Nous sommes maintenant prêts à démontrer deux résultats-clés.

Théorème 22. Soit $f : \mathbb{T}^1 \rightarrow \mathbb{C}$ une fonction continue. Alors

$$\|f - s_n(f, \cdot)\|_2 \rightarrow 0 \quad \text{quand } n \rightarrow \infty.$$

Rappelons que $s_n(f, x)$ est la somme partielle $\sum_{j=-n}^n c_j e_j$ des premiers termes de la série de Fourier de f évaluée en x .

Démonstration. Par le théorème 7³, pour tout $\epsilon > 0$, il existe un polynôme trigonométrique $P(t)$ tel que

$$|P(t) - f(t)| < \epsilon, \quad \forall t \in \mathbb{T}^1$$

avec $P(t) = \sum_{j=-m}^m a_j e_j$. Entre autres

$$\|P - f\|_2 = \left(\frac{1}{2\pi} \int_{\mathbb{T}^1} |P(t) - f(t)|^2 dt \right)^{\frac{1}{2}} < \left(\frac{\epsilon^2}{2\pi} \int_{\mathbb{T}^1} dt \right)^{\frac{1}{2}} = \epsilon.$$

Maintenant, si $n \geq m$, alors $s_n(P) = P$ et donc

$$\|f - s_n(f)\|_2 \leq \|f - P\|_2 + \underbrace{\|P - s_n(P)\|_2}_{=0} + \|s_n(P) - s_n(f)\|_2$$

par l'inégalité du triangle et donc

$$\begin{aligned} &= \|f - P\|_2 + \|s_n(P - f)\|_2 \\ &\leq 2\|f - P\|_2, \quad \text{par le thm. 20} \\ &< 2\epsilon \end{aligned}$$

et donc $\|f - s_n(f)\|_2 \rightarrow 0$ lorsque $n \rightarrow \infty$. □

Et enfin :

Théorème 23 (Parseval). Si $f \in \mathcal{C}(\mathbb{T}^1)$, alors $\|f\|_2^2 = \sum_{n \in \mathbb{Z}} |c_n|^2$.

identité de Parseval
Marc-Antoine Parseval
(1755-1836)

Démonstration. La preuve du théorème 20 a montré l'égalité

$$\|f - s_n(f)\|_2^2 = \|f\|_2^2 - \sum_{j=-n}^n |c_j|^2.$$

Par le théorème précédent $\|f - s_n(f)\|_2 \rightarrow 0$ et donc

$$\|f\|_2^2 = \lim_{n \rightarrow \infty} \sum_{j=-n}^n |c_j|^2,$$

c'est-à-dire $\|f\|_2^2 = \sum_{j \in \mathbb{Z}} |c_j|^2$. □

3. On remarquera que ce théorème repose également sur le théorème de Fejér.

Le corollaire suivant est une conséquence immédiate.

Corollaire 24. Si $f \in \mathcal{C}(\mathbb{T}^1)$, alors ses coefficients de Fourier c_r tendent vers 0 en module lorsque r tend vers l'infini.

Cette identité de Parseval permet d'évaluer plusieurs séries fameuses. Voici des exemples d'applications des différents théorèmes.

Exemple. Développons la fonction $f : [0, \pi] \rightarrow \mathbb{R}$ définie par $f(x) = x(\pi - x)$ en série de cosinus.

$$\begin{aligned} a_0 &= \frac{2}{\pi} \int_0^\pi x(\pi - x) dx = \frac{2}{\pi} \left(\frac{\pi x^2}{2} - \frac{x^3}{3} \right) \Big|_0^\pi \\ &= \frac{2}{\pi} \left(\frac{\pi^3}{2} - \frac{\pi^3}{3} \right) = \frac{\pi^2}{3} \end{aligned}$$

et

$$\begin{aligned} a_n &= \frac{2}{\pi} \int_0^\pi x(\pi - x) \cos nx \, dx, \quad \text{par parties} \\ &= \frac{2}{\pi} \left(x(\pi - x) \frac{\sin nx}{n} \Big|_0^\pi - \int_0^\pi (\pi - 2x) \frac{\sin nx}{n} dx \right) \\ &= \frac{2}{\pi} \left(\frac{\pi}{n} \frac{\cos nx}{n} \Big|_0^\pi + \frac{2}{n} \int_0^\pi x \sin nx \, dx \right), \quad \text{par parties encore} \\ &= \frac{2}{\pi} \left[\frac{\pi((-1)^n - 1)}{n^2} + \frac{2}{n} \left(-x \frac{\cos nx}{n} \Big|_0^\pi + \int_0^\pi \frac{\cos nx}{n} dx \right) \right] \\ &= \frac{2}{\pi} \left[\frac{\pi((-1)^n - 1)}{n^2} + \frac{2}{n} \left(\frac{-\pi(-1)^n}{n} + \frac{\sin nx}{n^2} \Big|_0^\pi \right) \right] \\ &= \frac{2\pi}{\pi n^2} ((-1)^n - 1 - 2(-1)^n) = -\frac{2}{n^2} (1 + (-1)^n) \\ &= \begin{cases} 0, & \text{si } n \text{ est impair,} \\ -\frac{4}{n^2}, & \text{si } n \text{ est pair.} \end{cases} \end{aligned}$$

En posant $n = 2j$ puisque seuls les termes n pairs interviennent, nous obtenons :

$$x(\pi - x) = \frac{\pi^2}{6} - \sum_{j=1}^{\infty} \frac{1}{j^2} \cos 2jx \quad \text{pour } x \in [0, \pi].$$

La suite $\{\frac{1}{n^2}\}$ est sommable. (Voir la preuve du théorème 12.) Ainsi le théorème 10 nous assure que cette série de Fourier converge uniformément et donc que le signe « = » est enfin justifié. L'identité de Parseval nous donne de plus une identité

remarquable :

$$\begin{aligned}
 \|f\|_2^2 &= \frac{2}{2\pi} \int_0^\pi x^2(\pi - x)^2 dx \\
 &= \frac{1}{\pi} \int_0^\pi (\pi^2 x^2 - 2\pi x^3 + x^4) dx \\
 &= \frac{1}{\pi} \left(\frac{\pi^2 x^3}{3} - \frac{2\pi x^4}{4} + \frac{x^5}{5} \right) \Big|_0^\pi \\
 &= \frac{\pi^4}{30}.
 \end{aligned}$$

alors que le membre de gauche est (voir exercice 24 ci-dessous)

$$\frac{a_0^2}{2} + \sum_{k=1}^{\infty} (a_k^2 + b_k^2) = \frac{1}{2} \frac{\pi^4}{9} + \sum_{j=1}^{\infty} \frac{1}{j^4}$$

et donc :

$$\sum_{j=1}^{\infty} \frac{1}{j^4} = \frac{\pi^4}{15} - \frac{\pi^4}{18} = \frac{\pi^4}{90}.$$

Puisque le théorème 10 nous assure l'égalité entre $f(x) = x(\pi - x)$, $x \in [0, \pi]$, et la série de Fourier que nous venons de calculer, nous pouvons évaluer cette égalité en n'importe quel point pour obtenir des identités spécifiques. Par exemple, évaluée en $x = 0$, cette somme nous donne :

$$\frac{\pi^2}{6} = \sum_{j=1}^{\infty} \frac{1}{j^2}.$$

Ce n'est pas tout ! Évaluée en $x = \frac{\pi}{2}$, elle donne

$$\frac{\pi^2}{12} = \sum_{j=1}^{\infty} \frac{(-1)^{j+1}}{j^2}.$$

Les suites qui convergent uniformément peuvent être intégrées terme à terme. En d'autres mots, l'opération d'intégration peut être intervertie avec celle de prendre la limite. Similairement, si la suite de fonctions dérivables $\{f_n\}$ converge (ponctuellement) vers la fonction f et est telle que les fonctions dérivées f'_n soient intégrables sur $[a, b]$ et convergent uniformément vers une fonction continue g , alors

$$f'(x) = \lim_{n \rightarrow \infty} f'_n(x).$$

Ces deux propriétés permettent donc de calculer de nouvelles séries de Fourier sans trop d'effort. Dans l'exemple ci-dessus, la série converge uniformément (par le théorème 10) et on peut intégrer terme à terme l'égalité

$$x(\pi - x) - \frac{\pi^2}{6} = - \sum_{j=1}^{\infty} \frac{1}{j^2} \cos 2jx. \quad (*)$$

Le membre de gauche est

$$\begin{aligned}\int \left(x(\pi - x) - \frac{\pi^2}{6} \right) dx &= -\frac{x^3}{3} + \frac{\pi x^2}{2} - \frac{\pi^2 x}{6} + \text{constante} \\ &= -\frac{1}{3}x\left(x - \frac{\pi}{2}\right)(x - \pi) + \text{constante}\end{aligned}$$

où « constante » est la constante d'intégration. Le membre de droite devient, par la convergence uniforme :

$$\begin{aligned}-\int \sum_{j=1}^{\infty} \frac{1}{j^2} \cos 2jx \, dx &= -\sum_{j=1}^{\infty} \int \frac{1}{j^2} \cos 2jx \, dx \\ &= -\sum_{j=1}^{\infty} \frac{1}{2j^3} \sin 2jx + \text{nouvelle constante}.\end{aligned}$$

Comment se débarrasser des constantes ? Il suffit d'évaluer les deux membres en un point pour les fixer. En $x = 0$ les deux membres s'annulent et les deux constantes peuvent donc être mises à zéro :

$$\frac{2}{3}x\left(x - \frac{\pi}{2}\right)(x - \pi) = \sum_{j=1}^{\infty} \frac{1}{j^3} \sin 2jx.$$

Peut-on dériver terme à terme l'égalité (*) ? Dans l'exemple présent, ce n'est pas évident. En effet la dérivée des sommes partielles est

$$\sum_{j=0}^n \frac{2}{j} \sin 2jx$$

dont la convergence uniforme est loin d'être évidente. (Rappelons que la série $\{a_n = \sum_{j=0}^n \frac{1}{j}\}$ ne converge pas.) Il faudrait pouvoir montrer la convergence uniforme des sommes partielles pour justifier la dérivation terme à terme.

Exercices

25. Montrez que l'identité de Parseval est, en termes des coefficients de Fourier trigonométriques :

$$\frac{1}{\pi} \int_{-\pi}^{\pi} |f(x)|^2 dx = \frac{a_0^2}{2} + \sum_{n=1}^{\infty} (a_n^2 + b_n^2).$$

26*. Lesquelles des fonctions $(\cdot, \cdot)$ suivantes définissent un produit scalaire sur l'espace des fonctions continues sur l'intervalle $[-\pi, \pi]$.

(a)

$$(f, g) = \left(\int_{-\pi}^{\pi} f(t)^2 (g(t)^2)^* dt \right)^{\frac{1}{2}}$$

(b)

$$(f, g) = \int_{-\pi}^{\pi} f(t)g(t)^* e^{-t^2} dt$$

(c)

$$(f, g) = \int_0^{\pi} f(t)g(t)^* dt$$

27*. Utiliser l'identité de Parseval pour évaluer les séries suivantes.

(a)

$$\sum_{n \in \mathbb{N}} \frac{1}{n^6}$$

(b)

$$\sum_{\substack{n=1 \\ n \text{ impair}}}^{\infty} \frac{1}{n^4}$$

(c)

$$\sum_{n \in \mathbb{N}} \frac{1}{(n^2 + 1)^2}$$

(d)

$$\frac{1}{1^2 3^2} + \frac{1}{3^2 5^2} + \frac{1}{5^2 7^2} + \frac{1}{7^2 9^2} + \dots$$

28*. Soit $L^1([0, 1])$ l'ensemble des fonctions intégrables au sens de Riemann sur l'intervalle $[0, 1]$.(a) Montrer que les fonctions f_{ij} , introduites à l'exercice 12 de la section 2.2, forment un ensemble orthonormé :

$$(f_{ij}, f_{kl}) = \delta_{ik} \delta_{jl}$$

où $(\cdot, \cdot)$ est le produit scalaire

$$(f, g) = \int_0^1 f(x)g(x) dx.$$

(b) Soit $f \in L^1$. Montrer comment trouver des coefficients ϕ_{ij} tels que sur $[0, 1]$ la série

$$\phi_{00} + \sum_{i=1}^{\infty} \sum_{j=0}^{2^{i-1}-1} \phi_{ij} f_{ij}$$

converge vers f au sens de Cesàro. (A cause des deux exercices 11 et 12 de la section 2.2, cet exercice devrait être facile.)

(c) Développer la fonction $f(x) = x$ en termes des fonctions f_{ij} .

(d) A l'aide d'un langage de manipulations symboliques, tracer le graphe des premiers termes de la série ci-dessus. Approximent-ils bien le graphe de x ?

29. (a) (Körner.) Montrer que les coefficients de Fourier c_r de la fonction en dent de scie $h(t) = t - \pi, 0 < t < 2\pi$ et $h(0) = 0$ sont

$$c_r = \frac{i}{r}, \quad r \neq 0 \quad \text{et} \quad c_0 = 0.$$

(b) Soient $P(t) = \sum_{j=0}^n a_j e^{-ijt}$, $Q(t) = \sum_{k=0}^n b_k e^{-ikt}$. Montrer que

$$\frac{1}{2\pi i} \int_{\mathbb{T}^1} P(t) Q(t) h(t) e^{-it} dt = \sum_{j=0}^n \sum_{k=0}^n \frac{a_j b_k}{j+k+1}.$$

(c) Utiliser le fait que $|h(t)e^{-it}| \leq \pi$ pour tout t et l'inégalité de Cauchy-Schwarz-Buniakowski pour $(P(t), (Q(t)h(t)e^{-it})^*)$ pour montrer que

$$\left| \frac{1}{2\pi} \int_{\mathbb{T}^1} P(t) Q(t) h(t) e^{-it} dt \right| \leq \pi \left(\sum_{j=0}^n |a_j|^2 \sum_{k=0}^n |b_k|^2 \right)^{\frac{1}{2}}.$$

(d) En déduire que :

$$\sum_{j=0}^n \sum_{k=0}^n \frac{|a_j| |b_k|}{j+k+1} \leq \pi \left(\sum_{j=0}^n |a_j|^2 \right)^{\frac{1}{2}} \left(\sum_{k=0}^n |b_k|^2 \right)^{\frac{1}{2}}.$$

(Se rappeler que l'identité prouvée en (c) est vraie pour *tous* polynômes P et Q .)

(e) Prouver que, si $\sum_{j=0}^n |a_j|^2$ et $\sum_{k=0}^n |b_k|^2$ convergent lorsque $n \rightarrow \infty$, alors $\sum_{j=0}^{\infty} \sum_{k=0}^{\infty} \frac{|a_j| |b_k|}{j+k+1}$ existe et est $\leq \pi (\sum_{j=0}^{\infty} |a_j|^2)^{\frac{1}{2}} (\sum_{k=0}^{\infty} |b_k|^2)^{\frac{1}{2}}$.

30. Montrer les affirmations suivantes.

(a) Si $z_1, z_2 \in \mathbb{C}$, alors

$$z_1 z_2^* = \frac{1}{4}(|z_1 + z_2|^2 - |z_1 - z_2|^2 + i|z_1 + iz_2|^2 - i|z_1 - iz_2|^2).$$

(b) Si $f, g \in \mathcal{C}(\mathbb{T}^1)$, alors

$$(f, g) = \frac{1}{4}(\|f + g\|_2^2 - \|f - g\|_2^2 + i\|f + ig\|_2^2 - i\|f - ig\|_2^2).$$

(c) Si $f, g \in \mathcal{C}(\mathbb{T}^1)$, alors

$$(f, g) = \sum_{j \in \mathbb{Z}} c_j d_j^*$$

où c_j et d_j sont les coefficients des développements en série de Fourier (d'exponentielles) de f et de g respectivement.

31. Cet exercice présente une preuve du théorème de Dirichlet (thm. 13).

preuve du
théorème de Dirichlet

(a) Dans l'exercice 9 de la section 2.2, nous avons réécrit les sommes partielles

$$s_n(x) = \sum_{k=-n}^n c_k e^{ikx}$$

sous la forme

$$s_n(x) = \frac{1}{2\pi} \int_{\mathbb{T}^1} f(x+t) d_n(t) dt$$

où $d_n(t)$ est le noyau de Dirichlet

$$d_n(t) = \frac{\sin(n + \frac{1}{2})t}{\sin \frac{t}{2}}.$$

Montrer qu'on peut réécrire également $s_n(x)$ sous la forme

$$s_n(x) = \frac{1}{2\pi} \int_0^\pi f(x+t) d_n(t) dt + \frac{1}{2\pi} \int_0^\pi f(x-t) d_n(t) dt.$$

Les deux termes de cette somme seront notés $s_n^+(x)$ et $s_n^-(x)$ respectivement. Nous allons montrer que $s_n^+(x) - f(x^+)/2$ et $s_n^-(x) - f(x^-)/2$ tendent vers 0 lorsque $n \rightarrow \infty$.

(b) En utilisant l'identité $\sin(n + \frac{1}{2})t = \cos nt \sin \frac{t}{2} + \sin nt \cos \frac{t}{2}$, montrer que $s_n^+(x) - f(x^+)/2$ peut être écrit sous la forme

$$\frac{1}{2\pi} \left(\int_0^\pi (f(x+t) - f(x^+)) \cos nt dt + \int_0^\pi (f(x+t) - f(x^+)) \cot \frac{t}{2} \sin nt dt \right).$$

(c) Montrer que, pour tout x , la fonction

$$g_1(t) = \begin{cases} (f(x+t) - f(x^+)), & 0 \leq t < \pi, \\ 0, & -\pi \leq t \leq 0 \end{cases}$$

est continue par morceaux.

(d) Montrer que, pour tout x , la fonction

$$g_2(t) = \begin{cases} (f(x+t) - f(x^+)) \cot \frac{t}{2}, & 0 < t < \pi, \\ 0, & -\pi \leq t \leq 0 \end{cases}$$

est continue par morceaux. (Vous devrez montrer que $\lim_{t \rightarrow 0^+} g_2(t)$ existe. Il en découlera que g_2 est bornée et donc intégrable.)

(e) Utiliser le lemme 11 pour montrer que les coefficients de Fourier de g_1 et g_2 tendent vers 0 lorsque $n \rightarrow \infty$. Ceci démontre que $s_n^+(x) - f(x^+)/2 \xrightarrow[n \rightarrow \infty]{} 0$. Une preuve similaire montre que $s_n^-(x) - f(x^-)/2 \xrightarrow[n \rightarrow \infty]{} 0$ et clôt la démonstration du théorème de Dirichlet.

Chapitre 3

Problème de Sturm-Liouville et fonctions spéciales

3.1 Rappel sur les équations différentielles

Un des exemples les plus simples d'équations différentielles est celui rencontré dans l'étude de la corde vibrante. Après séparation des variables, l'équation différentielle pour la partie spatiale X était :

$$X'' = -\lambda^2 X.$$

Pour obtenir cette équation, plusieurs étapes ont été suivies mais le point de départ était l'équation de base de la mécanique : $F = ma$. Ceci est très courant tant en physique qu'en mathématiques : certains principes de base (comme le principe d'action-réaction) mènent à des équations différentielles (comme l'équation de la corde vibrante). Résoudre l'équation ci-dessus peut être fait en répondant simplement à la question : quelle est la fonction qui, dérivée deux fois, redonne elle-même, à une constante négative près ? Beaucoup auront deviné que la réponse est les fonctions sinus et cosinus (ou mieux, une combinaison linéaire des deux). Cette équation était « facile ». Mais beaucoup d'équations en apparence aussi simple sont difficiles à résoudre : essayez $y'' = y^2$ ou $y'' = \sin y$ ou encore $(y'')^2 = y' + y$.

Si vous n'avez jamais étudié les techniques de base de solutions d'équations différentielles, il est sans doute utile de vous dire que ces techniques ressemblent un peu en nombre et en difficultés aux techniques pour trouver les primitives des fonctions. Cette section *n'a pas* pour but de vous enseigner un grand nombre de techniques. Tout d'abord, nous allons plutôt décrire plusieurs caractéristiques qui permettent de « classer » les équations différentielles. Puis nous donnerons quelques exemples de solutions d'équations différentielles ordinaires linéaires (voir ci-dessous) en expliquant la relation entre l'« ordre » de l'équation différentielle et le « nombre de solutions linéaires indépendantes ».

Nomenclature. Les équations différentielles lient des fonctions et leurs dérivées. Les fonctions sont souvent appelées *variables dépendantes* et les variables dont elles dépendent *variables indépendantes*. Si $y = y(x)$, l'équation différentielle $y'' + xy' + y = 0$ fait intervenir la variable dépendante y (la fonction) et la variable indépendante x .

édo et éd

Les équations différentielles sont divisées en équations différentielles ordinaires et équations différentielles aux dérivées partielles. Une équation différentielle est ordinaire si elle contient une ou des fonctions et leurs dérivées par rapport à une *unique* variable indépendante. Elle est aux dérivées partielles si elle fait intervenir les dérivées d'une ou plusieurs fonctions par rapport à plus d'une variable indépendante. Le jargon scientifique nomme familièrement ces deux classes *édo* et *édp* (ou en anglais *ode* et *pde*).

linéarité

Les équations différentielles (ordinaires et aux dérivées partielles) *linéaires* sont celles qui ont été le plus étudiées du XVIII^{ème} au XIX^{ème} siècle. Elles sont « plus simples » et c'est pour les équations aux dérivées partielles linéaires qu'a été développée l'analyse de Fourier. Une équation est linéaire si chacun de ses termes (de ses monômes) ne fait apparaître aucune des variables dépendantes ou les fait apparaître linéairement. Par exemple, si f et g sont deux variables dépendantes et t la variable indépendante, le monôme fg n'est pas linéaire, mais les termes f' et g'' le sont. Ainsi l'équation $(fg)' = f + g$ n'est pas linéaire (même si les monômes f et g le sont) et l'équation $f''' + g' = f + 3 \sin t$ l'est.

ordre

L'*ordre* d'une équation différentielle est l'ordre maximal des dérivées intervenant dans l'équation. Les équations $f''' - f' = \sin t$ et $f''(f')^2 - (f')^3 = 0$ sont respectivement du 3^{ème} et du second ordre.

Avec cette nomenclature, il est possible de dire pourquoi Fourier a inventé les séries qui portent son nom : pour ramener l'étude des équations aux dérivées partielles linéaires à celle des équations différentielles ordinaires linéaires.

homogénéité

Les équations différentielles ordinaires linéaires. Les édo linéaires se séparent elles-mêmes en édos *homogènes* et *non-homogènes*. Chacun des monômes d'une édo linéaire homogène fait apparaître linéairement la fonction (variable dépendante). Les non-homogènes contiennent des monômes qui ne contiennent pas la variable dépendante. Ainsi $xy'' + (1-x)y' + y = 0$ est homogène mais $y'' = \sin x$ ne l'est pas.

L'importance des édos linéaires homogènes est que leurs solutions forment un espace vectoriel. Par exemple les solutions de

$$x^2 y'' + xy' + (\lambda^2 x^2 - n^2)y = 0$$

où λ et n sont des constantes, forment un espace vectoriel. Si y est une solution, cy l'est aussi. Et si y_1 et y_2 sont des solutions :

$$\begin{aligned} x^2 y_1'' + xy_1' + (\lambda^2 x^2 - n^2)y_1 &= 0, \\ x^2 y_2'' + xy_2' + (\lambda^2 x^2 - n^2)y_2 &= 0, \end{aligned}$$

alors $(y_1 + y_2)$ l'est aussi :

$$\begin{aligned} x^2(y_1 + y_2)'' + x(y_1 + y_2)' + (\lambda^2 x^2 - n^2)(y_1 + y_2) \\ = x^2 y_1'' + x y_1' + (\lambda^2 x^2 - n^2)y_1 + x^2 y_2'' + x y_2' + (\lambda^2 x^2 - n^2)y_2 \\ = 0. \end{aligned}$$

Nous avons déjà rencontré une équation linéaire homogène :

$$y'' = -\lambda^2 y$$

et nous avons écrit ses solutions sous la forme

$$y = a \cos \lambda x + b \sin \lambda x$$

montrant bien que les solutions forment un espace vectoriel de dimension (au moins) 2. A quoi servent les constantes a et b dans la solution ci-dessus ? Elles jouent le même rôle que les constantes d'intégration : chaque fois qu'une intégrale est prise, une constante est ajoutée à la primitive trouvée. Dans le cas d'une édo du premier ordre (c'est-à-dire faisant intervenir qu'une dérivée première), nous devons intuitivement « intégrer » une seule fois et une seule constante apparaîtra. Si elle est du second ordre comme $y'' = -\lambda^2 y$, nous devons « intégrer » deux fois et deux constantes apparaîtront. Et ainsi de suite lorsque l'ordre de l'édo croît. Il semble donc raisonnable de prédire qu'une édo homogène d'ordre n nécessitera n constantes d'intégration. (L'argument intuitif peut-être démontré !) Une solution d'une édo d'ordre n faisant apparaître n constantes est appelée la solution générale de l'équation. Ainsi $y = a \cos \lambda x + b \sin \lambda x$ est la solution générale de l'équation du second ordre $y'' = -\lambda^2 y$.

Comment les constantes intervenant dans la solution générale sont-elles fixées ? Elles le sont par les conditions initiales ou par les conditions aux limites. (Si la variable indépendante peut être interprétée comme le temps, on parle de conditions initiales et si elle est l'espace, on parle plutôt de conditions aux limites.) Commençons par l'exemple simple d'une particule tombant en chute libre dans un champ de gravitation constant. Sa position vérifie l'équation $ma = mg$ où g est une constante, a l'accélération $\frac{d^2x}{dt^2}$ et m la masse. On peut intégrer une première fois

$$\frac{dx}{dt} = gt + c_1,$$

où c_1 est la première constante d'intégration, et une seconde fois

$$x(t) = \frac{gt^2}{2} + c_1 t + c_2.$$

Si on sait que la masse a été lâchée d'une hauteur de 10 unités à vitesse nulle, nous devons donc fixer c_1 et c_2 de façon, qu'en $t=0$, nous ayons :

$$x(0) = 10$$

et

$$\frac{dx}{dt}(0) = 0.$$

Ainsi la solution est

$$x(t) = \frac{gt^2}{2} + 10.$$

solution particulière

Lorsque les constantes de la solution générale sont fixées par les conditions initiales, la solution obtenue est dite une solution particulière.

Encore un exemple. Un ressort est soumis à une force de rappel dépendant de sa position : $F(x) = -kx$ où k est une constante positive. L'équation

$$ma = F$$

est donc

$$m \frac{d^2x}{dt^2} = -kx.$$

Cette équation a la même forme $x'' = \lambda^2 x$ et la solution générale est

$$x(t) = a \cos \sqrt{\frac{k}{m}} t + b \sin \sqrt{\frac{k}{m}} t.$$

Supposons que nous ayons mesuré la position de la masse à deux instants $t = 0$ et $t = \sqrt{\frac{m}{k}} \frac{\pi}{4}$ et que nous ayons obtenu :

$$x(0) = 1 \quad \text{et} \quad x\left(\sqrt{\frac{m}{k}} \frac{\pi}{4}\right) = \frac{1}{2}.$$

Alors ces deux conditions mènent à :

$$1 = a \quad \text{et} \quad \frac{1}{2} = \frac{a}{\sqrt{2}} + \frac{b}{\sqrt{2}};$$

donc la solution est

$$a = 1 \quad \text{et} \quad b = \left(\frac{1}{\sqrt{2}} - 1\right).$$

Exercices

1. Identifier les équations différentielles ordinaires et les équations différentielles aux dérivées partielles dans ce qui suit.

(a) $\frac{dy}{dx} + y \frac{dy}{dx^2} + 2x^3 - \sin(xy) = 0$

(b) $\frac{dy}{dx} = \frac{d}{dx} \arctan \frac{1}{\sqrt{x^2 - 1}}$

$$(c) \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} = \frac{1}{c^2} \frac{\partial^2 f}{\partial t^2}$$

$$(d) i\hbar \frac{\partial \psi}{\partial t} = -\nabla^2 \psi$$

$$(e) \frac{d^2 \vec{x}}{dt^2} = -A\vec{x} \text{ où } A \text{ est une matrice } n \times n \text{ et } \vec{x}(t) \in \mathbb{R}^n.$$

2. Identifier parmi les équations différentielles suivantes celles qui sont linéaires.

$$(a) y''y + 7y = 0$$

$$(b) y'' = -\lambda^2 y$$

$$(c) xy'' = 1$$

$$(d) y'' = y^2$$

$$(e) yy' - y'' = 0$$

$$(f) y'' + (1-x)y' + x^2y = 0$$

$$(g) y'' = \sin y$$

$$(h) r'' = -\frac{1}{r^2}$$

$$(i) \nabla^2 f = \frac{\partial^2 f}{\partial t^2}$$

$$(j) \frac{\partial \psi}{\partial t} - \frac{\partial \psi}{\partial x} = |\psi|^2 \psi$$

(k) Le système

$$\vec{\nabla} \cdot \vec{E} = \rho$$

$$\vec{\nabla} \cdot \vec{H} = 0$$

$$\vec{\nabla} \times \vec{E} + \vec{H} = 0$$

$$\vec{\nabla} \times \vec{H} - \vec{E} = \vec{J}$$

où $\vec{J}$ et ρ sont des fonctions connues, $\vec{E}$ et $\vec{H}$ sont six variables dépendantes et x, y, z et t sont les variables indépendantes.

3. Quel est l'ordre des équations différentielles suivantes ?

$$(a) y''y' + 7y = 0$$

$$(b) y' + y'' + x^7y = 0$$

$$(c) (y'')^2 + y = 0$$

$$(d) yy'y''' - y''y'''' = 0$$

$$(e) y'' = -\lambda^2 y$$

(f) $\nabla^2 \psi = \lambda \psi$

4. Parmi les équations différentielles ordinaires suivantes, quelles sont les homogènes ?

(a) $y'' + xy' + y = 0$

(b) $y'' = -\lambda^2 y$

(c) $y'' = -\lambda^2 y + x$

(d) $y'' = y \sin x$

(e) $y''' + x^2 y' = y$

5*. Pour les quatre équations suivantes, vérifier que la solution donnée est la solution générale et trouver la solution particulière qui satisfait les conditions initiales suggérées.

(a)

$$(x+1)y'' = y' \quad y = c \left(x + \frac{x^2}{2} \right) + d \quad y(0) = 1, y'(0) = 1$$

(b)

$$xy'' = (1 - 2x^2)y' \quad y = ce^{-x^2} + d \quad y(0) = 1, y'(1) = e^{-1}$$

(c)

$$4y'' - 4y' + y = 0 \quad y = ce^{x/2} + dxe^{x/2} \quad y(0) = 1, y'(0) = 1$$

(d)

$$4x^2 y'' + y = 0 \quad y = c\sqrt{x} + d\sqrt{x} \ln x \quad y(1) = 0, y'(1) = 1$$

3.2 Séparation de variables

Nous avons déjà rencontré quelques équations différentielles aux dérivées partielles :

équation de la corde vibrante	$\frac{1}{c^2} \frac{\partial^2 f}{\partial t^2} = \frac{\partial^2 f}{\partial x^2}$
équation du tambour	$\frac{1}{c^2} \frac{\partial^2 f}{\partial t^2} = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2}$
équation de la chaleur	$\frac{1}{a^2} \frac{\partial f}{\partial t} = \frac{\partial^2 f}{\partial x^2}$

D'autres sont également communes telles :

$$\begin{aligned} \text{équation d'onde} \quad & \frac{1}{c^2} \frac{\partial^2 f}{\partial t^2} = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} + \frac{\partial^2 f}{\partial z^2} \\ \text{équation de Schrödinger} \quad & i\hbar \frac{\partial \psi}{\partial t} = -\nabla^2 \psi + V(r)\psi \end{aligned}$$

où $V(r)$ est le potentiel dans lequel baigne la particule considérée. Dans les équations ci-dessus, les nombres c , a et $\hbar$ sont des constantes. Le "truc" de la séparation des variables que nous avons utilisé à la section 1.5 s'applique à toutes ces équations. Non seulement s'y applique-t-il en coordonnées cartésiennes mais aussi dans plusieurs autres systèmes de coordonnées. Le but de cette section est de présenter des exemples de séparation dans des systèmes de coordonnées autres que les cartésiennes.

Les deux exemples que nous ferons en un peu de détail sont l'équation des ondes en coordonnées cylindriques et sphériques. L'équation d'onde est

$$\frac{1}{c^2} \frac{\partial^2 f}{\partial t^2} = \nabla^2 f$$

où

$$\nabla^2 = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2}.$$

La constante c est la vitesse de propagation des ondes. Dans le cas de la lumière, cette vitesse est $c = 2.998 \times 10^8$ m/sec. Par la suite nous poserons cette constante c égale à 1. Comme précédemment, nous chercherons des solutions de la forme

$$f(\vec{r}, t) = T(t)\psi(\vec{r}).$$

(Contrairement à précédemment cependant, ψ dépend maintenant de 3 variables.) L'équation des ondes s'écrit alors

$$T''\psi = T\nabla^2\psi$$

ou encore

$$\frac{T''}{T} = \frac{1}{\psi} \nabla^2 \psi.$$

L'argument de séparabilité peut être à nouveau répété : puisque le membre de gauche ne dépend que de t et celui de droite que de $\vec{r}$, il faut donc que ces deux fonctions soient égales à une même constante, c'est-à-dire

$$T'' = \lambda T \quad \text{et} \quad \nabla^2 \psi = \lambda \psi.$$

La première équation $T'' = \lambda T$ a trois types de solutions :

$$\begin{aligned} T(t) &= ae^{\sqrt{\lambda}t} + be^{-\sqrt{\lambda}t}, & \text{si } \lambda > 0, \\ T(t) &= at + b, & \text{si } \lambda = 0, \\ T(t) &= a \cos(\sqrt{-\lambda}t) + b \sin(\sqrt{-\lambda}t), & \text{si } \lambda < 0. \end{aligned}$$

où a et b sont des constantes. Quoique les équations de propagation n'interdisent pas les deux premiers types (exponentiel et linéaire), le troisième est plus usuel. On le retrouve dans le cas de la corde vibrante (comme nous l'avons vu à la section 1.5), des ondes électromagnétiques dans le vide, etc. Nous choisirons donc une constante de séparation négative $\lambda = -k^2$ pour laquelle la solution prend la forme

$$T(t) = a \cos kt + b \sin kt.$$

équation d'Helmholtz
Hermann von Helmholtz
(1821-1894)

L'équation résultante pour $\psi(\vec{r})$ est alors l'équation de Helmholtz :

$$\nabla^2 \psi + k^2 \psi = 0.$$

C'est ici que notre chemin bifurque de la corde vibrante. Plutôt que d'obtenir une équation différentielle *ordinaire* pour la dépendance spatiale, l'équation demeure une équation différentielle aux dérivées partielles. Mais le truc de la séparation peut être employé à nouveau ! Nous le ferons dans deux cas : les systèmes de coordonnées cylindriques et sphériques. (Voir l'exercice 1.)

coordonnées cylindriques

Le système de coordonnées cylindriques est donné par

$$x = r \cos \theta$$

$$y = r \sin \theta$$

$$z = z$$

Le laplacien ∇^2 s'écrit en ces coordonnées sous la forme

$$\nabla^2 = \frac{\partial^2}{\partial r^2} + \frac{1}{r} \frac{\partial}{\partial r} + \frac{1}{r^2} \frac{\partial^2}{\partial \theta^2} + \frac{\partial^2}{\partial z^2}.$$

Supposons que la fonction ψ soit de la forme

$$\psi(r, \theta, z) = R(r)\Theta(\theta)Z(z).$$

L'équation d'Helmholtz sera

$$R''\Theta Z + \frac{1}{r}R'\Theta Z + \frac{1}{r^2}R\Theta''Z + R\Theta Z'' + k^2 R\Theta Z = 0$$

ou encore, après division par $R\Theta Z$:

$$\underbrace{\left(\frac{R''}{R} + \frac{1}{r} \frac{R'}{R} \right)}_{\text{ne dépend que de } r \text{ et } \theta} + \underbrace{\frac{1}{r^2} \frac{\Theta''}{\Theta}}_{\text{ne dépend que de } \theta} + \underbrace{\frac{Z''}{Z}}_{\text{ne dépend que de } z} + k^2 = 0.$$

L'expression ci-dessus dépend des trois variables r, θ et z . Même si la dépendance en r et θ n'est pas séparée (à cause du terme $\frac{1}{r^2} \frac{\Theta''}{\Theta}$ qui contient les deux variables), le terme $\frac{Z''}{Z}$ ne dépend que de z alors que le reste ne dépend que de r et θ . Il faut donc que ce terme soit constant. Alors il faut qu'il existe une constante λ_1 telle que :

$$\frac{Z''}{Z} = \lambda_1$$

et l'équation pour R et Θ devient alors (après multiplication par r^2)

$$\underbrace{r^2 \left(\frac{R''}{R} + \frac{1}{r} \frac{R'}{R} \right)}_{\text{ne dépend que de } r} + \underbrace{\frac{\Theta''}{\Theta}}_{\text{que de } \theta} + \underbrace{(k^2 + \lambda_1)r^2}_{\text{que de } r} = 0.$$

A cause de la multiplication par r^2 le terme qui dépendait des deux variables r et θ ne dépend plus maintenant que de θ et la dépendance est maintenant séparée. Une autre constante de séparation doit exister telle que

$$\frac{\Theta''}{\Theta} = \lambda_2$$

et finalement

$$r^2 \left(\frac{R''}{R} + \frac{1}{r} \frac{R'}{R} \right) + (\lambda_2 + (k^2 + \lambda_1)r^2) = 0$$

ou encore, en multipliant par R/r :

$$\frac{d}{dr} \left(r \frac{dR}{dr} \right) + \left(r(k^2 + \lambda_1) + \frac{\lambda_2}{r} \right) R = 0.$$

Nous avons introduit trois constantes de séparation : k , λ_1 et λ_2 . Dans le cas de la corde vibrante, la constante de séparation était fixée par les conditions aux limites. Ici la continuité jouera également un rôle. Par exemple, si la fonction $\psi = R\Theta Z$ est continue, il faut que la fonction Θ soit telle que $\Theta(0) = \Theta(2\pi)$. Une façon simple d'assurer cette contrainte est que la solution de $\Theta'' = \lambda_2 \Theta$ soit une combinaison linéaire de sin et de cos dont les arguments sont de la forme $m\theta$, $m \in \mathbb{N}$. Donc

$$\lambda_2 = -m^2$$

et

$$\Theta_m(\theta) = a \cos m\theta + b \sin m\theta, \quad m \in \mathbb{N}.$$

(Ce choix assure de plus que Θ est analytique sur le tore $\mathbb{T}^1$.) Nous allons demander (à titre exemple) que la fonction Z s'annule en $z = 0$ et L . Il faudra donc, pour que $Z'' = \lambda_1 Z$ soit satisfaite, que la constante λ_1 soit de la forme

$$\lambda_1 = -\frac{n^2 \pi^2}{L^2}, \quad n \in \mathbb{N}$$

et

$$Z_n(z) = \sin \frac{n\pi z}{L}.$$

(Cette étape est trop rapide ? Revoir le traitement des conditions aux limites pour la corde vibrante.) Alors l'équation pour R devient

$$\frac{d}{dr} \left(r \frac{dR}{dr} \right) + \left(r \left(k^2 - \frac{n^2 \pi^2}{L^2} \right) - \frac{m^2}{r} \right) R = 0.$$

Il est usuel d'introduire la constante

$$\kappa = \sqrt{k^2 - \frac{n^2\pi^2}{L^2}}.$$

Si on introduit une nouvelle variable $x = \kappa r$ et la fonction

$$y = y(x) = R\left(\frac{x}{\kappa}\right),$$

les dérivées deviennent :

$$\frac{dy}{dx} = \frac{d}{dx} R\left(\frac{x}{\kappa}\right) = \frac{1}{\kappa} R' \quad \text{et} \quad \frac{d^2y}{dx^2} = \frac{1}{\kappa^2} R''.$$

L'équation résultante est

$$\frac{d}{dx} \left(x \frac{dy}{dx} \right) + \left(x - \frac{m^2}{x} \right) y = 0$$

équation de Bessel
Friedrich Wilhelm Bessel
(1784-1846)

$Z(z) = 1$ et
l'équation du tambour

et elle joue souvent un rôle important dans les problèmes où la symétrie cylindrique est réalisée. Elle se nomme l'équation de Bessel.

Même si nous avons traité la propagation des ondes en trois dimensions à symétrie cylindrique, telle la propagation des ondes dans les câbles co-axiaux (la câblo-distribution des signaux de télévision se fait par de tels câbles), il est facile d'observer que, en posant $Z(z) = 1$, ce traitement s'applique à l'équation du tambour et donc les fonctions de Bessel sont la clef de la physique du tambour. (En posant $Z(z) = 1$, l'équation aux dérivées partielles perd sa dépendance en z et devient donc une équation aux dérivées partielles en trois variables : t, r et θ . Il s'agit de l'équation du tambour.)

coordonnées sphériques

Le système de coordonnées sphériques est donné par

$$x = r \sin \theta \cos \phi$$

$$y = r \sin \theta \sin \phi$$

$$z = r \cos \theta$$

où $\theta \in (0, \pi)$ et $\phi \in (0, 2\pi)$ (voir figure 3.1). Le laplacien $\nabla^2 \psi$, en ces coordonnées, se lit :

$$\nabla^2 \psi = \frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\partial \psi}{\partial r} \right) + \frac{1}{r^2 \sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial \psi}{\partial \theta} \right) + \frac{1}{r^2 \sin^2 \theta} \frac{\partial^2 \psi}{\partial \phi^2}.$$

A nouveau, après la séparation du temps dans l'équation des ondes, l'équation à résoudre est l'équation d'Helmholtz en trois dimensions qui se lit :

$$\frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\partial \psi}{\partial r} \right) + \frac{1}{r^2 \sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial \psi}{\partial \theta} \right) + \frac{1}{r^2 \sin^2 \theta} \frac{\partial^2 \psi}{\partial \phi^2} + k^2 \psi = 0.$$

Nous chercherons les solutions sous la forme

$$\psi(r, \theta, \phi) = R(r)\Theta(\theta)\Phi(\phi).$$

Notre traitement sera plus sommaire et nous ne donnerons ici que les étapes principales. (Vous devrez refaire les étapes intermédiaires. Exercice !) Après remplacement de ψ par $R\Theta\Phi$ et multiplication par $r^2 \sin^2 \theta / R\Theta\Phi$, l'équation se lit

$$\frac{\sin^2 \theta}{R} \frac{\partial}{\partial r} (r^2 R') + \frac{\sin \theta}{\Theta} \frac{\partial}{\partial \theta} (\Theta' \sin \theta) + \underbrace{\frac{1}{\Phi} \Phi''}_{\text{ne dépend que de } \phi} + k^2 r^2 \sin^2 \theta = 0.$$

Puisque le rapport $\frac{\Phi''}{\Phi}$ ne dépend que de ϕ alors que le reste de l'équation dépend des deux autres variables r et θ , une première constante de séparation est introduite :

$$\frac{\Phi''}{\Phi} = \lambda_1.$$

En divisant l'équation résiduelle par $\sin^2 \theta$, l'équation devient

$$\underbrace{\frac{1}{R} \frac{\partial}{\partial r} (r^2 R') + k^2 r^2}_{\text{ne dépend que de } r} + \underbrace{\frac{1}{\Theta \sin \theta} \frac{\partial}{\partial \theta} (\Theta' \sin \theta) + \frac{\lambda_1}{\sin^2 \theta}}_{\text{ne dépend que de } \theta} = 0.$$

Maintenant l'équation est séparée et une seconde constante de séparation λ_2 peut être introduite. L'équation pour Θ est

$$\frac{d}{d\theta} (\Theta' \sin \theta) + \frac{\lambda_1}{\sin \theta} \Theta - \lambda_2 \Theta \sin \theta = 0$$

et celle pour R

$$\frac{d}{dr} (r^2 R') + (k^2 r^2 + \lambda_2) R = 0.$$

Encore de nouvelles équations ! Des cas particuliers vont nous intéresser. A nouveau, il est naturel d'imposer la continuité de la fonction ψ qui sera satisfaite si Φ est périodique. Alors

$$\Phi'' = \lambda_1 \Phi$$

et

$$\Phi = a \cos m\phi + b \sin m\phi, \quad m \in \mathbb{N}$$

avec

$$\lambda_1 = -m^2.$$

L'équation pour θ devient alors

$$\frac{d}{d\theta} (\Theta' \sin \theta) - \frac{m^2}{\sin \theta} \Theta - \lambda_2 \Theta \sin \theta = 0.$$

Il est usuel d'introduire la nouvelle variable $x = \cos \theta$ et la fonction $y(x) = \Theta(\arccos x)$. Les dérivées sont obtenues par la règle de dérivation en chaîne

$$\frac{dy}{dx} = \frac{d\Theta}{d\theta} \frac{d\theta}{dx} = -\frac{1}{\sqrt{1-x^2}} \Theta'$$

et puisque $1 - x^2 = \sin^2 \theta$, pour $\theta \in [0, \pi]$ (qui est précisément le domaine de la coordonnée sphérique θ), la fonction $\sin \theta$ est toujours positive et

$$\Theta' = -\sin \theta \frac{dy}{dx} \quad \text{et} \quad \frac{d}{d\theta} = \frac{dx}{d\theta} \frac{d}{dx} = -\sin \theta \frac{d}{dx}.$$

En réécrivant l'équation différentielle en termes de ces nouvelles variables, on obtient l'équation de Legendre associée

$$\frac{d}{dx} \left((1 - x^2) \frac{dy}{dx} \right) - \frac{m^2}{1 - x^2} y - \lambda_2 y = 0.$$

Nous allons étudier le cas $m = 0$ à la section 3.4. Ceci correspond au cas où la solution ψ ne dépend pas de ϕ et alors l'équation se lit

$$\frac{d}{dx} \left((1 - x^2) \frac{dy}{dx} \right) = \lambda_2 y.$$

L'équation pour R est aussi une équation fameuse. Si la constante de séparation λ_2 a la forme $\lambda_2 = -l(l + 1)$, $l \in \mathbb{N} \cup \{0\}$, alors l'équation radiale se lit :

$$\frac{d}{dr} (r^2 R') + (k^2 r^2 - l(l + 1)) R = 0.$$

Ecrivant $x = kr$ et $y(x) = R(x/k)$, l'équation est alors

$$\frac{d}{dx} \left(x^2 \frac{dy}{dx} \right) + (x^2 - l(l + 1)) y = 0$$

qui porte le nom d'équation de Bessel sphérique.

Remarques. 1. Nous venons d'introduire beaucoup de nouvelles équations que nous avons baptisées des noms des mathématiciens qui les ont étudiés. Les solutions de certaines de ces équations sont caractérisées; tout comme $y'' = -\lambda y$ donne lieu aux solutions $\sin mx$ et $\cos mx$, $m \in \mathbb{N}$, les équations introduites donneront lieu à des familles de solutions à l'aide desquelles nous pourrions faire des séries de Fourier (ou plutôt des séries de Legendre, de Bessel, etc.). La tradition a donné le nom de "fonctions spéciales" aux fonctions de ces familles. Ce mot "spécial" ne veut pas dire grand chose si ce n'est que ces fonctions sont très utiles. En ce sens, les fonctions $\sin$, $\cos$ et $\exp$ sont aussi des fonctions spéciales.

2. Dans les exemples ci-dessus, on voit comment la séparation des variables dépend de la forme du laplacien dans les coordonnées étudiées. Cette séparation est assez délicate et on peut se demander s'il existe beaucoup d'autres systèmes de coordonnées dans lesquels l'équation d'Helmholtz se sépare. Par exemple un système de coordonnées (u, v) pour lequel le laplacien aurait la forme

$$\nabla^2 \psi = f(u, v) \frac{\partial^2 \psi}{\partial u^2} + g(u, v) \frac{\partial^2 \psi}{\partial u \partial v} + h(u, v) \frac{\partial^2 \psi}{\partial v^2}$$

équation de Legendre
associée
Adrien-Marie Legendre
(1752-1833)

équation de Legendre

équation de Bessel
sphérique

ne nous permettrait pas de jouer le jeu de la séparabilité à cause du terme

$$\frac{\partial^2 \psi}{\partial u \partial v}.$$

L'étude de la séparabilité d'équations différentielles est reliée au *groupe de symétrie* de cette équation. La théorie de Lie associée permet de classer les systèmes dits séparables.¹ Pour l'équation de Helmholtz, on les connaît tous; en deux dimensions, cette équation se sépare en quatre systèmes (les cartésien, polaire, parabolique et elliptique) et en trois dimensions en onze systèmes. Ces systèmes sont d'une grande utilité en mathématiques appliquées et en physique.

6. Vérifier que les laplaciens tri-dimensionnels en coordonnées cylindriques et sphériques sont bien **Exercices**

$$\nabla^2 f = \frac{\partial^2 f}{\partial r^2} + \frac{1}{r} \frac{\partial f}{\partial r} + \frac{1}{r^2} \frac{\partial^2 f}{\partial \theta^2} + \frac{\partial^2 f}{\partial z^2} \quad \text{cylindrique}$$

$$\nabla^2 f = \frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\partial f}{\partial r} \right) + \frac{1}{r^2 \sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial f}{\partial \theta} \right) + \frac{1}{r^2 \sin^2 \theta} \frac{\partial^2 f}{\partial \phi^2} \quad \text{sphérique}$$

7*. L'équation d'Helmholtz est séparable dans les 4 systèmes suivants

$$\begin{array}{ll} \text{cartésien} & \begin{cases} x \\ y \end{cases} \\ \text{polaire} & \begin{cases} x = r \cos \theta \\ y = r \sin \theta \end{cases} \\ \text{parabolique} & \begin{cases} x = \frac{1}{2}(\xi^2 - \eta^2) \\ y = \xi \eta \end{cases} \\ \text{elliptique} & \begin{cases} x = d \cosh \alpha \cos \beta \\ y = d \sinh \alpha \sin \beta \end{cases} \quad \text{où } d > 0. \end{array}$$

(a) Montrer que le laplacien $\nabla^2 f$ qui se lit

$$\nabla^2 f = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2}$$

1. Une référence sur la classification des systèmes séparables pour les équations de la physique mathématique est : Miller, W., *Symmetry and Separation of Variables*, Addison-Wesley (1977).

en coordonnées cartésiennes prend la forme suivante dans les autres systèmes :

$$\begin{aligned} \text{polaire} \quad \nabla^2 f &= \frac{\partial^2 f}{\partial r^2} + \frac{1}{r} \frac{\partial f}{\partial r} + \frac{1}{r^2} \frac{\partial^2 f}{\partial \theta^2} \\ \text{parabolique} \quad \nabla^2 f &= \frac{1}{\xi^2 + \eta^2} \left(\frac{\partial^2 f}{\partial \xi^2} + \frac{\partial^2 f}{\partial \eta^2} \right) \\ \text{elliptique} \quad \nabla^2 f &= \frac{1}{d^2(\cosh^2 \alpha - \cos^2 \beta)} \left(\frac{\partial^2 f}{\partial \alpha^2} + \frac{\partial^2 f}{\partial \beta^2} \right) \end{aligned}$$

(b) Trouver les équations différentielles ordinaires qui découlent de la séparation de l'équation d'Helmholtz dans ces quatre systèmes de coordonnées.

équation de Laplace
Pierre-Simon de Laplace
(1749-1827)

8. (a) Montrer que, si la fonction f satisfait l'équation de Laplace

$$\frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} = 0,$$

la fonction $g(x, y) = f(x^3 - 3xy^2, 3x^2y - y^3)$ en est aussi une solution.

(b) Montrer que l'équation d'Helmholtz n'est pas séparable en termes des variables $u(x, y) = x^3 - 3xy^2$ et $v(x, y) = 3x^2y - y^3$.

3.3 Le problème de Sturm-Liouville

Les équations que nous venons d'introduire dans la section précédente (Legendre, Legendre associé, Bessel, Bessel sphérique, et l'équation $y'' = -\lambda y$), peuvent toutes être mises sous la forme :

$$\frac{d}{dx} \left(p(x) \frac{dy}{dx} \right) - s(x)y + \lambda r(x)y = 0.$$

Cette équation porte le nom d'équation de Sturm-Liouville et un opérateur linéaire lui est naturellement associé :

$$\mathcal{L} = \frac{1}{r(x)} \left(\frac{d}{dx} \left[p(x) \frac{d}{dx} \right] - s(x) \right).$$

On doit comprendre cet opérateur comme agissant sur une fonction f de la façon suivante :

$$\mathcal{L}(f)(x) = \frac{1}{r(x)} \left(\frac{d}{dx} \left[p(x) \frac{df}{dx}(x) \right] - s(x)f(x) \right).$$

équation
de Sturm-Liouville

opérateur
de Sturm-Liouville

Charles Sturm
(1803-1855)
Joseph Liouville
(1809-1882)

Cette définition est quelque peu sommaire. Pour correctement définir une transformation linéaire (en algèbre linéaire), il faut toujours, en plus de donner la règle de transformation, donner l'espace vectoriel de départ (le domaine) et l'espace vectoriel d'arrivée (qui contient l'image). Quels sont les espaces qui permettent de correctement définir l'opérateur de Sturm-Liouville ? Ces espaces sont

donnés par le problème considéré. Nous avons répondu à cette question implicitement quand nous avons étudié la corde vibrante. Dans ce cas, l'opérateur était

$$\mathcal{L} = \frac{d^2}{dx^2}$$

et agissait sur les fonctions qui sont dérivables deux fois (sinon l'opérateur n'a pas de sens) et qui s'annulent aux extrémités de l'intervalle $[0, L]$. Cette dernière condition est d'origine physique et représente le fait que la corde est maintenue fixe en ses extrémités. Il est facile de vérifier que les fonctions dérivables deux fois et qui s'annulent en $x = 0$ et L forment un espace vectoriel. L'espace d'arrivée peut être beaucoup plus grand, par exemple l'ensemble des fonctions définies sur l'intervalle $[0, L]$. Cependant, dans les lignes qui suivent nous demanderons également que ces fonctions soient intégrables, une propriété qui sera vérifiée si le domaine contient des fonctions dont la seconde dérivée est continue.

Le but de cette section est d'étendre les quelques propriétés que nous avons observées dans l'étude de la corde vibrante aux opérateurs de Sturm-Liouville. Plusieurs des équations différentielles linéaires ont cette forme. Beaucoup viennent de la physique mathématique. Dans ce qui suit, nous allons suivre les étapes suivantes :

- une propriété élémentaire de $\mathcal{L}$: la linéarité
- les conditions aux limites : rôle et certains types
- $\mathcal{L}$ est un opérateur auto-adjoint
- « l'équation de Sturm-Liouville + conditions aux limites » = un problème aux valeurs propres

$\mathcal{L}$ est linéaire. L'opérateur de Sturm-Liouville est un opérateur *linéaire* sur un espace de fonctions. Rappelons qu'une transformation $T : V \rightarrow W$ est linéaire si pour tout $v_1, v_2 \in V$ et $c \in \mathbb{R}$

$$T(v_1 + cv_2) = T(v_1) + cT(v_2).$$

Dans le cas présent, la linéarité est immédiate :

$$\begin{aligned} \mathcal{L}(f_1 + cf_2) &= \frac{1}{r} \left[\frac{d}{dx} \left(p \frac{d}{dx} (f_1 + cf_2) \right) - s(f_1 + cf_2) \right] \\ &= \frac{1}{r} \left[\frac{d}{dx} \left(p \frac{df_1}{dx} \right) - sf_1 \right] + c \frac{1}{r} \left[\frac{d}{dx} \left(p \frac{df_2}{dx} \right) - sf_2 \right] \\ &= \mathcal{L}(f_1) + c\mathcal{L}(f_2). \end{aligned}$$

Conditions aux limites. Comme nous l'avons vu sur l'exemple de la corde vibrante, les solutions que nous cherchons aux équations différentielles sont soumises à des conditions aux limites, c'est-à-dire des contraintes auxquelles doit satisfaire la solution en les extrémités de l'intervalle sur lequel la solution est cherchée. (Les textes en anglais disent « *boundary conditions* » que l'on traduit souvent en français

conditions
aux limites

par « conditions frontières ». Les textes écrits s'en tiennent cependant aux mots « condition aux limites ».) C'est sûrement le cas de l'équation trigonométrique. Si l'intervalle définissant le domaine des solutions recherchées est $[a, b]$, les conditions aux limites les plus usuelles sont :

- Dirichlet : $f(a) = f(b) = 0$.
- Neumann : $f'(a) = f'(b) = 0$.
- intermédiaire : $f(a) + \alpha f'(a) = 0$ et $f(b) + \beta f'(b) = 0$ pour certains nombres α, β constants.
- $p(a) = p(b) = 0$.
- mixte : un des quatre types ci-dessus en a et un autre en b .

Il est donc possible de définir l'espace vectoriel sur lequel agit $\mathcal{L}$. Par exemple, pour les conditions de Dirichlet, il s'agit des fonctions dérivables continûment deux fois et qui s'annulent aux extrémités. La vérification de la structure d'espace vectoriel de ces ensembles est laissée en exercice. (Pour le cas $p(a) = p(b) = 0$, la condition sur la fonction est moins évidente. Voir un exemple à l'exercice 5.)

$\mathcal{L}$ est un opérateur auto-adjoint. Rappelons qu'une matrice X est symétrique si $X^T = X$, c'est-à-dire que sa transposée est égale à elle-même. Si $(x, y) = x^T y$ dénote le produit scalaire usuel sur $\mathbb{R}^n$, alors $(x, Xy) = x^T Xy = x^T X^T y = (Xx)^T y = (Xx, y)$. Si $x, y \in \mathbb{C}^n$ et X est une matrice hermitienne ($X^\dagger = X$), on vérifie de la même façon que

$$(x, Xy) = (Xx, y).$$

Une propriété semblable existe pour les opérateurs définis sur les espaces de fonctions.

Définition 17. Soit $\mathcal{C}^2([a, b])$ l'ensemble des fonctions (réelles) deux fois continûment différentiables sur l'intervalle $[a, b]$ et soit r une fonction positive et intégrable au sens de Riemann sur $[a, b]$. Pour $f, g \in \mathcal{C}^2([a, b])$, on définit le produit scalaire

$$(f, g) = \int_a^b r(x) f(x) g(x) dx.$$

opérateur auto-adjoint

L'opérateur $\mathcal{L}$ est dit auto-adjoint si :

$$(f, \mathcal{L}g) = (\mathcal{L}f, g)$$

pour toute paire de fonctions $f, g \in \mathcal{C}^2([a, b])$. L'ajout d'une condition aux limites peut restreindre l'espace vectoriel considéré et la contrainte $(f, \mathcal{L}g) = (\mathcal{L}f, g)$ sera alors restreinte aux fonctions de ce sous-espace.

(Exercice : Vérifier que $(\cdot, \cdot)$ est un produit scalaire.)

Dans notre cas, l'opérateur de Sturm-Liouville est auto-adjoint pour les conditions aux limites de Dirichlet, de Neumann, intermédiaires ou lorsque $p(a) =$

$p(b) = 0$:

$$\begin{aligned}
 (f, \mathcal{L}g) &= \int_a^b r f \frac{1}{r} \left(\frac{d}{dx} \left(p \frac{dg}{dx} \right) - sg \right) dx \\
 &= \underbrace{f p \frac{dg}{dx} \Big|_a^b}_{=0} - \int_a^b p \frac{df}{dx} \frac{dg}{dx} dx - \int_a^b s f g dx \\
 &= \underbrace{-p \frac{df}{dx} g \Big|_a^b}_{=0} + \int_a^b \frac{d}{dx} \left(p \frac{df}{dx} \right) g dx - \int_a^b s f g dx \\
 &= \int_a^b r \frac{1}{r} \left(\frac{d}{dx} \left(p \frac{df}{dx} \right) - sf \right) g dx \\
 &= \int_a^b r (\mathcal{L}f) g dx \\
 &= (\mathcal{L}f, g)
 \end{aligned}$$

Nous avons observé que la matrice caractérisant la corde plombée était symétrique. Similairement l'opérateur qui représente le problème équivalent pour la corde vibrante (et, en général, pour le problème de Sturm-Liouville) est auto-adjoint.

Le problème de Sturm-Liouville est un problème aux valeurs propres. L'équation de Sturm-Liouville est

$$\frac{d}{dx} \left(p(x) \frac{dy}{dx} \right) - s(x)y + \lambda r(x)y = 0.$$

Cette équation peut être réécrite à l'aide de l'opérateur de Sturm-Liouville sous la forme :

$$\mathcal{L}(f) = -\lambda f.$$

Si une condition aux limites est spécifiée, par exemple dans la liste ci-dessus, cette équation peut donc être interprétée comme une équation aux valeurs propres puisqu'alors $\mathcal{L}$ est un opérateur linéaire sur un espace vectoriel bien défini. Cette analogie entre algèbre linéaire et analyse va même plus loin !

Si le problème de Sturm-Liouville est un problème aux valeurs propres et l'opérateur linéaire est auto-adjoint, peut-on en conclure, comme on le fait pour les matrices symétriques, que les fonctions propres (vecteurs propres) correspondant à des valeurs propres distinctes sont orthogonales ? Et est-ce que les valeurs propres sont réelles ? En fait il suffit de revoir les preuves d'algèbre linéaire pour nous en convaincre. Nous allons tout d'abord montrer que, si f_m et f_n sont des fonctions propres de $\mathcal{L}$ avec valeurs propres λ_m et λ_n distinctes :

$$\begin{aligned}
 \mathcal{L}f_m &= -\lambda_m f_m \\
 \mathcal{L}f_n &= -\lambda_n f_n,
 \end{aligned}$$

alors elles sont orthogonales par rapport au produit scalaire :

$$(f, g) = \int_a^b r(x) f(x) g(x) dx.$$

Puisque $\mathcal{L}$ est auto-adjoint,

$$\lambda_m f_m = -(\mathcal{L} f_m, f_n) = -(f_m, \mathcal{L} f_n) = \lambda_n f_n$$

et donc $(\lambda_m - \lambda_n)(f_m, f_n) = 0$. Puisque les valeurs propres λ_m et λ_n sont distinctes, il faut que $(f_m, f_n) = 0$, c'est-à-dire que les fonctions f_m et f_n soient orthogonales pour le produit scalaire $(\cdot, \cdot)$. La dernière propriété observée lors de la solution de la corde plombée a donc son pendant ici également : les modes propres étaient orthogonaux et ici, les fonctions propres sont orthogonales.

Remarque : Contrairement au cas trigonométrique où nous avons utilisé les fonctions exponentielles pour étudier la convergence en moyenne, les exemples qui suivront seront pour des fonctions à valeurs réelles, c'est-à-dire $f : \mathbb{R} \rightarrow \mathbb{R}$. Vous pouvez retourner à la section 2.5 et vous assurer que, si r est positive sur l'intervalle $[a, b]$, le produit $(\cdot, \cdot)$ que nous venons de définir est un produit scalaire et satisfait l'inégalité de Cauchy-Schwarz-Buniakowski. (Parfois, il suffit d'absorber un signe dans la valeur propre pour rendre $r(x)$ positif.)

Dans ce qui suit nous allons faire les hypothèses suivantes :

- il y a une infinité de fonctions propres y_n de $\mathcal{L}$ dans $\mathcal{C}^2([a, b])$;
- chacun des sous-espaces propres est de dimension finie. Nous supposons les y_n orthogonales. Cette hypothèse n'a de l'importance que si la valeur propre associée λ_n est dégénérée ;
- toute fonction dans $\mathcal{C}^2([a, b])$ peut être représentée par une série de Fourier

$$f(x) = \sum_{n=1}^{\infty} a_n y_n, \quad \forall x \in [a, b].$$

(Il y a « assez » de fonctions propres.) Nous supposons de plus que cette convergence est uniforme et en tous les points de l'intervalle.

Les coefficients a_n de ce développement peuvent être aisément calculés :

$$\begin{aligned} \int_a^b r(x) f(x) y_m(x) dx &= \sum_{n=1}^{\infty} \int_a^b a_n r(x) y_n(x) y_m(x) dx \\ &= \int_a^b a_m r(x) y_m(x)^2 dx \end{aligned}$$

car y_n et y_m sont orthogonales et donc, de tous les termes de la somme, seul le terme $n = m$ demeure. Donc :

$$a_m = \frac{\int_a^b r(x) f(x) y_m(x) dx}{\int_a^b r(x) y_m(x)^2 dx} = \frac{(f, y_m)}{(y_m, y_m)}.$$

9. Vérifier que les ensembles des fonctions continûment différentiables deux fois et satisfaisant les conditions aux limites de Dirichlet, Neumann et intermédiaires définissent trois espaces vectoriels. **Exercices**

10*. Est-ce que la condition aux limites $f(a) = f(b) = 1$ définit un espace vectoriel ?

11*. Lesquels des opérateurs différentiels suivants sont linéaires ?

(a)

$$\mathcal{T}(f) = \frac{d^3 f}{dx^3}$$

(b)

$$\mathcal{T}(f) = x \frac{df}{dx} - f$$

(c)

$$\mathcal{T}(f) = f$$

(d)

$$\mathcal{T}(f) = f^2 - f$$

(e)

$$\mathcal{T}(f) = f \frac{df}{dx}$$

(f)

$$\mathcal{T}(f) = \frac{d(xf - f')}{dx}$$

(g)

$$\mathcal{T}(f) = \int_0^1 f(x) dx - \frac{df}{dx}$$

12*. (a) Déterminer les fonctions $p(x)$, $q(x)$ et $r(x)$ pour les équations de Legendre, Legendre associée, Bessel et Bessel sphérique.

(b) Pour les changements de variables que nous avons utilisés pour séparer l'équation d'Helmholtz, identifier le domaine des fonctions solutionnant les équations ci-dessus.

(c) Ecrire la forme bilinéaire par rapport à laquelle $\mathcal{L}$ est auto-adjoint pour ces mêmes équations.

13*. L'équation d'Hermite est la suivante :

$$e^{x^2} \frac{d}{dx} \left(e^{-x^2} \frac{dy}{dx} \right) + 2ny = 0.$$

(a) Quelles sont les fonctions $p(x)$, $q(x)$ et $r(x)$ du problème de Sturm-Liouville associé ? ($2n$ est ici la valeur propre.)

(b) Le domaine des solutions de l'équation d'Hermite est $(-\infty, +\infty)$. La preuve d'orthogonalité des solutions propres utilisait un intervalle borné (a, b) . Refaire la preuve en l'adaptant au cas présent. Pour répondre à cette question, vous devrez décrire l'espace vectoriel sur lequel agit l'opérateur $\mathcal{L}$ associé ? (Pour cette partie, vous devrez réfléchir.)

14*. Soit le problème de Sturm-Liouville

$$y'' + \lambda y = 0$$

où y vérifie les conditions aux limites $y(0) = 0$ et $y'(L) + \beta y(L) = 0$. (β est une constante réelle positive. On considèrera le cas de valeurs propres λ positives.)

(a) Trouver les fonctions propres du problème de Sturm-Liouville et montrer que les valeurs propres satisfont $\sqrt{\lambda} = -\beta \tan(\sqrt{\lambda}L)$.

(b) Montrer qu'il existe précisément une valeur propre λ_n dans chacun des intervalles

$$\lambda_n \in [(n - \frac{1}{2})^2 \pi^2 / L^2, (n + \frac{1}{2})^2 \pi^2 / L^2], n \in \mathbb{N}.$$

(c) Trouver un ensemble $\{\phi_n, n \in \mathbb{N}\}$ de fonctions propres telles que $(\phi_n, \phi_m) = \delta_{mn}$.

15. Soit $\{y_n, n \in \mathbb{N}\}$ l'ensemble des fonctions propres du problème de Sturm-Liouville $\mathcal{L}f = \lambda f$ plus certaines conditions aux limites. En supposant que la série

$$f(x) = \sum_{n=1}^{\infty} a_n y_n, \quad \forall x \in [a, b],$$

identité de Parseval

converge uniformément, démontrer l'identité de Parseval :

$$(f, f) = \sum_{n=1}^{\infty} a_n^2(y_n, y_n).$$

16. Voici des équations de Sturm-Liouville classiques dont certaines des fonctions propres sont des polynômes lorsque $n \in \mathbb{N} \cup \{0\}$.

nom	équation différentielle	intervalle
Legendre	$\frac{d}{dx} \left((1-x^2) \frac{df}{dx} \right) + n(n+1)f = 0$	$(-1, 1)$
Chebyshev	$\sqrt{1-x^2} \frac{d}{dx} \left(\sqrt{1-x^2} \frac{df}{dx} \right) + n^2 f = 0$	$(-1, 1)$
Hermite	$e^{x^2} \frac{d}{dx} \left(e^{-x^2} \frac{df}{dx} \right) + 2nf = 0$	$(-\infty, +\infty)$
Laguerre	$x^{-\alpha} e^x \frac{d}{dx} \left(x^{\alpha+1} e^{-x} \frac{df}{dx} \right) + nf = 0$	$(0, \infty)$

(a) Sachant que le polynôme propre associé à la valeur n est de degré n , trouver les trois premiers polynômes p_0, p_1 et p_2 pour ces équations.

(b) Par rapport à quelle forme bilinéaire les fonctions propres sont-elles orthogonales ?

(c) Vérifier que ces premiers polynômes sont mutuellement orthogonaux.

17*. Quelques polynômes de Chebyshev ont été construits à l'exercice précédent. Ils jouent un rôle central en analyse numérique car ils peuvent être utilisés efficacement pour approximer les fonctions définies sur $[-1, 1]$, ou plus généralement sur $[a, b]$. En plus de la relation d'orthogonalité par rapport à $(\cdot, \cdot)$ que vous avez décrite ci-dessus, les m premiers polynômes de Chebyshev T_i , $0 \leq i \leq m-1$, vérifient également une relation d'orthogonalité *discrète* : $\langle T_i, T_j \rangle_m = 0$ si $i \neq j$ et $i, j < m$. La forme $\langle \cdot, \cdot \rangle_m$ est définie par

$$\langle f, g \rangle_m = \sum_{k=1}^m f(x_k)g(x_k),$$

où les nombres $x_k, 1 \leq k \leq m$, sont les m zéros de T_m . Montrer que $\langle \cdot, \cdot \rangle_m$ est un produit scalaire sur l'espace des polynômes de degré égal ou inférieur à $(m-1)$.
Note : Prendre pour acquis que le polynôme de Chebyshev $T_i(x)$ a précisément i zéros réels tous distincts. L'exercice ne demande pas de montrer que les T_i et T_j sont orthogonaux par rapport à ce produit.

3.4 Polynômes de Legendre

Les deux dernières sections de ce chapitre sont consacrées à deux familles de fonctions spéciales ; les exercices en couvrent quelques-unes de plus. Il est utile de

faire le point avant de s'y engager. Deux idées maîtresses ont été découvertes au premier chapitre :

- une série de Fourier peut être mise en parallèle avec un vecteur écrit dans une base orthogonale comme en algèbre linéaire ;
- la séparation de variables permet d'écrire la solution de l'équation de la corde vibrante et du tambour (carré) comme une série de Fourier.

Or, depuis le début de ce chapitre, nous découvrons que les outils construits au chapitre 1 pourraient bien s'étendre aisément à une grande variété d'autres problèmes. En effet, nous avons montré que

- la séparation de variables peut être utilisée pour d'autres systèmes de coordonnées et d'autres équations ;
- les équations différentielles ordinaires que nous obtenons de ces séparations peuvent être réinterprétées, à l'aide du « truc » de Sturm-Liouville, comme des problèmes aux valeurs propres dont les fonctions propres sont orthogonales.

Ces résultats indiquent que l'interprétation du premier chapitre pourra fort probablement être étendue à ces nouvelles équations différentielles ordinaires et que plusieurs équations aux dérivées partielles linéaires pourront être résolues de la sorte.

Les deux prochaines sections auront le même but : trouver les fonctions propres d'un opérateur de Sturm-Liouville obtenu par séparation d'une équation aux dérivées partielles. Nous avons choisi l'équation de Legendre et celle de Bessel. Les deux auront un traitement légèrement différent car, dans le premier cas, les solutions propres que nous retiendrons seront des polynômes alors que, dans le second, ce n'en seront pas.

Quelque soit le type de solutions cherchées, les fonctions « spéciales » sont souvent étudiées à l'aide d'outils classiques. Certains d'entre eux vous sont probablement inconnus : fonction génératrice, formule de Rodrigues, relations de récurrence ; alors que d'autres vous sont familiers : valeurs des fonctions propres en des points spécifiques, orthogonalité et normalisation des fonctions propres, identité de Parseval. Nous les introduirons en discutant l'équation de Legendre.

L'équation de Legendre a été obtenue par séparation de l'équation de Helmholtz $\nabla^2 f = -k^2 f$ en coordonnées sphériques :

$$\frac{d}{dx} \left((1-x^2) \frac{dy}{dx} \right) = \lambda_2 y.$$

Elle représente la partie $\Theta(\theta)$ de la solution séparée $f(r, \theta, \phi) = R(r)\Theta(\theta)\Phi(\phi)$. Le changement de variables $y(x) = \Theta(\arccos x)$ relie y et Θ . Nous sommes donc intéressés à des solutions sur l'intervalle $[-1, 1]$ puisque $\arccos x$ n'a de sens que si $x \in [-1, 1]$. La section 3.2 a montré de plus que ce n'est pas l'équation la plus générale mais que la constante de séparation m a été posée à zéro pour l'obtenir. Puisque m représente le comportement de la fonction Φ

$$\Phi'' = -m^2 \Phi,$$

alors ce choix force un comportement azimutal trivial :

$$\Phi(\phi) = \text{constante.}$$

Ainsi seules les solutions de l'équation de Helmholtz avec comportement azimutal trivial pourront être écrites à l'aide des polynômes de Legendre.

La stratégie que nous emploierons pour définir les polynômes de Legendre serait due à Legendre lui-même. Elle repose sur l'idée suivante. Soit f une solution de l'équation de Laplace : $\nabla^2 f = 0$. Avec les hypothèses de complétude faites à la fin de la section précédente, il est possible d'écrire f comme une combinaison linéaire de solutions séparées $R(r)\Theta(\theta)\Phi(\phi)$ ou, si la dépendance azimutale est triviale, comme $f = \sum_n a_n R_n(r)\Theta_n(\theta)$ où n étiquette les valeurs possibles de la constante de séparation λ_2 . Et, pour un choix générique de f , un nombre infini des valeurs possibles de λ_2 apparaîtra dans la somme. (Rappelons que dans les exemples et exercices du chapitre 1, les séries de Fourier contenaient tous les $\sin mx$ et/ou tous les $\cos mx$.) Si les R_n sont connus, cette méthode permet donc de définir les fonctions Θ_n à une constante près.

Pour fixer les polynômes de Legendre, nous allons utiliser une solution connue de l'équation de Laplace, c'est-à-dire l'équation de Helmholtz avec $k = 0$:

$$\nabla^2 f = 0.$$

Conformément au paragraphe ci-dessus, cette solution aura un comportement azimutal trivial. Elle est :

$$f(r, \theta, \phi) = \frac{1}{\sqrt{1 - 2r \cos \theta + r^2}}.$$

En développant cette solution sous la forme

$$f(r, \theta, \phi) = \sum R_n(r)\Theta_n(\theta),$$

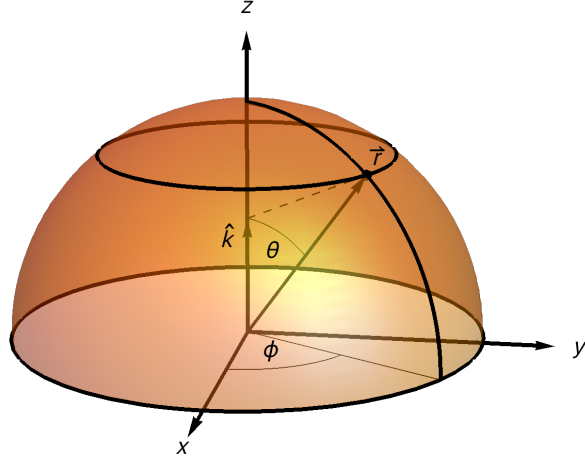
nous identifierons la partie Θ_n avec les polynômes de Legendre. Mais d'abord, qu'est-ce que cette fonction f ? C'est le potentiel électrique d'une particule ponctuelle située au point $\hat{k} = (0, 0, 1)$ le long de l'axe des z . Le potentiel électrique ne dépend que de la distance et donc

$$f(r, \theta, \phi) = \frac{1}{|\vec{r} - \hat{k}|}$$

et si on se reporte au graphique ci-contre :

$$\begin{aligned} |\vec{r} - \hat{k}| &= \sqrt{(\vec{r} - \hat{k}, \vec{r} - \hat{k})} \\ &= \sqrt{1 - 2r \cos \theta + r^2} \end{aligned}$$

Tel qu'annoncé, cette fonction ne dépend pas de ϕ . Mais est-ce que cette fonction f vérifie l'équation de Laplace ? Oui et c'est un court exercice que de le

Figure 3.1 – Les vecteurs $\vec{r}$ et $\hat{k}$ sont indiqués en gras, le vecteur $\vec{r} - \hat{k}$ en pointillé.

vérifier. (Deux façons de le faire : (1) utiliser le laplacien en coordonnées sphériques pour calculer $\nabla^2 f$ par force brute ou (2) remarquer que cette fonction est une translation de la fonction $f(r, \theta, \phi) = \frac{1}{r}$ et calculer le laplacien en coordonnées cartésiennes de $\frac{1}{r}$. Notons que les laplaciens d'une fonction ou de sa translatée sont égaux aux points correspondants.)

Puisque $k = 0$ (constante dans l'équation de Helmholtz) et $m = 0$ (dépendance azimutale triviale), les deux équations séparées qui demeurent sont :

$$\begin{aligned} \text{dépendance en } \theta & \quad \frac{d}{dx} \left((1 - x^2) \frac{dy}{dx} \right) - \lambda y = 0, \\ \text{dépendance radiale (en } r) & \quad \frac{d}{dr} \left(r^2 \frac{dR}{dr} \right) + \lambda R = 0. \end{aligned}$$

La prochaine étape est de restreindre les valeurs possibles de λ . Ceci utilise un résultat sur les équations différentielles ordinaires que nous citerons sans preuve. Si λ est quelconque, les solutions divergent ou possèdent un point de branchement en $x = +1$ ou $x = -1$ et donc en $\theta = 0$ et π *sauf* lorsque $\lambda = -l(l+1)$ avec l un entier non-négatif. Puisque f est analytique en $\theta = 0$ et π , il faudra donc se restreindre à ces valeurs :

$$\lambda = -l(l+1), \quad l \in \mathbb{N} \cup \{0\}.$$

L'équation pour R a alors deux solutions pour chacune des valeurs de l :

$$R_{l,1}(r) = r^l \quad \text{et} \quad R_{l,2} = \frac{1}{r^{l+1}}$$

comme on peut le vérifier directement pour $R_{l,1}$

$$\begin{aligned} \frac{d}{dr} \left(r^2 \frac{d}{dr} r^l \right) - l(l+1)r^l &= \frac{d}{dr} (r^2 l r^{l-1}) - l(l+1)r^l \\ &= l(l+1)r^l - l(l+1)r^l \\ &= 0 \end{aligned}$$

et pour $R_{l,2}$:

$$\begin{aligned} \frac{d}{dr} \left(r^2 \frac{d}{dr} \frac{1}{r^{l+1}} \right) - l(l+1) \frac{1}{r^{l+1}} &= - \frac{d}{dr} (r^2 (l+1) r^{-(l+2)}) - l(l+1) r^{-(l+1)} \\ &= l(l+1) r^{-(l+1)} - l(l+1) r^{-(l+1)} \\ &= 0. \end{aligned}$$

Puisque l'équation pour R est une équation différentielle ordinaire du second ordre, il existe précisément deux solutions linéairement indépendantes et nous pouvons les choisir comme étant $R_{l,1}$ et $R_{l,2}$. Et puisque le développement du potentiel f doit être régulier en r , nous ne conserverons que les r^l . (Ceci ressemble beaucoup à une condition aux limites : la fonction R ne doit pas avoir de singularité en $r = 0$.) Le développement sera donc de la forme

$$f(r, \theta, \phi) = \sum_{l=0}^{\infty} a_l r^l y_l(x), \quad \text{où } x = \cos \theta$$

et où y_l sont les solutions à l'équation de Legendre. Nous allons définir les polynômes de Legendre $P_l(x)$ comme les solutions de l'équation de Legendre qui font que toutes les constantes a_l du développement ci-dessus sont précisément égales à 1 :

$$= \sum_{l=0}^{\infty} r^l P_l(x).$$

Peut-on calculer les $P_l(x)$ avec une telle définition ? ! Oui, à l'aide de l'expression de f . Il s'agit de faire un développement en série de Taylor autour de $r = 0$ de la fonction f vue comme une fonction de la seule variable r . Alors les coefficients dépendront de $x = \cos \theta$ et ces coefficients seront les polynômes de Legendre. Puisque nous devons écrire

$$f(r, x) = \frac{1}{\sqrt{1 - 2rx + r^2}}$$

en série de r pour r petit, commençons par faire le développement de $g(\epsilon) = \frac{1}{\sqrt{1-\epsilon}}$

autour de $\epsilon = r(2x - r) = 0$. Mécaniquement, on obtient

$$\begin{aligned}
 g(\epsilon) &= (1 - \epsilon)^{-\frac{1}{2}}, & g(0) &= 1 \\
 g'(\epsilon) &= \frac{1}{2}(1 - \epsilon)^{-\frac{3}{2}}, & g'(0) &= \frac{1}{2} \\
 g''(\epsilon) &= \frac{1}{2} \frac{3}{2}(1 - \epsilon)^{-\frac{5}{2}}, & g''(0) &= \frac{1}{2} \frac{3}{2} \\
 g'''(\epsilon) &= \frac{1}{2} \frac{3}{2} \frac{5}{2}(1 - \epsilon)^{-\frac{7}{2}}, & g'''(0) &= \frac{1}{2} \frac{3}{2} \frac{5}{2} \\
 &\dots & &\dots \\
 g^{(n)}(\epsilon) &= \frac{1}{2} \frac{3}{2} \dots \frac{(2n-1)}{2} (1 - \epsilon)^{-\frac{(2n+1)}{2}}, & g^{(n)}(0) &= \frac{(2n)!}{2^{2n} n!}
 \end{aligned}$$

où la valeur de $g^{(n)}$ en zéro a été obtenue comme suit :

$$\begin{aligned}
 g^{(n)}(0) &= \frac{1}{2} \frac{3}{2} \dots \frac{(2n-1)}{2} \\
 &= \frac{1 \cdot 3 \cdot 5 \cdot \dots \cdot (2n-1)}{2^n} \\
 &= \frac{1 \cdot 2 \cdot 3 \cdot 4 \cdot \dots \cdot (2n-1) \cdot (2n)}{2^n \cdot 2 \cdot 4 \cdot \dots \cdot (2n)} \\
 &= \frac{(2n)!}{2^n 2^n n!} \\
 &= \frac{(2n)!}{2^{2n} n!}.
 \end{aligned}$$

Il est maintenant aisé d'obtenir le développement en série de Taylor de f :

$$\begin{aligned}
 f(r, x) &= \sum_{n=0}^{\infty} \frac{(2n)!}{2^{2n} n!} \frac{r^n (2x - r)^n}{n!}, \quad \text{développement autour de } \epsilon = r(2x - r) = 0 \\
 &= \sum_{n=0}^{\infty} \frac{(2n)!}{2^{2n} (n!)^2} \sum_{i=0}^n r^n (-r)^i (2x)^{n-i} \frac{n!}{i!(n-i)!}, \quad \text{binôme !} \\
 &= \sum_{n=0}^{\infty} \sum_{i=0}^n \frac{(-1)^i (2n)! r^{n+i} x^{n-i}}{2^{2n+i-n} n! i!(n-i)!} \\
 &= \sum_{l=0}^{\infty} r^l \sum_{i=0}^{\left[\frac{l}{2}\right]} \frac{(-1)^i (2l-2i)! x^{l-2i}}{2^l (l-i)! i! (l-2i)!}, \quad (\text{voir ci-dessous})
 \end{aligned}$$

polynôme
de Legendre

et donc le polynôme de Legendre est

$$P_l(x) = \sum_{i=0}^{\left[\frac{l}{2}\right]} \frac{(-1)^i (2l-2i)! x^{l-2i}}{2^l (l-i)! i! (l-2i)!}. \quad (\Delta)$$

Le symbole $[x]$ signifie « partie entière de x », c'est-à-dire le plus grand entier plus petit ou égal à x . Avant d'aller plus loin, expliquons le tour de magie qui a permis

d'écrire l'égalité entre les deux doubles sommes :

$$\sum_{n=0}^{\infty} \sum_{i=0}^n \cdots = \sum_{l=0}^{\infty} \sum_{i=0}^{\lfloor \frac{l}{2} \rfloor} \cdots$$

Cette étape est fort semblable au changement de variables d'intégration sur un domaine à deux dimensions. Représentons chacun des termes dans la première somme comme un simple point. Chaque point se voit donc étiqueté par sa valeur en n et en i :

$i =$	0	1	2	3	4	...
$n = 0$	•					
1	•	•				
2	•	•	•			
3	•	•	•	•		
4	•	•	•	•	•	
5	•	•	•	•	•	•
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$

Chacun des points ci-dessus peut être étiqueté par les deux entiers :

$$l = n + i \quad \text{et} \quad i.$$

Ce choix nous est dicté par le fait que nous voulons identifier le développement en série de r et donc remplacer l'exposant de r qui est $n + i$ par un entier l . Avec ce choix de l , les « • » qui sont sur des diagonales de pente unité dans le tableau ci-dessous possèdent les mêmes l . Ainsi

- si $l = 0$ les « • » sont ceux avec $(n, i) = (0, 0)$
- si $l = 1$ les « • » sont ceux avec $(n, i) = (1, 0)$
- si $l = 2$ les « • » sont ceux avec $(n, i) = (2, 0), (1, 1)$
- si $l = 3$ les « • » sont ceux avec $(n, i) = (3, 0), (2, 1)$
- si $l = 4$ les « • » sont ceux avec $(n, i) = (4, 0), (3, 1), (2, 2)$
- etc.

Un moment de réflexion vous convaincra que le nombre de « • » pour un l donné est précisément $\lfloor \frac{l}{2} \rfloor + 1$, comme le laisse deviner les cinq premières valeurs de l ci-dessus. Ceci explique donc le passage de la double somme sur n et i à une double somme sur $l = n + i$ et i .

Puisqu'il est possible de calculer *tous* les polynômes de Legendre à partir de la seule fonction f , on appelle f la *fonction génératrice* des polynômes de Legendre. Celle-ci peut être utilisée pour obtenir plusieurs relations fondamentales des polynômes de Legendre.

Les premiers polynômes de Legendre sont tracés sur la figure 3.2. Chacun peut être identifié comme suit : P_l a l zéros sur l'intervalle $[-1, 1]$.

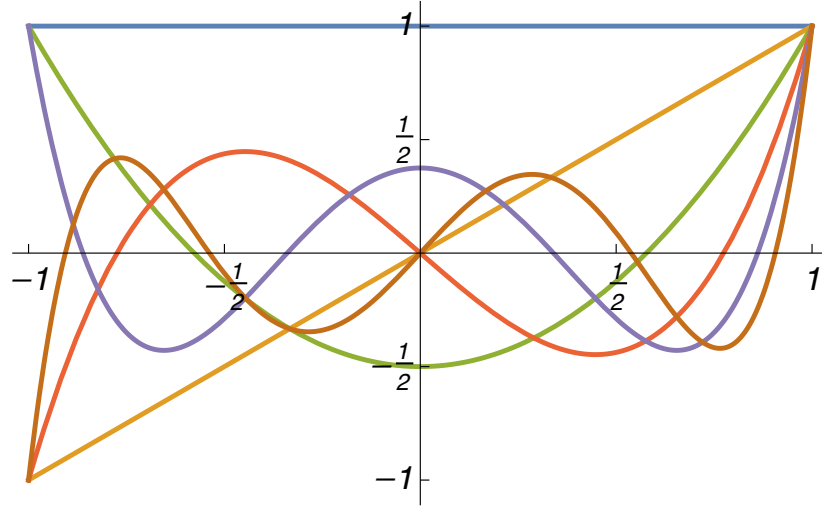


Figure 3.2 – Les six premiers polynômes de Legendre. Le polynôme P_l a l zéros dans l'intervalle $[-1, 1]$.

En dérivant f par rapport à r , on obtient d'une part :

$$\frac{\partial f}{\partial r} = \frac{\partial}{\partial r} \sum_{l=0}^{\infty} r^l P_l(x) = \sum_{l=0}^{\infty} l r^{l-1} P_l(x).$$

D'autre part la dérivée de l'expression explicite de f donne

$$\frac{\partial f}{\partial r} = -\frac{1}{2} \frac{(-2x + 2r)}{(\sqrt{1 - 2xr + r^2})^3} = \frac{x - r}{(\sqrt{1 - 2xr + r^2})^3} = \frac{x - r}{1 - 2xr + r^2} f.$$

Ainsi

$$(x - r)f = (1 - 2xr + r^2) \sum_{l=0}^{\infty} l r^{l-1} P_l$$

et donc

$$\sum_{l=0}^{\infty} x r^l P_l - \sum_{l=0}^{\infty} r^{l+1} P_l = \sum_{l=0}^{\infty} (l r^{l-1} P_l - 2x l r^l P_l + l r^{l+1} P_l)$$

ou encore

$$\sum_{l=0}^{\infty} x r^l P_l - \sum_{j=1}^{\infty} r^j P_{j-1} = \sum_{j=0}^{\infty} (j+1) r^j P_{j+1} - 2x \sum_{l=0}^{\infty} l r^l P_l + \sum_{j=1}^{\infty} (j-1) r^j P_{j-1}.$$

Les indices muets des sommes ont été translatés pour que tous les exposants des r soient précisément l'indice de la somme. En identifiant terme à terme les

coefficients d'une même puissance de r , les identités suivantes sont obtenues :

$$\begin{aligned} l = 0, & \quad xP_0 = P_1 \\ l \geq 1, & \quad xP_l - P_{l-1} = (l+1)P_{l+1} - 2lxP_l + (l-1)P_{l-1} \end{aligned}$$

et donc

$$(2l+1)xP_l = (l+1)P_{l+1} + lP_{l-1}, \quad \text{si } l \geq 0. \quad (R1) \quad \begin{array}{l} \text{première relation} \\ \text{de récurrence} \end{array}$$

Cette relation est appelée *relation de récurrence* pour les polynômes de Legendre.

Puisque

$$P_0 = 1 \quad \text{et} \quad P_1 = x$$

par la définition explicite (Δ), cette relation permet de calculer les autres. Par exemple, pour $l = 1$:

$$3xP_1 = 2P_2 + P_0$$

et donc

$$P_2 = \frac{1}{2}(3xP_1 - P_0) = \frac{3}{2}x^2 - \frac{1}{2}.$$

De façon similaire, la dérivée de la fonction génératrice f par rapport à x donne lieu à

$$\frac{\partial f}{\partial x} = \sum_{l=0}^{\infty} r^l P'_l(x)$$

et puisque

$$\frac{\partial f}{\partial x} = \frac{r}{(\sqrt{1-2xr+r^2})^3} = \frac{r}{1-2xr+r^2} f,$$

on a donc

$$\begin{aligned} \sum_{l=0}^{\infty} r^{l+1} P_l &= \sum_{l=0}^{\infty} (1-2xr+r^2) r^l P'_l(x) \\ &= \sum_{l=0}^{\infty} (r^l P'_l - 2xr^{l+1} P'_l + r^{l+2} P'_l). \end{aligned}$$

Effectuant les changements appropriés sur les indices de sommation :

$$\sum_{l=1}^{\infty} r^l (P_{l-1} + 2xP'_{l-1}) = \sum_{l=0}^{\infty} r^l P'_l + \sum_{l=2}^{\infty} r^l P'_{l-2}.$$

A nouveau, en regroupant les coefficients d'une même puissance de r :

$$\begin{aligned} l = 0 & \quad P'_0 = 0 \\ l = 1 & \quad P_0 + 2xP'_0 = P'_1, \quad \text{c'est-à-dire } P'_1 = P_0 \\ l \geq 2 & \quad P_{l-1} + 2xP'_{l-1} = P'_l + P'_{l-2} \end{aligned}$$

ou encore

$$P_l = P'_{l+1} - 2xP'_l + P'_{l-1}, \quad l \geq 1. \quad (R2) \quad \begin{array}{l} \text{seconde relation} \\ \text{de récurrence} \end{array}$$

Ceci est la seconde relation de récurrence pour les polynômes de Legendre.

Quelques valeurs spéciales des polynômes de Legendre peuvent être obtenues toujours à partir de la fonction génératrice. Par exemple

$$f(r, 1) = \frac{1}{\sqrt{1-2r+r^2}} = \frac{1}{1-r} = \sum_{l=0}^{\infty} r^l.$$

Puisque $f(r, 1) = \sum_{l=0}^{\infty} r^l P_l(1)$, on obtient

$$P_l(1) = 1, \quad \text{pour tout } l.$$

De même, en $x = -1$

$$f(r, -1) = \frac{1}{1+r} = \sum_{l=0}^{\infty} (-r)^l$$

et alors

$$P_l(-1) = (-1)^l, \quad \text{pour tout } l.$$

Ceci donne un exemple de valeurs particulières des polynômes de Legendre.

On peut écrire les polynômes de Legendre à l'aide de l'expression

$$P_l = \frac{1}{2^l l!} \frac{d^l}{dx^l} (x^2 - 1)^l.$$

Cette relation s'appelle la formule de Rodrigues pour les polynômes de Legendre. Sa preuve est directe

$$\begin{aligned} \frac{1}{2^l l!} \frac{d^l}{dx^l} (x^2 - 1)^l &= \frac{1}{2^l l!} \frac{d^l}{dx^l} \sum_{k=0}^l (-1)^k x^{2(l-k)} \frac{l!}{k!(l-k)!} \\ &= \frac{1}{2^l l!} \sum_k (-1)^k x^{2l-2k-l} \frac{(2l-2k)(2l-2k-1) \dots (2l-2k-l+1)l!}{k!(l-k)!} \end{aligned}$$

où la somme ne porte que sur les k tels que $2l-2k-l \geq 0$ (puisque la dérivée de x^0 est nulle) et donc sur k de 0 à $\lfloor \frac{l}{2} \rfloor$. Donc

$$\begin{aligned} &= \frac{1}{2^l l!} \sum_{k=0}^{\lfloor \frac{l}{2} \rfloor} (-1)^k x^{l-2k} \frac{(2l-2k)!!}{(l-2k)!k!(l-k)!} \\ &= \sum_{k=0}^{\lfloor \frac{l}{2} \rfloor} \frac{(-1)^k x^{l-2k} (2l-2k)!}{2^l (l-2k)!k!(l-k)!} \end{aligned}$$

qui est bien la relation de définition.

La fonction $r(x)$ qui, dans le problème de Sturm-Liouville, définit le produit scalaire est ici simplement la constante 1, alors que la fonction $p(x)$ est $(1-x^2)$. Le produit est donc

$$(f, g) = \int_{-1}^1 f(x)g(x)dx.$$

valeurs
particulières

formule de Rodrigues
Olinde Rodrigues
(1795-1851)

A nouveau, la variable $x = \cos \theta$ prend ses valeurs sur l'intervalle $[-1, 1]$ et, puisque la fonction p s'annule aux extrémités du domaine, $p(-1) = p(1) = 0$, l'opérateur de Sturm-Liouville est auto-adjoint et ses fonctions propres seront orthogonales. Nous offrons une autre preuve de l'orthogonalité, basée elle sur la formule de Rodrigues. En supposant $l \geq m$:

$$2^l l! 2^m m! (P_l, P_m) = \int_{-1}^1 \frac{d^l}{dx^l} (x^2 - 1)^l \frac{d^m}{dx^m} (x^2 - 1)^m dx$$

Intégrant par parties :

$$\begin{aligned} &= \frac{d^m}{dx^m} (x^2 - 1)^m \frac{d^{l-1}}{dx^{l-1}} (x^2 - 1)^l \Big|_{-1}^1 \\ &\quad - \int_{-1}^1 \frac{d^{l-1}}{dx^{l-1}} (x^2 - 1)^l \frac{d^{m+1}}{dx^{m+1}} (x^2 - 1)^m dx \end{aligned}$$

Remarquons que

$$\frac{d^{l-1}}{dx^{l-1}} (x^2 - 1)^l = \frac{d^{l-1}}{dx^{l-1}} (x - 1)^l (x + 1)^l.$$

A chaque dérivée, un des facteurs $(x - 1)$ ou $(x + 1)$ disparaît. On enlève donc, en faisant $\frac{d^{l-1}}{dx^{l-1}}$, précisément $(l - 1)$ des $2l$ facteurs. Chacun des termes du résultat contiendra donc au moins un facteur $(x - 1)$ et un facteur $(x + 1)$. Donc :

$$\frac{d^{l-1}}{dx^{l-1}} (x^2 - 1)^l \Big|_{-1}^1 = 0.$$

Ainsi

$$2^l l! 2^m m! (P_l, P_m) = - \int_{-1}^1 \frac{d^{l-1}}{dx^{l-1}} (x^2 - 1)^l \frac{d^{m+1}}{dx^{m+1}} (x^2 - 1)^m dx$$

Répétant cette première étape $(l - 1)$ autres fois, on aura :

$$= (-1)^l \int_{-1}^1 \left(\frac{d^{m+l}}{dx^{m+l}} (x^2 - 1)^m \right) (x^2 - 1)^l dx.$$

Mais $(x^2 - 1)^m$ est un polynôme de degré $2m$ que nous dérivons $m + l \geq 2m$ fois. Ainsi, si $l > m$, alors

$$(P_l, P_m) = 0, \quad \text{si } l \neq m,$$

et si $l = m$

$$\begin{aligned} (P_l, P_l) &= \frac{(-1)^l}{(2^l l!)^2} \int_{-1}^1 (2l)! (x^2 - 1)^l dx \\ &= \frac{(-1)^l (2l)!}{(2^l l!)^2} \int_{-1}^1 (x^2 - 1)^l dx \end{aligned}$$

Pour évaluer cette dernière intégrale, on fait :

$$\int_{-1}^1 (x^2 - 1)^l dx = \int_{-1}^1 (x - 1)^l (x + 1)^l dx$$

et par parties :

$$= \frac{(x - 1)^l (x + 1)^{l+1}}{l + 1} \Big|_{-1}^1 - \frac{l}{l + 1} \int_{-1}^1 (x - 1)^{l-1} (x + 1)^{l+1} dx$$

et encore $(l - 1)$ autres fois :

$$\begin{aligned} &= (-1)^l \frac{(l!)^2}{(2l)!} \int_{-1}^1 (x + 1)^{2l} dx \\ &= (-1)^l \frac{(l!)^2}{(2l)!} \frac{(x + 1)^{2l+1}}{2l + 1} \Big|_{-1}^1 \\ &= \frac{(-1)^l (l!)^2 2^{2l+1}}{(2l)! (2l + 1)} \end{aligned}$$

et alors

$$(P_l, P_l) = \frac{(-1)^l (2l)!}{2^{2l} (l!)^2} \frac{(-1)^l (l!)^2 2^{2l+1}}{(2l)! (2l + 1)} = \frac{2}{2l + 1}.$$

orthogonalité

Voilà ! Ainsi les polynômes de Legendre sont orthogonaux

$$(P_l, P_m) = 0, \quad l \neq m$$

normalisation

et leur norme au carré est

$$(P_l, P_l) = \frac{2}{2l + 1}.$$

Exemple. Voici un exemple de calcul de série de Legendre. Le problème que nous allons formuler peut être formulé de deux façons différentes :

- Quel est le potentiel à l'intérieur d'une sphère vide de rayon 1 si une différence de potentiel de 1 Volt est appliquée entre ses deux hémisphères nord et sud ?
- Développer la fonction saut :

$$f(x) = \begin{cases} 0, & -1 \leq x \leq 0 \\ 1, & 0 < x \leq 1 \end{cases}$$

en termes des polynômes de Legendre.

Dans le premier problème, il faut trouver le potentiel $V(r, \theta, \phi)$ à l'intérieur de la sphère si les conditions aux limites sont

$$V(1, \theta, \phi) = \begin{cases} 1, & 0 \leq \theta \leq \frac{\pi}{2}, \\ 0, & \frac{\pi}{2} < \theta \leq \pi. \end{cases}$$

Remarquons que les conditions aux limites ne dépendent pas de l'angle ϕ et que les polynômes de Legendre sont les fonctions propres à utiliser. En termes du développement de ces fonctions propres, la condition ci-dessus s'écrit

$$V(r=1, \theta, \phi) = \sum_{l=0}^{\infty} a_l P_l(\cos \theta) = \begin{cases} 1, & 0 \leq \theta \leq \frac{\pi}{2} \\ 0, & \frac{\pi}{2} < \theta \leq \pi, \end{cases}$$

qui n'est rien d'autre que la fonction f . Il faut donc trouver les coefficients a_l pour que cette égalité tienne, ce qui est clairement équivalent au second énoncé.

Rappelons la formule pour trouver les coefficients d'un développement en termes des fonctions propres d'un problème de Sturm-Liouville :

$$a_l = \frac{(f, P_l)}{(P_l, P_l)}.$$

Le coefficient de normalisation a été obtenu ci-dessus et alors :

$$\begin{aligned} &= \frac{2l+1}{2} \int_{-1}^1 f(x) P_l(x) dx \\ &= \frac{2l+1}{2} \int_0^1 P_l(x) dx \\ &= \frac{2l+1}{2} \frac{1}{2l+1} \int_0^1 (P'_{l+1} - P'_{l-1}) dx, \end{aligned}$$

par les relations de récurrence (voir ci-dessous)

$$\begin{aligned} &= \frac{1}{2} (P_{l+1} - P_{l-1})|_0^1 \\ &= -\frac{1}{2} (P_{l+1}(0) - P_{l-1}(0)), \quad \text{si } l \geq 1. \end{aligned}$$

(La relation de récurrence utilisée est une combinaison linéaire des deux relations de récurrence obtenues dans le texte. En faisant $\frac{d}{dx}(R1) - (l+1)(R2)$, on obtient

$$lP_l = xP'_l - P'_{l-1} \quad (R3)$$

et la relation utilisée ci-dessus est obtenue en faisant $2(R3) + (R2)$. Quant à a_0 , son calcul est aisé

$$a_0 = \frac{1}{2} \int_0^1 dx = \frac{1}{2}.$$

La formule (Δ) permet d'écrire plus explicitement le coefficient général. Si l est impair, $P_l(0) = 0$ puisque les polynômes de Legendre ont la parité de l . Si l est pair, seul le terme en x^0 demeure dans la somme :

$$P_l(0) = \frac{(-1)^{\frac{l}{2}} (2l-l)!}{2^l (l-\frac{l}{2})! \frac{l}{2}! (l-l)!}$$

et donc

$$P_{2m}(0) = \frac{(-1)^m (2m)!}{2^{2m} (m!)^2}.$$

Le coefficient a_l avec l pair est donc nul et celui avec $l = 2m + 1$ est :

$$\begin{aligned} a_{2m+1} &= \frac{1}{2} \left(\frac{(-1)^m (2m)!}{2^{2m} (m!)^2} + \frac{(-1)^m (2m+2)!}{2^{2(m+1)} ((m+1)!)^2} \right) \\ &= \frac{1}{2} \frac{(-1)^m (2m)!}{2^{2m} (m!)^2} \left(1 + \frac{(2m+2)(2m+1)}{4(m+1)(m+1)} \right) \\ &= \frac{(-1)^m (2m)!}{2^{2m+1} (m!)^2} \frac{4m+3}{2m+2}. \end{aligned}$$

Ainsi la fonction f peut s'écrire

$$f(x) = \frac{1}{2} + \sum_{m=0}^{\infty} \frac{(-1)^m (2m)!}{2^{2m+1} (m!)^2} \frac{4m+3}{2m+2} P_{2m+1}(x)$$

et le potentiel désiré est

$$V(r, \theta, \phi) = \frac{1}{2} + \sum_{m=0}^{\infty} \frac{(-1)^m (2m)!}{2^{2m+1} (m!)^2} \frac{4m+3}{2m+2} r^{2m+1} P_{2m+1}(\cos \theta).$$

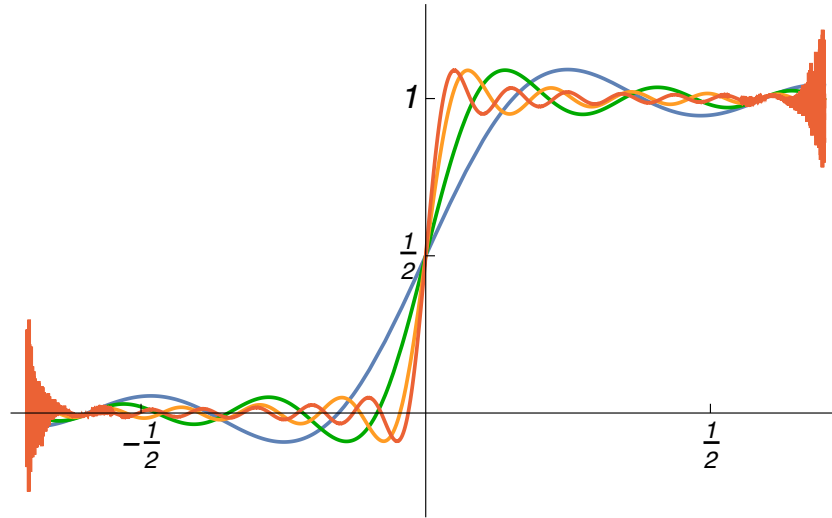


Figure 3.3 – La somme des cent premiers termes de la série de Legendre pour f

La figure ci-contre trace le graphe de la somme partielle des cent premiers termes. Le domaine a été restreint à $[-0.7, 0.7]$. La fonction semble tendre vers la fonction saut f , mais après cent termes on n'aurait pu s'attendre à mieux. De

plus le graphe aux extrémités de cet intervalle semble indiquer un désastre. Un exercice ci-dessous étudiera le pourquoi de ce désastre apparent.

$P_0(x) = 1$	$P_4(x) = \frac{1}{8}(35x^4 - 30x^2 + 3)$
$P_1(x) = x$	$P_5(x) = \frac{1}{8}(63x^5 - 70x^3 + 15x)$
$P_2(x) = \frac{1}{2}(3x^2 - 1)$	$P_6(x) = \frac{1}{16}(231x^6 - 315x^4 + 105x^2 - 5)$
$P_3(x) = \frac{1}{2}(5x^3 - 3x)$	$P_7(x) = \frac{1}{16}(429x^7 - 693x^5 + 315x^3 - 35x)$

18*. Plusieurs des caractéristiques que nous avons obtenues pour les polynômes de Legendre vous sont familières pour les fonctions trigonométriques. Les exercices suivants vous en convaincront. Attention : certains sont *vraiment* triviaux ; ils ont pour seul but de vous faire remarquer que les propriétés obtenues pour les polynômes de Legendre sont similaires aux propriétés que nous utilisons tous les jours pour les fonctions trigonométriques.

Exercices

- (a) Ecrire les relations de récurrence exprimant $\sin nx$ et $\cos nx$ en termes des $\sin mx$ et $\cos mx$ avec $m < n$.
- (b) Ecrire les relations de récurrence exprimant les dérivées de $\sin nx$ et $\cos nx$ en termes des $\sin mx$ et $\cos mx$ avec $m \leq n$.
- (c) Obtenir la fonction génératrice pour les fonctions $\sin nx$ et $\cos nx$, c'est-à-dire calculer :

$$S(x, z) = \sum_{n=1}^{\infty} z^n \sin nx$$

$$C(x, z) = \sum_{n=0}^{\infty} z^n \cos nx$$

Note : Ecrire les fonctions trigonométriques sous forme d'exponentielles et sommer.

- (d) Donner les valeurs particulières de $\sin nx$ et $\cos nx$ en $x = 0$ et $\frac{\pi}{2}$.
- (e) Montrer que les $\sin nx$ et $\cos mx$ sont orthogonaux deux à deux sur l'intervalle $[-\pi, \pi]$ et calculer leur normalisation.

19*. Réorganiser l'ordre de sommation pour les doubles sommes suivantes. (En d'autres mots, remplacer les ? par leur valeur précise.)

- (a)

$$\sum_{n=0}^{\infty} \sum_{m=0}^{\infty} A_{nm} = \sum_{m=?}^? \sum_{n=?}^? A_{nm}$$

(b)

$$\sum_{n=0}^{\infty} \sum_{m=0}^n B_{nm} = \sum_{m=0}^{\infty} \sum_{n=?}^? B_{nm}$$

(c)

$$\sum_{n=0}^{\infty} \sum_{m=0}^n C_{nm} = \sum_{l=0}^{\infty} \sum_{n=?}^? C_{n,l-n}$$

(d)

$$\sum_{n=0}^{\infty} \sum_{m=0}^n D_{nm} = \sum_{l=0}^{\infty} \sum_{m=?}^? D_{l+m,m}$$

20*. (a) Ecrire l'identité de Parseval pour les séries de polynômes de Legendre.

(b) Pour le restant de cet exercice, la fonction f est celle traitée en exemple dans la section. Prendre pour acquis l'égalité entre f et la série de polynômes de Legendre aux points où f est continue. Evaluer cette relation en $x = \frac{1}{2}, 1, -\frac{1}{2}$ et -1 et vérifier le résultat obtenu à l'aide d'un langage de manipulations symboliques.

(c) A l'aide de la formule de Stirling (voir les équations (Γ_1) et (Γ_2)), montrer que le coefficient a_{2m+1} se comporte asymptotiquement comme

$$a_{2m+1} \sim \frac{(-1)^m}{\sqrt{\pi m}}, \quad m \rightarrow \infty.$$

En conclure que $a_{21} \approx 10a_{2001}$. A ce rythme tortuesque, ce n'est pas surprenant que le graphe de la figure 3.3 soit si désappointant. (Si vous n'êtes pas convaincu que ce rapport a_{2001}/a_{21} est encore très proche de 1, il suffit de calculer l'ordre de grandeur du rapport $21!/2001!$ qui est le rapport des 2001ième et 21ième coefficients du développement en série de Taylor de la fonction exponentielle : $21!/2001! \approx 7 \times 10^{-5720}$.)

(d) Ecrire l'identité de Parseval pour cette égalité. Vérifier l'accord numérique.

21. (a) Développer la fonction $f : [-1, 1] \rightarrow \mathbb{R}$ donnée par $f(x) = x^3$ en termes des polynômes de Legendre et vérifier que la somme $\sum_{l=0}^{\infty} a_l P_l$ redonne bien f .

(b) Vérifier que l'identité de Parseval est satisfaite.

22. Utiliser la fonction génératrice f pour déterminer la normalisation des polynômes de Legendre. Pour cela développer les deux membres de l'égalité

$$\int_{-1}^1 f(x)^2 dx = \int_{-1}^1 \sum_{m=0}^{\infty} \sum_{n=0}^{\infty} P_m P_n r^{m+n} dx$$

en prenant pour acquis que les P_n sont orthogonaux. Cet exercice donne un autre usage de la fonction génératrice.

23. (a) Montrer que le développement de la fonction $f(x) = |x|$ en une série de polynômes de Legendre :

$$f(x) = \sum_{l=0}^{\infty} a_l P_l(x)$$

est donné par

$$a_l = \begin{cases} \frac{1}{2}, & l = 0 \\ \frac{l+1}{2l+3} (P_l(0) - P_{l+2}(0)) + \frac{l}{2l-1} (P_{l-2}(0) - P_l(0)), & l \geq 2, \text{ pair} \\ 0, & l \text{ impair.} \end{cases}$$

(Suggestion : Relations de récurrence !)

(b) Pour les braves : montrer que

$$a_{2l} = \frac{(-1)^{l-1} (2l-2)! (4l+1)}{2^{2l} (l-1)! (l+1)!}.$$

(Suggestion : Sauter cet exercice si vous n'avez pas beaucoup de temps.)

(c) Montrer que l'erreur relative entre le membre de gauche et les trois premiers termes non-nuls du membre de droite de l'identité de Parseval est $< \frac{1}{250}$. (Cette convergence est donc beaucoup plus rapide que celle de l'exemple traité dans la section. Voir l'exercice **3.** ci-dessus.)

24*. Les polynômes d'Hermite sont définis par

$$H_n(x) = (-1)^n e^{x^2} \frac{d^n}{dx^n} e^{-x^2}, \quad n = 0, 1, 2, \dots$$

polynômes d'Hermite
Charles Hermite
(1822-1901)

(a) Obtenir les 3 premiers polynômes d'Hermite.

(b) Utiliser la définition ci-dessus pour montrer que

$$H_{n+1}(x) = 2xH_n(x) - 2nH_{n-1}(x).$$

Note : Les deux premières étapes de ce calcul sont :

$$\frac{d^{n+1}}{dx^{n+1}} e^{-x^2} = \frac{d^n}{dx^n} (-2xe^{-x^2}) = -2 \frac{d^{n-1}}{dx^{n-1}} e^{-x^2} - 2 \frac{d^{n-1}}{dx^{n-1}} \left(x \frac{d}{dx} e^{-x^2} \right).$$

(c) Montrer que : $H'_n = 2nH_{n-1}$.

(d) Montrer que les polynômes d'Hermite ($n = 0, 1, 2, \dots$) satisfont l'équation différentielle : $y'' - 2xy' + 2ny = 0$.

(e) Montrer que cette équation différentielle peut être réécrite sous la forme

$$e^{x^2} \frac{d}{dx} \left(e^{-x^2} \frac{d}{dx} y \right) + 2ny = 0.$$

(f) Donner la fonction poids et l'intervalle d'intégration du produit $(\cdot, \cdot)$ pour lequel les polynômes H_n et H_m , $m \neq n$, sont orthogonaux.

(g) Utilisant (c), montrer que $(H_n, H_n) = 2^n n! \sqrt{\pi}$.

(h) Montrer que la fonction génératrice pour H_n est e^{2tx-t^2} , c'est-à-dire :

$$e^{2tx-t^2} = \sum_{n=0}^{\infty} \frac{t^n H_n(x)}{n!}.$$

Note : Montrer tout d'abord que :

$$\left(\frac{\partial^n}{\partial t^n} e^{-(t-x)^2} \right) \Big|_{t=0} = (-1)^n \frac{d^n}{dx^n} e^{-x^2}.$$

Remarque : Les polynômes d'Hermite sont au cœur de la description de l'oscillateur harmonique en mécanique quantique.

polynômes de Laguerre
Edmond Laguerre
(1834-1886)

25*. Les polynômes de Laguerre $L_n(x)$ satisfont l'équation différentielle de Laguerre :

$$xy'' + (1-x)y' + ny = 0, \quad n \in \mathbb{N} \cup \{0\}.$$

(a) Montrer que la solution polynomiale de degré n qui vérifie $L_n(0) = 1$ est

$$L_n(x) = \sum_{i=0}^n \frac{(-1)^i n! x^i}{i! i! (n-i)!}.$$

Note : Ecrire $L_n(x) = \sum_{i=0}^n l_i x^i$ et trouver une équation récursive pour les l_i à l'aide de l'équation différentielle.

(b) Ecrire explicitement les 4 premiers polynômes de Laguerre.

(c) Montrer la formule de Rodrigues :

$$L_n(x) = \frac{e^x}{n!} \frac{d^n}{dx^n} (x^n e^{-x}).$$

Note : Montrer d'abord que

$$\frac{d^n}{dx^n}(f \cdot g) = \sum_{i=0}^n \binom{n}{i} \frac{d^i f}{dx^i} \frac{d^{n-i} g}{dx^{n-i}}.$$

(d) Utilisant la formule de Rodrigues (ou autrement), montrer que

$$(L_n, L_m) = \int_0^\infty e^{-x} L_n(x) L_m(x) dx = \delta_{mn}.$$

Note : Se convaincre d'abord que

$$\left. \frac{d^i}{dx^i}(x^n e^{-x}) \right|_0 = \left. \frac{d^i}{dx^i}(x^n e^{-x}) \right|_\infty = 0$$

si $i < n$.

(e) Montrer que la fonction génératrice des polynômes de Laguerre est :

$$F(x, z) = \frac{e^{-xz/(1-z)}}{1-z} = \sum_{n=0}^{\infty} L_n(x) z^n.$$

Note : Pour obtenir le développement en série de $(1-z)^{-j-1}$, dériver j fois le développement de $\frac{1}{1-z}$. (Est-ce une méthode rigoureuse ?)

(f) En dérivant la fonction génératrice par rapport à z , obtenir la relation de récurrence :

$$(n+1)L_{n+1} = (2n+1-x)L_n - nL_{n-1}, \quad n \geq 0.$$

(g) Montrer que la fonction génératrice satisfait l'identité :

$$x \frac{\partial F}{\partial x} = z(1-z) \frac{\partial F}{\partial z} - zF.$$

En déduire la relation de récurrence :

$$xL'_n = nL_n - nL_{n-1}.$$

3.5 Fonctions de Bessel

A la section 3.2, nous avons montré que la séparation de variables pour le laplacien (ou l'équation de Helmholtz) en coordonnées cylindriques donne lieu

à trois équations différentielles ordinaires : si $\psi(r, \theta, z) = R(r)\Theta(\theta)Z(z)$ vérifie $\nabla^2\psi + k^2\psi = 0$, alors :

$$\frac{Z''}{Z} = \lambda_1,$$

$$\frac{\Theta''}{\Theta} = \lambda_2$$

et

$$\frac{d}{dr} \left(r \frac{dR}{dr} \right) + \left(r(k^2 + \lambda_1) + \frac{\lambda_2}{r} \right) R = 0.$$

Cette dernière équation est réécrite, après le changement de variable :

$$y(x) = R(r) \quad \text{où} \quad x = \kappa r \quad \text{et} \quad \kappa = \sqrt{k^2 - \frac{n^2\pi^2}{L^2}},$$

équation de Bessel

sous la forme

$$\frac{d}{dx} \left(x \frac{dy}{dx} \right) + \left(x - \frac{m^2}{x} \right) y = 0$$

Bessel, Friedrich
(1784, 1846)

où $\lambda_1 = -n^2\pi^2/L^2$ et $\lambda_2 = -m^2$. (Voir section 3.2.) Si Θ doit être continue, m doit être un entier. Mais, dans ce qui suit, nous allons nous intéresser au cas où m^2 est un nombre positif et nous le noterons plutôt ν^2 pour éviter l'allusion à un entier. L'équation peut alors être réécrite

$$x \frac{d^2y}{dx^2} + \frac{dy}{dx} + \left(x - \frac{\nu^2}{x} \right) y = 0$$

ou encore

$$y'' + \frac{1}{x}y' + \left(1 - \frac{\nu^2}{x^2} \right) y = 0.$$

Nous allons étudier les solutions de cette équation différentielle en procédant, cette fois-ci, à partir de la solution explicite de l'équation différentielle. (Je suis le traitement de Schwartz.²)

méthode de Frobenius
Ferdinand Georg Frobenius
(1849-1917)

Posons une solution de la forme

$$y = x^\lambda \sum_{k=0}^{\infty} a_k x^k.$$

Nous devons identifier les constantes λ et a_k , $k \in \mathbb{N} \cup \{0\}$. Les dérivées de cette série sont

$$y' = x^\lambda \sum_{k=0}^{\infty} (\lambda + k) a_k x^{k-1}$$

$$y'' = x^\lambda \sum_{k=0}^{\infty} (\lambda + k)(\lambda + k - 1) a_k x^{k-2}.$$

2. Le livre de Geroge Neville Watson *Treatise on the theory of Bessel functions* publié en 1922 et réédité en 1995 par Cambridge University Press, constitue la monographie la plus complète sur les fonctions de Bessel.

L'équation différentielle donne

$$\begin{aligned}
 0 &= y'' + \frac{1}{x}y' + \left(1 - \frac{\nu^2}{x^2}\right)y \\
 &= x^\lambda \left(\sum_{k=0}^{\infty} (\lambda + k)(\lambda + k - 1)a_k x^{k-2} \right. \\
 &\quad \left. + \sum_{k=0}^{\infty} (\lambda + k)a_k x^{k-2} + \sum_{k=0}^{\infty} a_k x^k - \nu^2 \sum_{k=0}^{\infty} a_k x^{k-2} \right) \\
 &= x^\lambda \left(\sum_{k=0}^{\infty} [(\lambda + k)^2 - (\lambda + k) + (\lambda + k) - \nu^2] x^{k-2} a_k + \sum_{k=0}^{\infty} a_k x^k \right) \\
 &= x^\lambda \left(\sum_{k=0}^{\infty} [(\lambda + k)^2 - \nu^2] x^{k-2} a_k + \sum_{k=0}^{\infty} a_k x^k \right).
 \end{aligned}$$

En identifiant les coefficients des puissances de x , les équations sont

$$\begin{aligned}
 x^{\lambda-2} : & \quad ((\lambda + 0)^2 - \nu^2)a_0 = 0 \\
 x^{\lambda-1} : & \quad ((\lambda + 1)^2 - \nu^2)a_1 = 0 \\
 x^{\lambda+k-2} : & \quad ((\lambda + k)^2 - \nu^2)a_k + a_{k-2} = 0, \quad k \geq 2.
 \end{aligned}$$

Si on suppose $a_0 \neq 0$, alors $\lambda = \pm\nu$. Ainsi $(\lambda + 1)^2 - \nu^2 = \nu^2 \pm 2\nu + 1 - \nu^2 = \pm 2\nu + 1$
et si $\lambda = \nu$:

cas $\lambda = \nu$

$$\begin{aligned}
 (2\nu + 1)a_1 &= 0 \\
 (2\nu k + k^2)a_k + a_{k-2} &= 0.
 \end{aligned}$$

Il faut donc que $a_1 = 0$ (et donc que tous les a_i avec i impair soient également nuls) et que :

$$\begin{aligned}
 a_k &= -\frac{a_{k-2}}{k(2\nu + k)} \\
 &= \frac{(-1)^2 a_{k-4}}{k(k-2)(2\nu + k)(2\nu + k-2)} \\
 &= \dots \\
 &= \frac{(-1)^{\frac{k}{2}} a_0}{[k(k-2)\dots 2][(2\nu + k)(2\nu + k-2)\dots (2\nu + 2)]} \\
 &= \frac{(-1)^{\frac{k}{2}} a_0}{2^{\frac{k}{2}} \frac{k}{2}! 2^{\frac{k}{2}} (\nu + \frac{k}{2})(\nu + \frac{k}{2} - 1)\dots (\nu + 1)} \\
 &= \frac{(-1)^{\frac{k}{2}} a_0 \Gamma(\nu + 1)}{2^k \frac{k}{2}! \Gamma(\nu + \frac{k}{2} + 1)}.
 \end{aligned}$$

Choisissons

$$a_0 = \frac{1}{2^\nu \Gamma(\nu + 1)}.$$

Alors

$$a_k = \frac{(-1)^{\frac{k}{2}}}{2^{\nu+k} \frac{k}{2}! \Gamma(\nu + \frac{k}{2} + 1)}$$

ou, posant $k = 2n$:

$$a_{2n} = \frac{(-1)^n}{2^{\nu+2n} n! \Gamma(\nu + n + 1)},$$

J_ν

et nous venons de trouver une solution :

$$J_\nu(x) = \left(\frac{x}{2}\right)^\nu \sum_{n=0}^{\infty} \frac{(-1)^n}{n! \Gamma(\nu + n + 1)} \left(\frac{x}{2}\right)^{2n}.$$

cas $\lambda = -\nu$: $J_{-\nu}$

Similairement, pour $\lambda = -\nu$, on trouve :

$$J_{-\nu}(x) = \left(\frac{x}{2}\right)^{-\nu} \sum_{n=0}^{\infty} \frac{(-1)^n}{n! \Gamma(-\nu + n + 1)} \left(\frac{x}{2}\right)^{2n}, \quad \nu \notin \mathbb{N}.$$

Ces deux séries convergent uniformément (en x) sur tout intervalle borné. (Est-ce vrai ? Comparer avec le développement en série de Taylor de $e^{-\frac{x^2}{2}}$.)

Si ν n'est pas un entier positif ou nul, alors J_ν et $J_{-\nu}$ sont deux solutions linéairement indépendantes. En effet, en $x \rightarrow 0$, les comportements sont :

$$J_\nu \rightarrow \left(\frac{x}{2}\right)^\nu \frac{1}{\Gamma(\nu + 1)} \quad \text{et} \quad J_{-\nu} \rightarrow \left(\frac{x}{2}\right)^{-\nu} \frac{1}{\Gamma(-\nu + 1)}$$

qui sont visiblement non-proportionnels. Si on a $\nu = 0$, on a qu'une solution et si $\nu = p \in \mathbb{N}$, alors

$$J_{-p} = \left(\frac{x}{2}\right)^{-p} \sum_{n=0}^{\infty} \frac{(-1)^n}{n! \Gamma(n - p + 1)} \left(\frac{x}{2}\right)^{2n}.$$

Mais $\Gamma(n - p + 1) = \pm\infty$ pour $n - p + 1 \leq 0$. Alors on peut *formellement* sauter ces termes dans la somme et commencer à $n = p$ (exercice : comment rendre cet argument rigoureux ?) :

$$\begin{aligned} &= \left(\frac{x}{2}\right)^{-p} \sum_{n=p}^{\infty} \frac{(-1)^n}{n! \Gamma(n - p + 1)} \left(\frac{x}{2}\right)^{2n} \\ &= \left(\frac{x}{2}\right)^{-p} \sum_{m=0}^{\infty} \frac{(-1)^{m+p}}{(m+p)! \Gamma(m+1)} \left(\frac{x}{2}\right)^{2(m+p)}, \quad \text{avec } m = n - p \\ &= (-1)^p \left(\frac{x}{2}\right)^p \sum_{m=0}^{\infty} \frac{(-1)^m}{\Gamma(m+p+1) m!} \left(\frac{x}{2}\right)^{2m} \\ &= (-1)^p J_p(x). \end{aligned}$$

Donc, dans ce cas, les deux solutions J_p et J_{-p} ne sont pas linéairement indépendantes. Il nous manque une solution.

Un mot sur comment sont obtenues les solutions manquantes. La méthode de Frobenius peut être étendue pour les construire mais pour les fonctions de Bessel, il est usuel de définir les fonctions de Neumann :

fonctions de
Neumann
Franz Ernst Neumann
(1798-1895)

$$N_\nu(x) = \frac{J_\nu(x) \cos \nu\pi - J_{-\nu}(x)}{\sin \nu\pi}, \quad \nu \notin \mathbb{Z}$$

qui sont simplement une combinaison linéaire des deux solutions que nous venons d'obtenir. C'est donc une solution de l'équation de Bessel linéairement indépendante de $J_\nu(x)$. En $\nu = p$, le rapport n'est cependant pas défini puisqu'il prend la forme

$$\frac{(-1)^p J_p - J_{-p}(x)}{\sin(p\pi)}.$$

Remarquons que, pour $\nu \notin \mathbb{Z}$, N_ν est une fonction analytique de ν puisque la fonction

$$\nu \mapsto J_\nu(x) = \left(\frac{x}{2}\right)^\nu \sum_{n=0}^{\infty} \frac{(-1)^n}{n! \Gamma(n + \nu + 1)} \left(\frac{x}{2}\right)^{2n}$$

vue comme fonction de ν l'est. En effet, la série converge uniformément (en ν) sur tout intervalle $[a, b]$ (puisque $1/\Gamma(n + \nu + 1)$ est analytique en ν , fait que nous n'avons cependant pas démontré). Donc il est possible d'obtenir N_p en utilisant la règle de l'Hospital :

$$N_p(x) = \lim_{\nu \rightarrow p} \frac{\frac{\partial}{\partial \nu} (J_\nu(x) \cos \nu\pi - J_{-\nu}(x))}{\frac{\partial}{\partial \nu} (\sin \nu\pi)}.$$

Pour obtenir $N_p(x)$ comme une série, il faut savoir dériver les $\Gamma(x)$, ce que nous n'avons pas discuté dans l'appendice A. Donc nous prendrons pour acquis que la fonction $N_p(x)$ existe, est non-nulle et linéairement indépendante de $J_p(x)$, $p \in \mathbb{N} \cup \{0\}$. Elle est donnée sous forme de série par

$$\begin{aligned} \pi N_p(x) = & 2 \left[\ln \frac{x}{2} + \gamma - \frac{1}{2} \sum_{n=1}^p \frac{1}{n} \right] J_p(x) \\ & - \sum_{n=0}^{\infty} \left(\frac{(-1)^n}{n!(n+p)!} \left(\frac{x}{2}\right)^{p+2n} \cdot \sum_{k=1}^n \left(\frac{1}{k} + \frac{1}{k+p} \right) \right) \\ & - \sum_{n=0}^{p-1} \frac{(p-n-1)!}{n!} \left(\frac{x}{2}\right)^{-p+2n} \end{aligned}$$

où le nombre $\gamma = 0,5772156\dots$ est la constante d'Euler définie par la limite $\gamma = \lim_{n \rightarrow \infty} (\sum_{i=1}^n \frac{1}{i} - \ln n)$. Quelque peu impressionnant... Remarquons qu'à cause du terme $(\frac{x}{2})^{-p}$, son comportement est singulier en $x = 0$. Nous nous concentrerons par la suite sur les fonctions J_ν .

transformée de
Fourier

Notre prochaine étape dans l'étude des fonctions de Bessel est de les représenter par une intégrale. On se rappellera que c'est précisément de cette façon que nous avons introduit la fonction Γ . Pour cela nous introduisons une opération « $\hat{}$ » qui agit sur les fonctions $f : \mathbb{R} \rightarrow \mathbb{C}$ qui ont la propriété que $\int_{-\infty}^{\infty} |f(t)| dt < \infty$. L'opération transforme une telle fonction f en une fonction « $\hat{f}(x)$ » définie par

$$(\hat{f})(x) = \int_{-\infty}^{\infty} f(t) e^{-itx} dt.$$

linéarité

L'opération « $\hat{}$ » s'appelle la transformée de Fourier et le chapitre 4 est consacré à son étude. Les deux seules propriétés dont nous aurons besoin sont la linéarité³

$$\begin{aligned} (f_1 + cf_2)^{\hat{}}(x) &= \int_{-\infty}^{\infty} (f_1(t) + cf_2(t)) e^{-itx} dt \\ &= \int_{-\infty}^{\infty} f_1(t) e^{-itx} dt + c \int_{-\infty}^{\infty} f_2(t) e^{-itx} dt \\ &= \hat{f}_1(x) + c\hat{f}_2(x) \end{aligned}$$

 $\frac{d}{dx}$

et la possibilité de dériver à l'intérieur du signe d'intégration :

$$\frac{d}{dx} \hat{f}(x) = \int_{-\infty}^{\infty} f(t) (-it) e^{-itx} dt = -i(\widehat{tf})(x).$$

(Nous reviendrons, au prochain chapitre, sur cette possibilité d'intervertir la dérivée et l'intégration pour en spécifier les limites d'application.) Remarquons que cette opération ressemble beaucoup au calcul des coefficients d'une série de Fourier :

$$\begin{aligned} f(t) &\xrightarrow{\text{série de Fourier}} c_n, & n \in \mathbb{Z}, \\ f(t) &\xrightarrow{\text{transformée de Fourier}} \hat{f}(x), & x \in \mathbb{R}. \end{aligned}$$

Soit la fonction

$$f_{\nu}(t) = \begin{cases} (1 - t^2)^{\nu - \frac{1}{2}}, & |t| < 1, \\ 0, & \text{autrement.} \end{cases}$$

Pour $\nu + \frac{1}{2} > 0$, f_{ν} est intégrable⁴ sur $\mathbb{R}$. Soit

$$\begin{aligned} \hat{f}_{\nu}(x) &= \int_{-\infty}^{\infty} f_{\nu}(t) e^{-itx} dt \\ &= \int_{-1}^1 (1 - t^2)^{\nu - \frac{1}{2}} e^{-itx} dt. \end{aligned}$$

3. On appose parfois le « $\hat{}$ » après le symbole de la fonction pour alléger la notation. Par exemple, on notera $(f + g)^{\hat{}}$ plutôt que $\widehat{f + g}$.

4. La fonction f n'est intégrable au sens de Riemann que si $\nu \geq \frac{1}{2}$ car, si $-\frac{1}{2} < \nu < \frac{1}{2}$, elle n'est pas bornée sur $[-1, 1]$. Cependant, pour ces dernières valeurs de ν , l'intégrale de f sur l'intervalle $(-1 + \epsilon, 1 - \epsilon)$ existe pour tout $0 < \epsilon < 1$ et la limite de cette intégrale lorsque $\epsilon \rightarrow 0$ existe et est finie.

La fonction $\widehat{f}_\nu$ est de $\mathbb{R}$ dans $\mathbb{R}$. (Que ses valeurs soient réelles n'est pas tout à fait évident puisque l'intégrand prend des valeurs complexes. Comment peut-on se convaincre que le résultat est réel pour x et ν réel ?) Nous allons montrer que $Z_\nu(x) = x^\nu \widehat{f}_\nu(x)$ est une solution de l'équation de Bessel. Notons d'abord que

$$\begin{aligned}\frac{d}{dx} \widehat{f}_\nu &= \int_{-1}^1 (1-t^2)^{\nu-\frac{1}{2}} (-it) e^{-itx} dt = (-it f_\nu)^\wedge \\ \frac{d^2}{dx^2} \widehat{f}_\nu &= \int_{-1}^1 (1-t^2)^{\nu-\frac{1}{2}} (-t^2) e^{-itx} dt = (-t^2 f_\nu)^\wedge.\end{aligned}$$

Donc

$$\begin{aligned}Z_\nu'' + \frac{1}{x} Z_\nu' + \left(1 - \frac{\nu^2}{x^2}\right) Z_\nu &= \\ &= \frac{d^2}{dx^2} (x^\nu \widehat{f}_\nu) + \frac{1}{x} \frac{d}{dx} (x^\nu \widehat{f}_\nu) + \left(1 - \frac{\nu^2}{x^2}\right) x^\nu \widehat{f}_\nu \\ &= \nu(\nu-1) x^{\nu-2} \widehat{f}_\nu + 2\nu x^{\nu-1} (\widehat{f}_\nu)' + x^\nu (\widehat{f}_\nu)'' \\ &\quad + \nu x^{\nu-2} \widehat{f}_\nu + x^{\nu-1} (\widehat{f}_\nu)' \\ &\quad + x^\nu \widehat{f}_\nu - \nu^2 x^{\nu-2} \widehat{f}_\nu \\ &= x^\nu \widehat{f}_\nu + (2\nu+1) x^{\nu-1} (\widehat{f}_\nu)' + x^\nu (\widehat{f}_\nu)'' \\ &= x^\nu ((1-t^2) f_\nu)^\wedge + (2\nu+1) x^{\nu-1} (\widehat{f}_\nu)'.\end{aligned}$$

Par la définition même de f_ν , le premier terme ci-dessus est $((1-t^2) f_\nu)^\wedge = (f_{\nu+1})^\wedge$. Il reste donc à montrer que $x^\nu \widehat{f}_{\nu+1} = -(2\nu+1) x^{\nu-1} (\widehat{f}_\nu)'$ pour terminer la preuve. Notons tout d'abord que :

$$\widehat{f'_{\nu+1}} = \int_{-1}^1 f'_{\nu+1} e^{-itx} dt$$

qui, par intégration par parties, se réécrit

$$\begin{aligned}&= f_{\nu+1} e^{-itx} \Big|_{-1}^1 + ix \int_{-1}^1 f_{\nu+1} e^{-itx} dt \\ &= ix (\widehat{f_{\nu+1}}), \quad \text{si } \nu + \frac{1}{2} > 0.\end{aligned}$$

La dérivée de $f_{\nu+1}$ peut être aisément obtenue :

$$f'_{\nu+1} = -2t(\nu + \frac{1}{2})(1-t^2)^{\nu-\frac{1}{2}} = -(2\nu+1) t f_\nu.$$

Ces deux remarques permettent donc de réécrire le terme $x^\nu \widehat{f_{\nu+1}}$:

$$\begin{aligned}x^\nu \widehat{f_{\nu+1}} &= -ix^{\nu-1} (\widehat{f'_{\nu+1}}) = i(2\nu+1) x^{\nu-1} (t f_\nu)^\wedge \\ &= -(2\nu+1) x^{\nu-1} (-it f_\nu)^\wedge = -(2\nu+1) x^{\nu-1} (\widehat{f}_\nu)'. \end{aligned}$$

D'où :

$$\begin{aligned} Z''_\nu + \frac{1}{x} Z'_\nu + \left(1 - \frac{\nu^2}{x^2}\right) Z_\nu &= \\ &= x^\nu \widehat{f_{\nu+1}} + (2\nu + 1)x^{\nu-1}(\widehat{f_\nu})' \\ &= 0. \end{aligned}$$

Ainsi $Z_\nu(x)$ est une solution de l'équation de Bessel et il est possible de trouver une combinaison linéaire de J_ν et N_ν qui reproduise cette solution. Lorsque $x = 0$, l'intégrale ci-dessus est :

$$\begin{aligned} \widehat{f_\nu}(0) &= \int_{-1}^1 (1-t^2)^{\nu-\frac{1}{2}} dt \\ &= 2 \int_0^1 (1-t^2)^{\nu-\frac{1}{2}} dt \end{aligned}$$

et avec le changement de variable $s = t^2$, $ds = 2t dt$

$$\begin{aligned} &= \int_0^1 (1-s)^{\nu-\frac{1}{2}} s^{-\frac{1}{2}} ds \\ &= \frac{\Gamma(\nu + \frac{1}{2})\Gamma(\frac{1}{2})}{\Gamma(\nu + 1)} \end{aligned}$$

où nous avons utilisé la fonction $B(p, q)$. Donc

$$Z_\nu(x) \xrightarrow{x \rightarrow 0} \frac{x^\nu \sqrt{\pi} \Gamma(\nu + \frac{1}{2})}{\Gamma(\nu + 1)}.$$

Ce comportement lorsque x approche 0 exclut donc N_ν qui, elle, diverge quand $x \rightarrow 0$. Donc Z_ν et J_ν sont proportionnelles : $Z_\nu = a_\nu J_\nu$. Puisque

$$J_\nu(x) \xrightarrow{x \rightarrow 0} \left(\frac{x}{2}\right)^\nu \frac{1}{\Gamma(\nu + 1)},$$

représentation
intégrale des $J_\nu(x)$

il faut que $a_\nu = 2^\nu \sqrt{\pi} \Gamma(\nu + \frac{1}{2})$ et donc que

$$J_\nu(x) = \frac{\left(\frac{x}{2}\right)^\nu}{\sqrt{\pi} \Gamma(\nu + \frac{1}{2})} \int_{-1}^1 (1-t^2)^{\nu-\frac{1}{2}} e^{-itx} dt, \quad \nu + \frac{1}{2} > 0.$$

En montrant que Z_ν est une solution de l'équation de Bessel, nous avons démontré la relation

$$\widehat{f_{\nu+1}} = -\frac{2\nu+1}{x}(\widehat{f_\nu})'.$$

Cette relation est presque une relation entre $J_{\nu+1}$ et J_ν . Pour en faire une véritable relation de récurrence, rappelons que pour passer des $Z_\nu(x) = x^\nu \widehat{f_\nu}(x)$ aux J_ν , il faut diviser par a_ν . Notons que

$$a_{\nu+1} = 2^{\nu+1} \sqrt{\pi} \Gamma(\nu + \frac{3}{2}) = 2 \cdot 2^\nu \sqrt{\pi} (\nu + \frac{1}{2}) \Gamma(\nu + \frac{1}{2}) = (2\nu + 1) a_\nu.$$

Ainsi

$$\frac{\widehat{f_{\nu+1}}}{a_{\nu+1}} = -\frac{1}{x} \left(\frac{\widehat{f_{\nu}}}{a_{\nu}} \right)'$$

et donc

$$x^{-\nu-1} J_{\nu+1} = -\frac{1}{x} \frac{d}{dx} (x^{-\nu} J_{\nu})$$

qu'on peut réécrire également

$$x^{-\nu-1} J_{\nu+1} = \frac{\nu}{x} x^{-\nu-1} J_{\nu} - x^{-\nu-1} J'_{\nu},$$

c'est-à-dire

$$J'_{\nu} = \frac{\nu}{x} J_{\nu} - J_{\nu+1}, \quad \nu + \frac{1}{2} > 0. \quad (R1) \text{ première relation de récurrence}$$

Une seconde relation similaire existe

$$J'_{\nu} = -\frac{\nu}{x} J_{\nu} + J_{\nu-1}, \quad \nu - \frac{1}{2} > 0.$$

La preuve de cette seconde relation procède en deux étapes. Pour la première, notons que

$$\begin{aligned} (2\nu - 1)f_{\nu} - tf'_{\nu} &= (2\nu - 1)(1 - t^2)^{\nu-\frac{3}{2}}(1 - t^2) - t(\nu - \frac{1}{2})(-2t)(1 - t^2)^{\nu-\frac{3}{2}} \\ &= (2\nu - 1)(1 - t^2)^{\nu-\frac{3}{2}}(1 - t^2 + t^2) \\ &= (2\nu - 1)f_{\nu-1}. \end{aligned}$$

Ainsi

$$\frac{2\nu \widehat{f_{\nu}} - (f_{\nu} + tf'_{\nu})^{\wedge}}{a_{\nu}} = \frac{(2\nu - 1)\widehat{f_{\nu-1}}}{(2\nu - 1)a_{\nu-1}} = \frac{\widehat{f_{\nu-1}}}{a_{\nu-1}}. \quad (\diamond)$$

La seconde étape consiste en réécrire le terme $(f_{\nu} + tf'_{\nu})^{\wedge}$:

$$\begin{aligned} (f_{\nu} + tf'_{\nu})^{\wedge} &= \left(\frac{d}{dt}(tf_{\nu}) \right)^{\wedge} = \int_{-1}^1 \frac{d}{dt}(tf_{\nu}) e^{-itx} dt \\ &= tf_{\nu} e^{-itx} \Big|_{-1}^1 + x \int_{-1}^1 f_{\nu} i t e^{-itx} dt \\ &= -x \frac{d}{dt} \widehat{f_{\nu}}, \quad \text{si } \nu - \frac{1}{2} > 0 \end{aligned}$$

d'où

$$(f_{\nu} + tf'_{\nu})^{\wedge} = -x \frac{d}{dx} \widehat{f_{\nu}}.$$

L'équation $(\diamond)$ devient alors

$$2\nu \frac{\widehat{f_{\nu}}}{a_{\nu}} + x \frac{d}{dx} \frac{\widehat{f_{\nu}}}{a_{\nu}} = \frac{\widehat{f_{\nu-1}}}{a_{\nu-1}}$$

$$2\nu x^{-\nu} J_\nu + x \frac{d}{dx} (x^{-\nu} J_\nu) = x^{-\nu+1} J_{-\nu+1}$$

$$2\nu x^{-\nu} J_\nu - \nu x^{-\nu} J_\nu + x^{-\nu+1} J'_\nu = x^{-\nu+1} J_{\nu-1}$$

seconde relation
de récurrence

ou encore

$$J'_\nu = -\frac{\nu}{x} J_\nu + J_{\nu-1}, \quad \nu - \frac{1}{2} > 0. \quad (R2)$$

(Ces relations valent en fait pour tout ν en autant que les fonctions J qui y apparaissent n'aient pas de pôles au point x étudié. Nos preuves ne sont valides que pour le domaine de ν donné.)

Nous avons vu qu'il est possible de définir d'un coup toute une famille de fonctions spéciales à l'aide de leur fonction génératrice. A la section précédente, la famille des polynômes de Legendre a été discutée et, en exercice, plusieurs autres familles. Dans tous ces cas, les polynômes apparaissent comme coefficients dans le développement de Taylor de la fonction génératrice. Pour les fonctions de Bessel d'indice entier, un jeu semblable peut être joué ; maintenant le développement sera plutôt une série de Fourier.

Soit la série de Fourier de

$$e^{ix \sin \theta} = \sum_{n \in \mathbb{Z}} c_n(x) e^{in\theta}.$$

La fonction $e^{ix \sin \theta}$ est une fonction périodique en θ , et ce pour tout x , et analytique en θ puisque la composition de fonctions analytiques l'est également. Elle est donc sûrement différentiable continûment deux fois sur le tore et la convergence de la série sera donc absolue et uniforme. (Voir le théorème 12.)

Calculons le coefficient de Fourier :

$$c_n(x) = \frac{1}{2\pi} \int_{-\pi}^{\pi} e^{ix \sin \theta} e^{-in\theta} d\theta$$

$$= \frac{1}{2\pi} \int_{-\pi}^{\pi} \sum_{k=0}^{\infty} \frac{i^k x^k \sin^k \theta}{k!} e^{-in\theta} d\theta$$

et puisque la série exponentielle converge uniformément sur un domaine borné

$$= \frac{1}{2\pi} \sum_{k=0}^{\infty} \int_{-\pi}^{\pi} \left(\frac{x}{2}\right)^k \frac{1}{k!} (e^{i\theta} - e^{-i\theta})^k e^{-in\theta} d\theta$$

et par le développement binomial

$$= \frac{1}{2\pi} \sum_{k=0}^{\infty} \left(\frac{x}{2}\right)^k \frac{1}{k!} \int_{-\pi}^{\pi} \sum_{l=0}^k \binom{k}{l} (-1)^l e^{i(k-l)\theta} e^{-il\theta} e^{-in\theta} d\theta$$

$$= \frac{1}{2\pi} \sum_{k=0}^{\infty} \left(\frac{x}{2}\right)^k \frac{1}{k!} \int_{-\pi}^{\pi} \sum_{l=0}^k \binom{k}{l} (-1)^l e^{i(k-2l-n)\theta} d\theta$$

dont l'intégrand sera non-nul que lorsque $k = 2l + n$. Si n et l sont fixés avec $n \geq 0$, il y aura précisément une valeur de k qui rend cette intégrale non-nulle. Donc les deux sommes

$$\sum_{k=0}^{\infty} \sum_{l=0}^k$$

seront une somme unique sur les k de la parité de n , ou encore, une somme sur l , $0 \leq l < \infty$, où les k seront remplacés par $2l + n$. Donc, si $n \geq 0$:

$$\begin{aligned} c_n(x) &= \sum_{l=0}^{\infty} \left(\frac{x}{2}\right)^{n+2l} \frac{1}{(n+2l)!} \frac{(n+2l)!}{(n+2l-l)!l!} (-1)^l \\ &= \left(\frac{x}{2}\right)^n \sum_{l=0}^{\infty} \left(\frac{x}{2}\right)^{2l} \frac{(-1)^l}{(n+l)!l!} \\ &= J_n(x). \end{aligned}$$

Si $-n \in \mathbb{N}$, alors la condition $k = 2l + n$ force une autre condition sur k . En effet $-n = 2l - k$ et il faut qu'il existe un l , $0 \leq l \leq k$, tel que cette égalité soit satisfaite. Puisque $2l - k$ prend les valeurs de $-k$ à k , il faut donc que $k \geq -n$, c'est-à-dire $k \geq |n|$. Le restant du calcul est alors identique. Ainsi

fonction génératrice
pour les fonctions de Bessel

$$e^{ix \sin \theta} = \sum_{n \in \mathbb{Z}} J_n(x) e^{in\theta}.$$

Comme pour la fonction $\Gamma(x)$, nous allons décrire le comportement asymptotique de $J_\nu(x)$ lorsque $x \rightarrow \infty$. Nous n'en donnerons cependant pas de preuve. Nous pouvons cependant le justifier. Ecrivons pour cela l'équation différentielle satisfaite par la fonction

$$u_\nu(x) = \sqrt{x} J_\nu(x).$$

Puisque

$$\begin{aligned} J'_\nu &= \frac{u'_\nu}{\sqrt{x}} - \frac{1}{2} \frac{u_\nu}{x^{3/2}} \\ J''_\nu &= \frac{u''_\nu}{\sqrt{x}} - \frac{u'_\nu}{x^{3/2}} + \frac{3}{4} \frac{u_\nu}{x^{5/2}}, \end{aligned}$$

l'équation de Bessel devient

$$\begin{aligned} 0 &= J''_\nu + \frac{1}{x} J'_\nu + \left(1 - \frac{\nu^2}{x^2}\right) J_\nu \\ &= \left(\frac{u''_\nu}{\sqrt{x}} - \frac{u'_\nu}{x^{3/2}} + \frac{3}{4} \frac{u_\nu}{x^{5/2}}\right) + \left(\frac{u'_\nu}{x^{3/2}} - \frac{1}{2} \frac{u_\nu}{x^{5/2}}\right) + \frac{u_\nu}{\sqrt{x}} - \frac{\nu^2 u_\nu}{x^{5/2}} \\ &= \frac{1}{\sqrt{x}} \left(u''_\nu + \left(1 + \frac{\frac{1}{4} - \nu^2}{x^2}\right) u_\nu\right) \end{aligned}$$

et donc

$$0 = u''_{\nu} + \left(1 + \frac{\frac{1}{4} - \nu^2}{x^2}\right) u_{\nu}.$$

Lorsque $x \rightarrow \infty$, le terme en $\frac{1}{x^2}$ devient négligeable et il est raisonnable de croire que u_{ν} vérifiera l'équation différentielle pour les fonctions trigonométriques : $u''_{\nu} + u_{\nu} \approx 0$. Le comportement exact est

$$J_{\nu}(x) = \sqrt{\frac{2}{\pi x}} \cos\left(x - (2\nu + 1)\frac{\pi}{4}\right) + \mathcal{O}\left(\frac{1}{|x|^{\frac{3}{2}}}\right).$$

comportement
asymptotique de J_{ν}

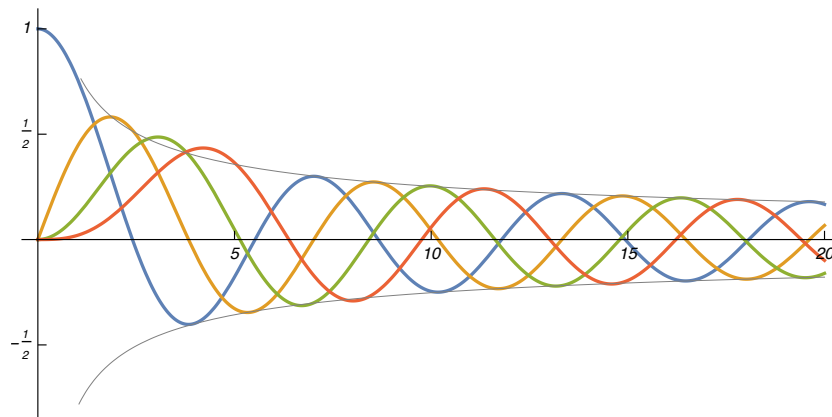


Figure 3.4 – Les fonctions de Bessel J_0 , J_1 , J_2 et J_3

Les premières fonctions de Bessel sont tracées sur la figure ci-contre. Les quatre fonctions de Bessel J_0 , J_1 , J_2 et J_3 y sont tracées ainsi qu'en trait plus fin, l'enveloppe asymptotique $\pm \sqrt{\frac{2}{\pi x}}$. La fonction J_0 est égale à 1 en $x = 0$. Pour distinguer les trois autres, il suffit de savoir que leur comportement autour de $x = 0$ est caractérisé par une pente d'autant plus faible que l'indice du J est grand. On remarquera que pour x suffisamment élevé, l'enveloppe asymptotique agit vraiment comme une « gaine » et que les zéros des J tracées tombent très proches de leurs valeurs prédites puisque les premiers zéros de $\cos(x + \frac{\pi}{4})$ sont approximativement $\{0.79, 3.9, 7.1, 10.21, 13.35, 16.49, 19.63\}$ alors que ceux de $\cos(x + \frac{3\pi}{4})$ sont $\{2.36, 5.50, 8.64, 11.78, 14.92, 18.06\}$.

Exemple. Nous allons clore cette section par la solution de l'équation du tambour rond de rayon 1. Ce problème simple, en apparence, fera appel à tout ce que nous avons vu depuis le début de ce cours : la séparation des variables pour une équation aux dérivées partielles, la solution des équations ordinaires obtenues que nous

traiterons comme problèmes de Sturm-Liouville et le développement de la solution générale en termes des fonctions propres de ces problèmes, en l'occurrence, les fonctions trigonométriques et de Bessel. Nous avons déjà franchi quelques étapes. En effet, l'équation du tambour est l'équation des ondes en deux dimensions :

$$\frac{\partial^2 u}{\partial t^2} = a^2 \left(\frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} \right)$$

où u mesure le déplacement transversal du tambour et a est une constante que nous poserons par la suite à 1. La variable de temps est t et les variables d'espace x et y . Puisque le tambour est rond, les coordonnées polaires s'imposent et l'équation prend la forme :

$$\frac{\partial^2 u}{\partial t^2} = \frac{\partial^2 u}{\partial \rho^2} + \frac{1}{\rho} \frac{\partial u}{\partial \rho} + \frac{1}{\rho^2} \frac{\partial^2 u}{\partial \theta^2}$$

où ρ joue le rôle du rayon et θ de l'angle polaire. Nous avons vu à la section 3.2 que l'équation d'onde se sépare en coordonnées cylindriques (en trois dimensions spatiales). Ici l'équation est en deux coordonnées spatiales mais les calculs faits précédemment s'appliquent pratiquement sans changements, seule la dépendance en z est ignorée. Si la fonction $u(t, \rho, \theta)$ prend la forme $T(t)R(\rho)\Theta(\theta)$, les résultats précédents donnent les trois équations séparées suivantes :

$$\begin{aligned} T'' &= -\lambda^2 T \\ \Theta'' &= -\mu^2 \Theta \\ \rho^2 R'' &= -\rho R' - (\lambda^2 \rho^2 - \mu^2) R \end{aligned}$$

où λ et μ sont les constantes de séparation. Chacune de ces équations peut être interprétées comme un problème de Sturm-Liouville. Les deux premiers sont cependant communs, menant les deux à des solutions trigonométriques :

$$\begin{aligned} T &= a_1 \cos \lambda t + a_2 \sin \lambda t \\ \Theta &= b_1 \cos \mu \theta + b_2 \sin \mu \theta. \end{aligned}$$

Dans le second cas, pour assurer la continuité de la solution Θ (et donc de u), il faudra restreindre $\mu \in \mathbb{N} \cup \{0\}$. Par la suite, nous utiliserons la lettre $m = \mu$ pour cet entier. Reste donc

$$\rho R'' + R' + \left(\lambda^2 \rho - \frac{m^2}{\rho} \right) R = 0$$

ou encore, pour mieux épouser la forme de Sturm-Liouville :

$$\frac{d}{d\rho} \left(\rho \frac{dR}{d\rho} \right) + \left(\lambda^2 \rho - \frac{m^2}{\rho} \right) R = 0$$

et donc les fonctions p, r et s introduites à la section 3.3 sont

$$p(\rho) = \rho, \quad r(\rho) = \rho \quad \text{et} \quad s(\rho) = \frac{m^2}{\rho}.$$

Ainsi, pour chaque valeur de m intervenant dans l'équation pour Φ , nous avons un problème de Sturm-Liouville différent pour les fonctions propres radiales. (Ceci est une possibilité que nous rencontrons pour la première fois.)

Quelles sont les solutions de l'équation radiale qui ressemble tant à l'équation de Bessel ? Essayons :

$$R(\rho) = J_m(\lambda\rho).$$

Alors

$$\frac{dR}{d\rho} = \lambda J'_m(\lambda\rho) \quad \text{et} \quad \frac{d^2R}{d\rho^2} = \lambda^2 J''_m(\lambda\rho)$$

et alors

$$\begin{aligned} R''(\rho) + \frac{1}{\rho}R'(\rho) + \left(\lambda^2 - \frac{m^2}{\rho^2}\right)R(\rho) &= \\ &= \lambda^2 J''_m(\lambda\rho) + \frac{\lambda}{\rho}J'_m(\lambda\rho) + \left(\lambda^2 - \frac{m^2}{\rho^2}\right)J_m(\lambda\rho) \\ &= \lambda^2 \left(J''_m(\lambda\rho) + \frac{1}{(\lambda\rho)}J'_m(\lambda\rho) + \left(1 - \frac{m^2}{(\lambda\rho)^2}\right)J_m(\lambda\rho) \right) \\ &= 0, \end{aligned}$$

les fonctions
propres radiales

la dernière étape venant du fait que la ligne précédente est précisément l'équation de Bessel évaluée au point $(\lambda\rho)$. Donc les solutions sont

$$R(\rho) = J_m(\lambda\rho).$$

Quel est le spectre ? Il sera fixé par les conditions aux limites ! Nous devons avoir

la fonction R est régulière en $\rho = 0$

car la peau du tambour est lisse en $\rho = 0$ (elle ne va pas à l'infini...) et

$$R(\rho = 1) = 0$$

car la peau du tambour est fixé en son pourtour. La première condition permet de rejeter les fonctions de Neumann N_m qui sont singulières en $\rho = 0$, chose que nous avons déjà fait implicitement en écrivant $R(\rho) = J_m(\lambda\rho)$ plutôt que $R(\rho) = aJ_m(\lambda\rho) + bN_m(\lambda\rho)$. La seconde donne une relation implicite pour les valeurs propres de l'équation aux valeurs propres (qu'est le problème de Sturm-Liouville) :

$$R(1) = J_m(\lambda) = 0.$$

Si on numérote les zéros positifs de J_m par α_{mn} , $n \in \mathbb{N}$, alors les fonctions propres du problème de Sturm-Liouville sont

$$R_{mn}(\rho) = J_m(\alpha_{mn}\rho).$$

Le problème de Sturm-Liouville nous assure que, deux solutions R_1 et R_2 de l'équation radiale correspondant à des valeurs propres distinctes λ_1 et λ_2 seront orthogonales par rapport à :

$$(R_1, R_2) = \int_0^1 \rho R_1(\rho) R_2(\rho) d\rho = 0$$

puisque le rayon du tambour est 1. La preuve d'orthogonalité des fonctions propres à la section 3.3 reposait sur certains types de conditions aux limites. Les fonctions de Bessel offrent un exemple des conditions que nous avons appelées *mixtes*. En $\rho = 0$, la fonction $p(\rho) = \rho$ s'annule. En $\rho = 1$, ce sont les fonctions propres qui s'annulent : $R_{mn}(\rho = 1) = 0$. Et donc

$$(R_{mn}, R_{ml}) = \int_0^1 J_m(\alpha_{mn}\rho) J_m(\alpha_{ml}\rho) \rho d\rho = 0, \quad \text{si } n \neq l.$$

Le spectre du tambour, c'est-à-dire l'ensemble des fréquences du tambour, est donc spectre du
tambour rond

$$\{\alpha_{mn}, m \in \mathbb{N} \cup \{0\}, n \in \mathbb{N}\}$$

c'est-à-dire l'ensemble de tous les zéros positifs des fonctions de Bessel J_m à indice m entier positif ou nul.

Quelle est la normalisation des fonctions propres $R_{mn}(\rho) = J_m(\alpha_{mn}\rho)$? Afin de répondre à cette question, nous prendrons un détour ; nous allons d'abord calculer le produit de $y_1(\rho) = J_m(\lambda\rho)$ et de $y_2 = J_m(\mu\rho)$, même si λ et μ ne sont pas des zéros de J_m : $(y_1, y_2) = \int_0^1 \rho J_m(\lambda\rho) J_m(\mu\rho) d\rho$. Pour faire cette intégrale, employons le truc que nous avons employé pour montrer l'orthogonalité des fonctions propres à la section 3.3. Ecrivons les équations satisfaites par chacune de ces fonctions et multiplions-les par l'autre fonction :

$$\begin{aligned} y_2 \times & \left(\frac{d}{d\rho}(\rho y_1') + (\lambda^2 \rho - \frac{m^2}{\rho}) y_1 = 0 \right) \\ y_1 \times & \left(\frac{d}{d\rho}(\rho y_2') + (\mu^2 \rho - \frac{m^2}{\rho}) y_2 = 0 \right) \end{aligned}$$

Soustrayons et intégrons de 0 à 1 :

$$\int_0^1 \left(y_2 \frac{d}{d\rho}(\rho y_1') - y_1 \frac{d}{d\rho}(\rho y_2') + \rho(\lambda^2 - \mu^2) y_1 y_2 \right) d\rho = 0.$$

Donc, si $\lambda \neq \mu$:

$$\int_0^1 \rho y_1 y_2 d\rho = \frac{1}{\mu^2 - \lambda^2} \int_0^1 \left(y_2 \frac{d}{d\rho}(\rho y_1') - y_1 \frac{d}{d\rho}(\rho y_2') \right) d\rho$$

qui, par intégration par parties, devient

$$\begin{aligned} &= \frac{1}{\mu^2 - \lambda^2} \left[(\rho y_2 y_1' - \rho y_1 y_2') \Big|_0^1 - \int_0^1 \rho (y_1' y_2' - y_2' y_1') d\rho \right] \\ &= \frac{1}{\mu^2 - \lambda^2} (\lambda J_m(\mu) J_m'(\lambda) - \mu J_m(\lambda) J_m'(\mu)). \end{aligned}$$

(Les λ et μ sont apparus car $y_1'(\rho) = \lambda J_m'(\lambda\rho)$.) Il est clair alors que si λ et μ sont des zéros distincts de J_m , y_1 et y_2 sont orthogonales. Pour trouver la normalisation de $J_m(\lambda\rho)$, notons que $J_m(\lambda\rho)J_m(\mu\rho)$ sont des fonctions continues de λ et μ puisque que ces fonctions sont définies par séries de Taylor qui, sur des intervalles bornés, convergent uniformément. Un théorème de Lebesgue-Riemann dit que

$$\int_0^1 \lim_{\lambda \rightarrow \mu} \rho y_1 y_2 d\rho = \lim_{\lambda \rightarrow \mu} \int_0^1 \rho y_1 y_2 d\rho$$

si l'intégrand est borné. Ainsi

$$\int_0^1 \rho J_m(\mu\rho)^2 d\rho = \lim_{\lambda \rightarrow \mu} \frac{1}{\mu^2 - \lambda^2} (\lambda J_m(\mu) J_m'(\lambda) - \mu J_m(\lambda) J_m'(\mu))$$

et par la règle de l'Hospital :

$$= \frac{1}{2\mu} (\mu J_m'(\mu)^2 - J_m(\mu) J_m'(\mu) - \mu J_m(\mu) J_m''(\mu)).$$

Et si, comme dans le cas qui nous intéresse, μ est un zéro de J_m , alors :

$$\begin{aligned} \int_0^1 \rho J_m(\alpha_{mn}\rho)^2 d\rho &= \frac{1}{2} J_m'(\alpha_{mn})^2 \\ &= \frac{1}{2} J_{m+1}(\alpha_{mn})^2 \end{aligned}$$

où nous avons utilisé la relation de récurrence (R1). Ouf ! Avec cette normalisation, on peut donc écrire des séries de fonctions de Bessel J_m . Par exemple, il est possible d'écrire

$$f(x) = \sum_{n=1}^{\infty} a_n J_m(\alpha_{mn}x),$$

où a_n est déterminé par la relation obtenue à la section 3.3

$$a_n = \frac{(f, y_n)}{(y_n, y_n)}$$

avec $y_n(x) = J_m(\alpha_{mn}x)$ et donc :

$$= \frac{2}{J_{m+1}(\alpha_{mn})^2} \int_0^1 x f(x) J_m(\alpha_{mn}x) dx.$$

Nous pouvons mettre maintenant les pièces du casse-tête ensemble pour écrire la solution de l'équation du tambour. Si la membrane est lâchée sans vitesse, il faut que $T'(t=0) = 0$ et donc que

$$T'(0) = (-a_1 \lambda \sin \lambda t + a_2 \cos \lambda t)|_0 = 0,$$

c'est-à-dire $a_2 = 0$. La solution générale sera de la forme

normalisation
des J_m

solution générale

$$u(\rho, \theta, t) = \sum_{m=0}^{\infty} \sum_{p=1}^{\infty} \cos(\alpha_{mp} t) (a_{mp} \cos m\theta + b_{mp} \sin m\theta) J_m(\alpha_{mp} \rho).$$

Comment sont déterminés les a_{mp} et b_{mp} si en $t = 0$

$$u(\rho, \theta, 0) = f(\rho, \theta)?$$

Pour un ρ fixé, on a

$$f(\rho, \theta) = \sum_{m=0}^{\infty} \sum_{p=1}^{\infty} (a_{mp} \cos m\theta + b_{mp} \sin m\theta) J_m(\alpha_{mp} \rho)$$

qui est une fonction de θ de période 2π (quelque soit ρ) et alors

$$\frac{1}{1 + \delta_{m0}} \frac{1}{\pi} \int_{-\pi}^{\pi} f(\rho, \theta) \cos m\theta d\theta = \sum_{p=1}^{\infty} a_{mp} J_m(\alpha_{mp} \rho), \quad m \geq 0$$

où le facteur $\frac{1}{1 + \delta_{m0}}$ est 1 pour tout $m > 0$ et $\frac{1}{2}$ pour $m = 0$, la normalisation de $\cos m\theta$ étant différente lorsque $m = 0$. Similairement

$$\frac{1}{\pi} \int_{-\pi}^{\pi} f(\rho, \theta) \sin m\theta d\theta = \sum_{p=1}^{\infty} b_{mp} J_m(\alpha_{mp} \rho), \quad m \geq 1.$$

Si nous utilisons maintenant l'orthogonalité des fonctions propres radiales, les coefficients a_{mp} et b_{mp} peuvent être obtenus à partir d'une autre intégration :

$$a_{mp} = \frac{2}{\pi J_{m+1}(\alpha_{mp})^2} \frac{1}{1 + \delta_{m0}} \int_0^1 \rho d\rho \int_{-\pi}^{\pi} d\theta f(\rho, \theta) \cos m\theta J_m(\alpha_{mp} \rho)$$

pour $m \geq 0, p \geq 1$ et

$$b_{mp} = \frac{2}{\pi J_{m+1}(\alpha_{mp})^2} \int_0^1 \rho d\rho \int_{-\pi}^{\pi} d\theta f(\rho, \theta) \sin m\theta J_m(\alpha_{mp} \rho)$$

pour $m > 0, p \geq 1$. Voilà !

Quelques commentaires moins mathématiques maintenant. Nous avons vu que les zéros de la fonction de Bessel J_m se comportent asymptotiquement comme $\sim n\pi + (m+1)\frac{\pi}{2} + \frac{\pi}{4}$. Les fréquences du tambour rond sont donc (approximativement) espacées par un certain multiple de π . C'est donc un spectre qui ressemble qualitativement au spectre de la corde vibrante qui est lui (précisément) $\{n\omega_0, n \in \mathbb{N}\}$ où ω_0 est la fréquence la plus basse dite la fondamentale. La différence repose donc dans le mot « approximativement ». Le fait qu'il soit difficile de chanter le son d'un tambour est une conséquence du fait que les fréquences naturelles ne sont pas précisément des multiples de la fréquence fondamentale, $\alpha_{01} \approx 2.405 \dots$ lorsque toutes les constantes physiques ont été posées à 1 comme nous l'avons fait ici.

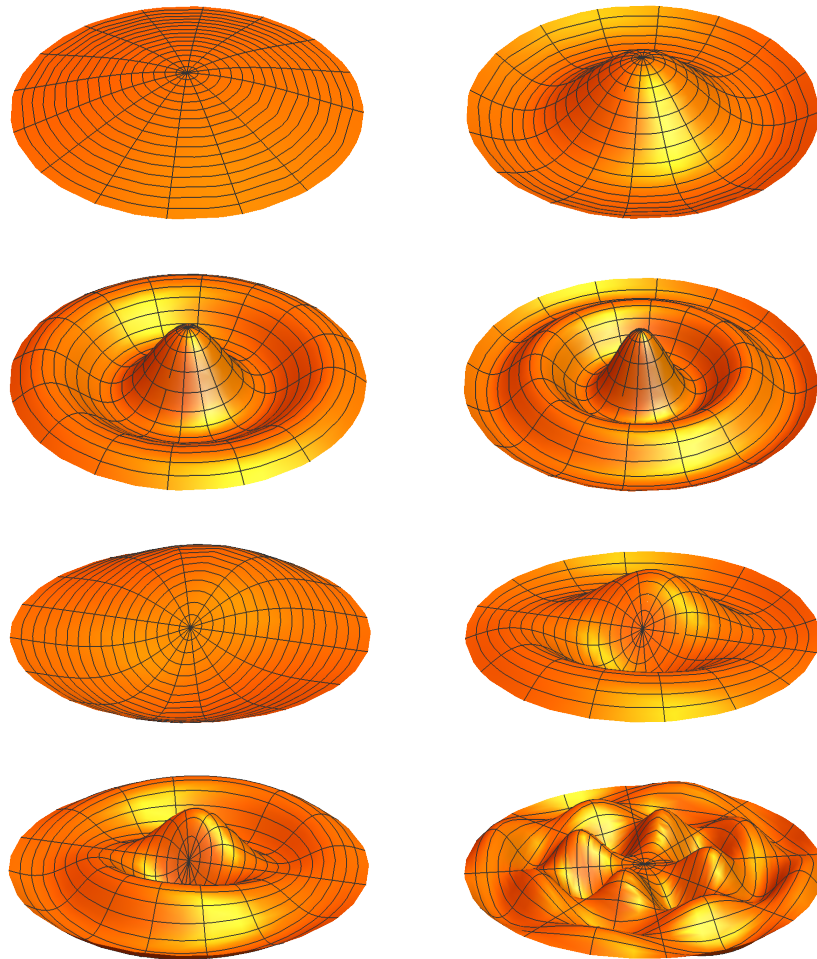


Figure 3.5 – Quelques modes propres du tambour rond

Les modes propres du tambour rond sont donc le produit des fonctions propres des équations pour T , pour Θ et pour R . Il est donc possible de dessiner pour t fixé le graphe d'un mode donné. C'est ce que font les dessins de la figure 3.5. Les modes sont étiquetés par la paire d'entiers (m, p) . Ceux représentés sont $(m, p) = (0, 1), (0, 2), (0, 3), (0, 4), (1, 1), (1, 2), (1, 3), (1, 4)$ lus de gauche à droite puis de haut en bas. Le premier indice réfère à la fonction de Bessel J_m . La fonction $\Theta(\theta)$ correspondant à m est une combinaison de sin et de cos d'argument $m\theta$. Ainsi, plus m est grand plus le nombre d'oscillations le long d'un cercle de rayon constant est grand. Les quatre premiers dessins n'ont pas de dépendance angulaire ($m = 0$). Le long d'un cercle, la fonction propre est constante. Lorsque $m = 1$,

la fonction propre le long d'un cercle aura précisément un creux et une bosse. Le dernier dessin a $m = 4$ et pour ce mode, il y aura 4 creux alternant avec 4 bosses. Le second indice p compte le nombre de zéros le long d'un rayon (à θ constant). Ainsi plus ce nombre est grand, plus il y aura d'oscillations entre le centre et le bord du tambour.

26*. Obtenir les solutions de l'équation $y'' = -m^2 y$ sous la forme $y = \sum_{i=0}^{\infty} a_i x^i$ par la méthode de Frobenius. Vous devriez trouver deux solutions : une pour laquelle $a_0 = 1$ et $a_1 = 0$ et une autre pour laquelle $a_0 = 0$ et $a_1 = 1$. **Exercices**

27*. (a) Montrer que $J_{\frac{1}{2}}$ a la forme simple suivante

$$J_{\frac{1}{2}}(x) = \sqrt{\frac{2}{\pi x}} \sin x.$$

(b) Montrer que

$$J_{-\frac{1}{2}}(x) = \sqrt{\frac{2}{\pi x}} \cos x$$

et, plus généralement, que les $J_{\pm(n+\frac{1}{2})}$, $n \in \mathbb{N} \cup \{0\}$, peuvent être exprimés comme une somme finie de fonctions $\cos x$ et $\sin x$ pondérées par des puissances demi-entières de x .

(c) Obtenir (rigoureusement) le comportement asymptotique de $J_{n+\frac{1}{2}}$ pour $n \in \mathbb{N} \cup \{0\}$.

28. En se rappelant que $J_p(x) = (-1)^p J_{-p}(x)$, montrer que

$$\begin{aligned} \cos(x \sin \theta) &= J_0(x) + 2 \sum_{n=1}^{\infty} J_{2n}(x) \cos 2n\theta \\ \sin(x \sin \theta) &= 2 \sum_{n=0}^{\infty} J_{2n+1}(x) \sin(2n+1)\theta. \end{aligned}$$

29*. Ecrire l'identité de Parseval pour les fonctions de Bessel intervenant dans l'analyse du tambour rond.

30. En utilisant l'identité de Parseval pour les séries trigonométriques, montrer que

$$1 = J_0(x)^2 + 2 \sum_{k=1}^{\infty} J_k(x)^2$$

pour tout x . (Ceci a comme conséquence immédiate que $|J_0(x)| \leq 1$ et $|J_k(x)| \leq \frac{1}{\sqrt{2}}$, $k \in \mathbb{N}$, pour tout x .)

31*. Montrer que

$$\frac{\sin x}{x} = \int_0^{\frac{\pi}{2}} J_0(x \cos \theta) \cos \theta d\theta$$

et

$$\frac{1 - \cos x}{x} = \int_0^{\frac{\pi}{2}} J_1(x \cos \theta) d\theta.$$

32. (a) Montrer que

$$x = 2 \sum_{p=1}^{\infty} \frac{J_1(\lambda_p x)}{\lambda_p J_2(\lambda_p)}$$

pour $x \in (0, 1)$ où λ_p sont les racines positives de $J_1(\lambda) = 0$. Suggestion : relation de récurrence !

(b) Montrer que

$$\sum_{p=1}^{\infty} \frac{1}{\lambda_p^2} = \frac{1}{8}.$$

33. Dans la solution du tambour rond, nous avons montré que les fonctions propres $J_m(\alpha_{mn}\rho)$ et $J_m(\alpha_{ml}\rho)$ avec α_{mn} et α_{ml} deux zéros distincts de J_m sont orthogonales. Est-ce que la fonction $J_p(\alpha_{pq}\rho)$ (qui est également propre du problème de Sturm-Liouville, mais un $p \neq m$) est orthogonale à $J_m(\alpha_{mn}\rho)$? Cette question est importante. Si ces fonctions sont orthogonales, il faut le montrer. Si elles ne le sont pas, il semble y avoir un problème : les modes propres du tambour ne seront pas orthogonaux ?

34. Vérifier les relations de récurrence (R1) et (R2) en utilisant le développement en série des J_ν .

35*. Montrer que, quels que soient $m \in \mathbb{N}$ et $x \in [0, \infty)$, il existe $0 < M < \infty$ tel que

$$|J_n(x)| \leq \frac{M}{n^m}, \quad \forall n \in \mathbb{N} \cup \{0\}.$$

La constante M peut dépendre de x . Est-ce que ce résultat contredit le comportement asymptotique décrit dans le texte ?

36*. Evaluer l'intégrale

$$\int_0^{\infty} e^{-at} J_0(bt) dt.$$

37*. (a) Une plaque circulaire de rayon 1 est isolée thermiquement, sauf en son pourtour où elle est maintenue à température nulle. Si, au temps $t = 0$, la température du disque est donnée par la fonction $f(\rho, \theta)$, calculer l'évolution de la

température $u(t, \rho, \theta)$ en fonction du temps. (L'évolution est décrite par l'équation aux dérivées partielles : $\frac{\partial u}{\partial t} = \nabla^2 u$.)

(b) En quoi la solution ci-dessus se simplifie-t-elle si la température initiale ne dépend que de ρ : $f = f(\rho)$?

38. (a) Un cylindre de rayon et de hauteur unité a sa température initiale décrite par $f(\rho, \theta, z)$. Soudainement sa surface est mise à température nulle. Décrire sa température $u(t, \rho, \theta, z)$ au cours du temps.

(b) Trouver la solution stationnaire de l'équation de la chaleur (c'est-à-dire une solution ne dépendant pas du temps) pour le cylindre de l'exercice précédent si toute sa surface est maintenue à température nulle à l'exception de sa surface supérieure (le « couvercle ») qui est à température u_0 .

(c) Supposons que la température initiale du cylindre soit décrite par $f(\rho, \theta, z)$ et que sa surface soit maintenue à température nulle sauf pour le disque supérieur qui est à température u_0 . Ecrire l'évolution de la température $u(t, \rho, \theta, z)$ pour ces conditions initiales et aux limites.

Note : Ecrire $u(t, \rho, \theta, z) = v(t, \rho, \theta, z) + w(\rho, z)$ où v est une solution du problème (a) et w est stationnaire.

39*. Soit une cavité cylindrique de rayon a et longueur l . Pour les ondes transverses électriques, la composante longitudinale B du champ magnétique satisfait au problème aux valeurs propres

$$\nabla^2 B + \omega^2 B = 0$$

avec conditions aux limites

$$B(\rho, \theta, z = 0) = B(\rho, \theta, z = l) = 0 \quad \text{et} \quad \frac{\partial B}{\partial \rho}(\rho = a, \theta, z) = 0.$$

Montrer que les valeurs propres peuvent être étiquetées par trois entiers et ont la forme :

$$\omega_{mnp} = \sqrt{\frac{\beta_{mn}^2}{a^2} + \frac{p^2 \pi^2}{l^2}}.$$

Dire ce que sont les β_{mn} et préciser les valeurs possibles de m, n et p .

40*. Montrer que, pour tout entier n :

$$J_n(a + b) = \sum_{m \in \mathbb{Z}} J_m(a) J_{n-m}(b).$$

Chapitre 4

La transformation de Fourier

4.1 Un passage à la limite et exemples de calcul

Existe-t-il un équivalent des séries de Fourier pour les fonctions qui ne sont pas périodiques ? La réponse à cette question est « oui » et le but de ce chapitre est de développer la réponse. La transformation de Fourier sera introduite de façon intuitive en prenant la limite des expressions définissant les séries de Fourier lorsque la période tend vers l'infini.

Rappelons que si

$$f : [-T/2, T/2] \subset \mathbb{R} \rightarrow \mathbb{C}$$

est continue et dérivable continûment deux fois, alors l'égalité suivante est valide :

$$f(x) = \sum_{n=-\infty}^{\infty} c_{n,T} e^{2i\pi nx/T}, \quad x \in [-T/2, T/2]$$

où

$$c_{n,T} = \frac{1}{T} \int_{-T/2}^{T/2} f(x) e^{-2i\pi nx/T} dx.$$

J'ai ajouté un indice « T » supplémentaire aux coefficients c_n car nous nous apprêtons à varier la période T et il y aura donc plusieurs familles de coefficients c_n .

Notre but est de prendre la limite $T \rightarrow \infty$. Si cela est fait un peu trop rapidement, la série de Fourier devient une somme sur les seuls coefficients $c_{n,T}$ puisque l'exponentielle voit son argument devenir zéro. De plus, les coefficients $c_{n,T}$ deviennent eux-mêmes zéro si la fonction f est intégrable puisque

$$c_{n,T} = \frac{1}{T} \int_{-T/2}^{T/2} f(x) e^{-2i\pi nx/T} dx \xrightarrow{T \rightarrow \infty} \frac{1}{T} \int_{-\infty}^{\infty} f(x) e^0 dx = \frac{M}{T}$$

pour un certain nombre fini M . Ainsi cette façon de faire est inutile (et fausse).

Pour résoudre cette difficulté, remplaçons l'indice de sommation n par un nombre qui demeurera fini durant le processus limite. Ceci est justifié puisque $n \in \mathbb{Z}$ et, puisque n prend des valeurs aussi grandes que désiré, l'argument de l'exponentielle demeure non-nul même si T est très grand. Considérons donc la quantité

$$\lambda = \frac{n}{T}$$

qui devra être également sommée de $-\infty$ à $+\infty$, quelque soit T . Regroupons les termes qui possèdent des λ très voisins :

$$2\pi\lambda \leq \frac{2\pi}{T}n \leq 2\pi(\lambda + \Delta\lambda).$$

Le nombre de valeurs de l'indice n dans cet intervalle est donc

$$\Delta n \approx T\Delta\lambda.$$

Pour des λ très voisins (c'est-à-dire un $\Delta\lambda \approx 0$), la somme sur les valeurs de n dans l'intervalle ci-dessus peut-être réécrite comme

$$\sum_{2\pi\lambda \leq \frac{2\pi n}{T} \leq 2\pi(\lambda + \Delta\lambda)} c_{n,T} e^{\frac{2\pi i n x}{T}} \approx c_\lambda e^{2\pi i \lambda x} T \Delta\lambda$$

où

$$c_\lambda = \frac{1}{T} \int_{-T/2}^{T/2} f(x) e^{-\frac{2\pi i n x}{T}} dx.$$

La forme approximée de cette somme peut être justifiée par le fait que, sur un petit intervalle $\Delta\lambda$ et pour T suffisamment grand, les coefficients c_λ et les exponentielles $e^{2\pi i n x/T}$ devraient être à peu près égales. En posant

$$c(\lambda) = T c_\lambda = \int_{-T/2}^{T/2} f(x) e^{-\frac{2\pi i n x}{T}} dx,$$

la somme donnant $f(x)$ devient

$$f(x) = \sum_{\lambda=-\infty}^{\infty} c(\lambda) e^{2\pi i \lambda x} \Delta\lambda \xrightarrow{\Delta\lambda \rightarrow 0} \int_{-\infty}^{\infty} c(\lambda) e^{2\pi i \lambda x} d\lambda$$

avec

$$c(\lambda) = \int_{-\infty}^{\infty} f(x) e^{-2\pi i \lambda x} dx.$$

On peut faire un changement de variables pour supprimer les 2π . En posant $\omega = 2\pi\lambda$, on obtient dans l'intégrale donnant f :

$$f(x) = \frac{1}{2\pi} \int_{-\infty}^{\infty} c\left(\frac{\omega}{2\pi}\right) e^{i\omega x} d\omega$$

et il est alors naturel de définir un nouveau coefficient $\hat{f}(\omega)$ par

$$\hat{f}(\omega) = c \left(\frac{\omega}{2\pi} \right) = \int_{-\infty}^{\infty} f(x) e^{-i\omega x} dx.$$

Cette forme est celle que nous utiliserons pour définir la transformée de Fourier de la fonction $f : \mathbb{R} \rightarrow \mathbb{R}$ qui, puisque la période est maintenant infinie, a comme domaine l'axe des x en entier. Notre définition sera donc

$$f(x) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\lambda) e^{i\lambda x} d\lambda \quad \text{où} \quad \hat{f}(\lambda) = \int_{-\infty}^{\infty} f(x) e^{-i\lambda x} dx. \quad (\mathcal{F})$$

transformation
de Fourier

Nous avons dit « notre » définition car, contrairement aux séries de Fourier qui sont les mêmes dans tous les livres, les conventions utilisées pour la définition de la transformation de Fourier sont nombreuses. Les définitions varient sur la valeur du coefficient numérique devant les deux intégrales (dont les valeurs les plus communes sont 1 , $\frac{1}{2\pi}$ et $\frac{1}{\sqrt{2\pi}}$) et la position du signe « $-$ » dans l'argument des exponentielles (se trouve-t-il dans la définition de f ou de $\hat{f}$). Le lecteur doit être conscient de ces différences s'il fait des exercices dans plusieurs livres.¹ Un autre commentaire est utile, cette fois de nature linguistique. Vous trouverez souvent les mots « transformation » et « transformée » de Fourier dans ce qui suit. Ils ne sont pas tout à fait synonymes ; le premier est à l'intégration ce que le second est à l'intégrale, c'est-à-dire le premier est l'opération, le second le résultat.

Nous pourrions immédiatement soulever les questions que nous avons posées pour les séries de Fourier :

- si $\hat{f}(\lambda)$ est donné par l'intégrale ci-dessus, est-ce que $f(x)$ égale

$$\frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\lambda) e^{i\lambda x} d\lambda?$$

- donné $\hat{f}(\lambda)$, peut-on reconstruire $f(x)$?
- et existe-t-il un concept de distance comme pour les séries de Fourier qui permette de mesurer la distance entre f et l'intégrale où $\hat{f}(\lambda)$ apparaît ?

En plus, d'autres questions s'ajoutent à ces premières :

- les intégrales $\int_{-\infty}^{\infty} \dots$ convergent-elles ?
- $\hat{f}$ est-elle continue ? sous quelles conditions ?
- $\hat{f}$ est-elle dérivable ? Et si oui, peut-on obtenir sa dérivée en passant la dérivée $\frac{d}{d\lambda}$ sous le signe de l'intégrale ?
- et plus viscéralement, à quoi sert la transformation de Fourier ?

1. Notre convention est celle des livres de Körner et de Spiegel (série Schaum). Elle est différente de celle de Schwartz et Arfken.

Nous ne répondrons pas à toutes ces questions. Seules quelques-unes recevront une réponse. Mais avant, voici des exemples de calcul !

Exemple 1. Calculer la transformée de Fourier de la fonction

$$f(x) = \begin{cases} 1, & |x| \leq 1 \\ 0, & \text{autrement.} \end{cases}$$

Le calcul est direct :

$$\begin{aligned} \hat{f}(\omega) &= \int_{-\infty}^{\infty} f(x) e^{-i\omega x} dx \\ &= \int_{-1}^1 e^{-i\omega x} dx \\ &= \left. \frac{e^{-i\omega x}}{-i\omega} \right|_{-1}^1 \\ &= -\frac{1}{\omega} \frac{e^{-i\omega} - e^{i\omega}}{i} \\ &= \frac{2}{\omega} \sin \omega. \end{aligned}$$

Ce calcul n'est évidemment pas valide en $\omega = 0$. Cependant il est aisé d'obtenir cette valeur particulière : $\hat{f}(0) = \int_{-\infty}^{\infty} f(x) dx = \int_{-1}^1 dx = 2$ qui est la limite de $\hat{f}(\omega)$ lorsque $\omega \rightarrow 0$.

Exemple 2. Si $f(x) = e^{-a|x|}$ avec a positif, alors

$$\begin{aligned} \hat{f}(\omega) &= \int_{-\infty}^{\infty} e^{-a|x|} e^{-i\omega x} dx \\ &= \int_{-\infty}^0 e^{-a|x|} e^{-i\omega x} dx + \int_0^{\infty} e^{-a|x|} e^{-i\omega x} dx \\ &= \int_0^{\infty} e^{-a|y|} e^{i\omega y} dy + \int_0^{\infty} e^{-a|x|} e^{-i\omega x} dx \\ &= 2 \operatorname{Re} \int_0^{\infty} e^{-ax} e^{-i\omega x} dx = 2 \operatorname{Re} \int_0^{\infty} e^{-x(a+i\omega)} dx \\ &= 2 \operatorname{Re} \left. -\frac{e^{-x(a+i\omega)}}{a+i\omega} \right|_0^{\infty} = 2 \operatorname{Re} \frac{1}{a+i\omega} \cdot \frac{a-i\omega}{a-i\omega} \\ &= \frac{2a}{a^2 + \omega^2}. \end{aligned}$$

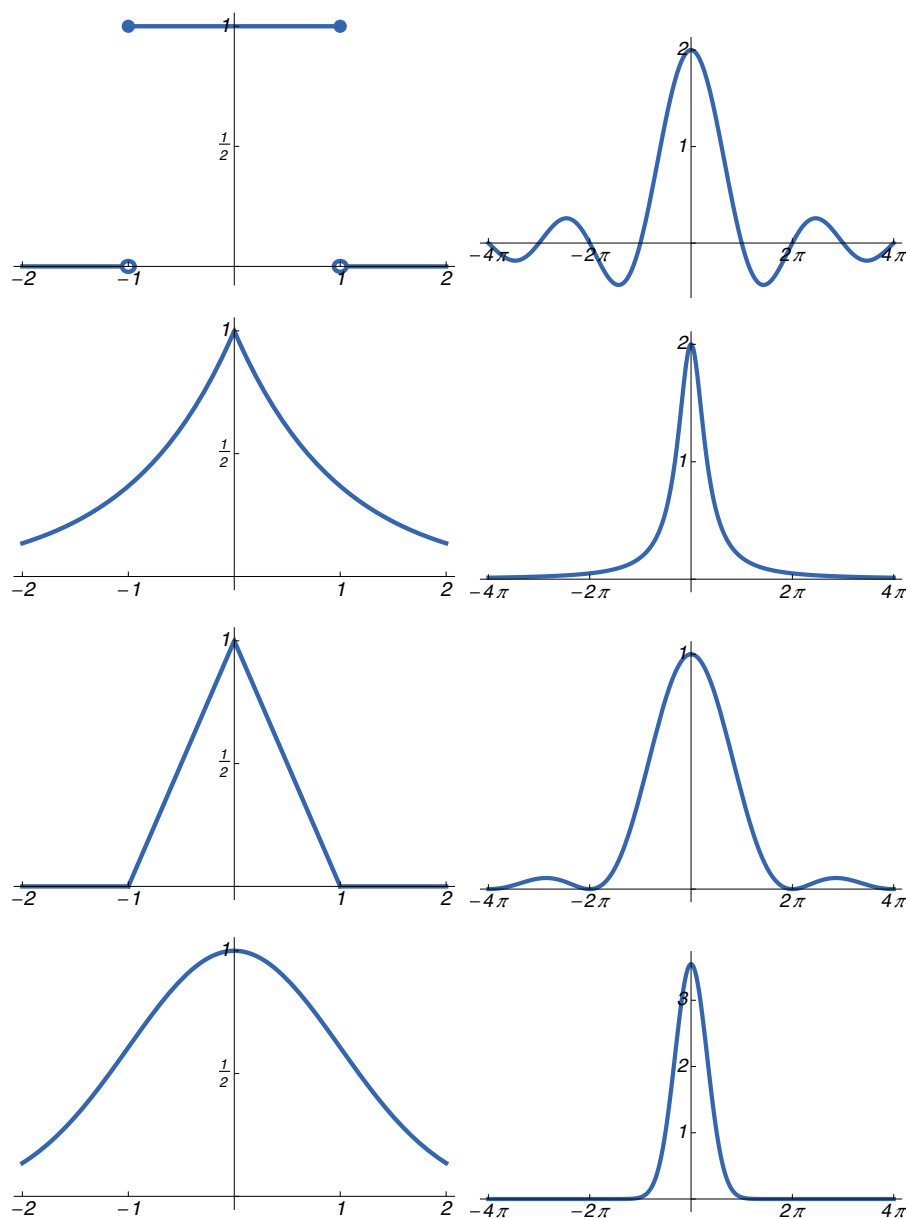


Figure 4.1 – Le graphe des fonctions des quatre exemples de cette section (à gauche) et de leur transformée de Fourier (à droite). De bas en haut, les exemples 1, 2, 3 et 4. Les constantes a et h ont été posées à 1.

Exemple 3. Voici encore une fonction dont la transformée est relativement aisée à calculer :

$$f(x) = \begin{cases} h(1 - a|x|), & |x| < \frac{1}{a}, \\ 0, & \text{autrement.} \end{cases}$$

Alors :

$$\begin{aligned} \hat{f}(\omega) &= \int_{-\frac{1}{a}}^{\frac{1}{a}} h(1 - a|x|) e^{-i\omega x} dx \\ &= 2h \operatorname{Re} \int_0^{\frac{1}{a}} (1 - ax) e^{-i\omega x} dx \\ &= 2h \operatorname{Re} \left[\frac{e^{-i\omega x}}{-i\omega} \Big|_0^{\frac{1}{a}} - a \int_0^{\frac{1}{a}} x e^{-i\omega x} dx \right] \end{aligned}$$

et par intégration par parties

$$\begin{aligned} &= 2h \operatorname{Re} \left[\frac{e^{-i\omega/a}}{-i\omega} + \frac{1}{i\omega} - a \left(\frac{x e^{-i\omega x}}{-i\omega} \Big|_0^{\frac{1}{a}} + \frac{1}{i\omega} \int_0^{\frac{1}{a}} e^{-i\omega x} dx \right) \right] \\ &= 2h \operatorname{Re} \left[\frac{e^{-i\omega/a}}{-i\omega} + \frac{1}{i\omega} - a \left(\frac{1}{a} \frac{e^{-i\omega/a}}{-i\omega} + \frac{1}{i\omega} \frac{e^{-i\omega x}}{-i\omega} \Big|_0^{\frac{1}{a}} \right) \right] \\ &= 2h \operatorname{Re} \left[\frac{1}{i\omega} - \frac{a}{\omega^2} (e^{-i\omega/a} - 1) \right] \\ &= \frac{2ha}{\omega^2} \left(1 - \cos \frac{\omega}{a} \right). \end{aligned}$$

Remarquons que, à nouveau, cette fonction est continue partout, même en $\omega = 0$, si on remplace l'expression ci-dessus par sa limite en $\omega \rightarrow 0$, c'est-à-dire par h/a .

transformée de Fourier
de la gaussienne

Exemple 4. Le prochain exemple joue un rôle important tant en statistique qu'en mécanique quantique. Il s'agit de la fonction gaussienne (ou encore, de la distribution normale) :

$$f(x) = e^{-\frac{x^2}{2}}.$$

Sa transformée de Fourier peut être obtenue par intégration dans le plan complexe ou à l'aide du truc suivant. Observons tout d'abord que nous connaissons une valeur de la fonction $\hat{f}(\omega)$. En effet, en $\omega = 0$, l'intégrale devient

$$\hat{f}(0) = \int_{-\infty}^{\infty} e^{-\frac{x^2}{2}} dx$$

et en remplaçant x par $\sqrt{2}u$:

$$\begin{aligned} &= \sqrt{2} \int_{-\infty}^{\infty} e^{-u^2} du \\ &= \sqrt{2} \Gamma\left(\frac{1}{2}\right) \\ &= \sqrt{2\pi}. \end{aligned}$$

Plutôt que calculer directement la transformée, nous allons obtenir une équation différentielle pour $\hat{f}(\omega)$. Calculons d'abord $\hat{f}'$. En admettant que l'on puisse intervertir l'ordre de la dérivation et de l'intégration :

$$\begin{aligned} \frac{d\hat{f}}{d\omega} &= \frac{d}{d\omega} \int_{-\infty}^{\infty} e^{-\frac{x^2}{2}} e^{-i\omega x} dx \\ &= \int_{-\infty}^{\infty} e^{-\frac{x^2}{2}} (-ix) e^{-i\omega x} dx \end{aligned}$$

qui peut être intégré par parties. Prenant $dv = -xe^{-\frac{x^2}{2}} dx$ et donc $v = e^{-\frac{x^2}{2}}$, nous obtenons :

$$= i \left(e^{-\frac{x^2}{2}} e^{-i\omega x} \Big|_{-\infty}^{\infty} + i\omega \int_{-\infty}^{\infty} e^{-\frac{x^2}{2}} e^{-i\omega x} dx \right).$$

Le terme à évaluer en $x = \pm\infty$ s'annule et l'intégrale résiduelle est précisément la définition de $\hat{f}(\omega)$. Donc

$$= -\omega \hat{f}(\omega).$$

Cette équation différentielle

$$\frac{d\hat{f}}{d\omega} = -\omega \hat{f}(\omega)$$

est particulièrement aisée à résoudre. La solution est simplement

$$\hat{f}(\omega) = \text{constante} \times e^{-\frac{\omega^2}{2}}.$$

La constante doit être choisie de façon à ce que $\hat{f}$ vérifie la valeur en $\omega = 0$ que nous connaissons. Avec cette valeur, nous fixons complètement $\hat{f}$:

$$\hat{f}(\omega) = \sqrt{2\pi} e^{-\frac{\omega^2}{2}}.$$

Remarquons, qu'à une constante près, la dépendance fonctionnelle de $\hat{f}$ en son argument est exactement celle de f ! (Est-ce la seule fonction avec cette propriété ? Voir les exercices.)

Ces quelques exemples suggèrent deux observations. Premièrement, la fonction $\hat{f}$ se comporte beaucoup mieux que f dans le sens suivant.

Lemme 25. Si f est intégrable au sens de Riemann sur tout intervalle $[a, b]$ et l'intégrale $\int_{-\infty}^{\infty} |f(x)| dx$ existe, alors $\hat{f}$ est (uniformément) continue.

En effet, le premier exemple était celui d'une fonction avec un saut en deux points, alors que les second et troisième exemples avec des dérivées non-continues. Pourtant, leurs transformées de Fourier sont continues partout.

Deuxièmement, la fonction $\hat{f}$ se comporte beaucoup moins bien que f dans le sens suivant.

Lemme 26. Même si $\int_{-\infty}^{\infty} |f(x)| dx < \infty$, l'intégrale $\int_{-\infty}^{\infty} |\hat{f}(w)| dw$ peut diverger.

Voici les preuves de ces deux affirmations.

Démonstration. Pour le premier lemme, on doit montrer que, pour tout $\epsilon > 0$, il existe $\delta > 0$, tel que si $\zeta, \eta \in \mathbb{R}$ et $|\zeta - \eta| < \delta$, alors $|\hat{f}(\zeta) - \hat{f}(\eta)| < \epsilon$. Puisque $\int_{-\infty}^{\infty} |f(x)| dx$ existe, $\hat{f}$ existe. De plus, pour tout $\epsilon > 0$, il existe $R(\epsilon)$ tel que

$$\int_{|x| \geq R(\epsilon)} |f(x)| dx \leq \frac{\epsilon}{4}.$$

Et puisque f est intégrable au sens de Riemann sur $[-R, R]$, elle est bornée et soit $\kappa(\epsilon) = \sup_{|x| \leq R(\epsilon)} |f(x)|$. Alors

$$\begin{aligned} |\hat{f}(\zeta) - \hat{f}(\eta)| &= \left| \int_{-\infty}^{\infty} f(x) (e^{-i\zeta x} - e^{-i\eta x}) dx \right| \\ &\leq \int_{-\infty}^{\infty} |f(x)| |e^{-i\zeta x} - e^{-i\eta x}| dx \\ &\leq 2 \int_{|x| \geq R(\epsilon)} |f(x)| dx + \int_{-R}^R |f(x)| |e^{-i\zeta x} - e^{-i\eta x}| dx \\ &\leq \frac{\epsilon}{2} + 2\kappa(\epsilon)R(\epsilon) \sup_{|x| \leq R(\epsilon)} |e^{-i\zeta x} - e^{-i\eta x}|. \end{aligned}$$

Il est possible de borner supérieurement le supremum comme suit. Changeant la variable x pour $y = (\zeta - \eta)x$, nous obtenons :

$$\begin{aligned} \sup_{|x| \leq R(\epsilon)} |e^{-i\zeta x} - e^{-i\eta x}| &= \sup_{|x| \leq R(\epsilon)} |e^{-i\zeta x}| |1 - e^{i(\zeta - \eta)x}| \\ &= \sup_{|y| \leq |\zeta - \eta|R(\epsilon)} |1 - e^{iy}| \\ &= \sup_{|y| \leq |\zeta - \eta|R(\epsilon)} |(1 - \cos y) - i \sin y| \\ &= \sup_{|y| \leq |\zeta - \eta|R(\epsilon)} \sqrt{2(1 - \cos y)} \\ &= 2 \sup_{|y| \leq |\zeta - \eta|R(\epsilon)} \left| \sin \frac{y}{2} \right| \end{aligned}$$

et puisque $|x| \geq |\sin x|$ pour tout $x \in \mathbb{R}$,

$$\leq |\zeta - \eta| R(\epsilon).$$

Ainsi

$$|\hat{f}(\zeta) - \hat{f}(\eta)| \leq \frac{\epsilon}{2} + 2\kappa(\epsilon) R^2(\epsilon) |\zeta - \eta|.$$

Pour terminer la preuve, il suffit donc de choisir δ de façon à ce que l'expression ci-dessus soit $< \epsilon$. Ceci est possible en posant

$$|\zeta - \eta| < \delta = \frac{\epsilon}{4\kappa(\epsilon) R^2(\epsilon) + 1},$$

le « 1 » ayant été ajouté au cas où κ serait nul. □

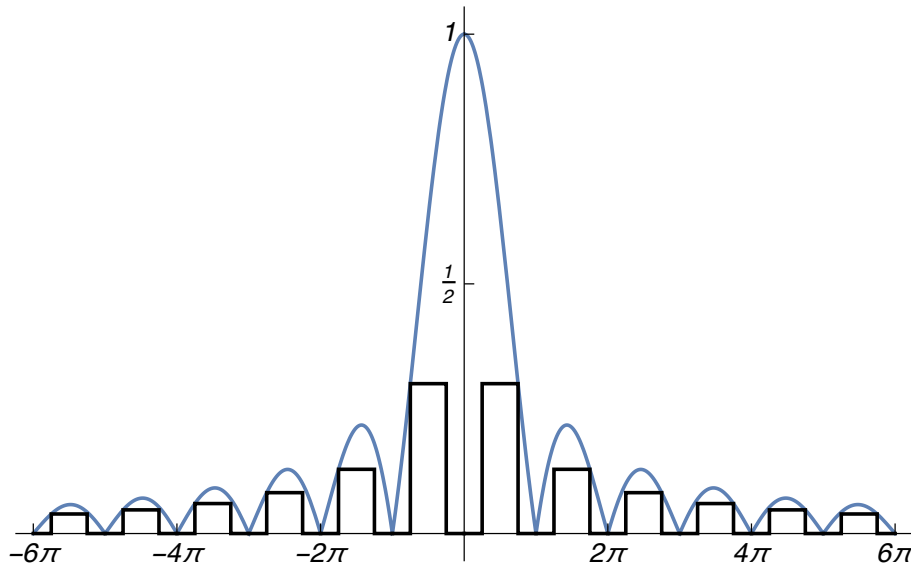


Figure 4.2 – La fonction $\sin \omega / \omega$ et l'approximation utilisée

Démonstration. Le second lemme est plus aisé à démontrer. La transformée de l'exemple 1 ci-dessus est du type désiré : notons que

$$\int_{-(N+1)\pi}^{(N+1)\pi} \left| \frac{2 \sin \omega}{\omega} \right| d\omega > \sum_{n=0}^N 4 \int_{n\pi + \frac{\pi}{4}}^{n\pi + \frac{3\pi}{4}} \left| \frac{\sin \omega}{\omega} \right| d\omega.$$

Cette majoration devrait être évidente : voir le graphique ci-dessus. Donc :

$$\begin{aligned}
 &> 4 \sum_{n=0}^N \int_{n\pi + \frac{\pi}{4}}^{n\pi + \frac{3\pi}{4}} \frac{1}{\sqrt{2}\omega} d\omega \\
 &> 2\sqrt{2} \sum_{n=0}^N \frac{\pi}{2} \frac{1}{\pi(n+1)} \\
 &\xrightarrow{N \rightarrow \infty} \infty.
 \end{aligned}$$

□

Exercices

1*. Obtenir les transformées de Fourier des fonctions suivantes.

(a)

$$f(x) = \begin{cases} |x|, & |x| \leq 1, \\ 0, & |x| > 1. \end{cases}$$

(b)

$$f(x) = \begin{cases} 1 - x^2, & |x| \leq 1, \\ 0, & |x| > 1. \end{cases}$$

(c)

$$f(x) = \begin{cases} 1, & 0 \leq x \leq 1 \\ -1, & -1 \leq x < 0 \\ 0, & \text{autrement.} \end{cases}$$

(d)

$$f(x) = \cos x^2.$$

Note : L'identité suivante pourrait être utile :

$$\int_0^\infty \cos ax^2 \cos 2bx \, dx = \frac{1}{2} \frac{\sqrt{\pi}}{\sqrt{2a}} \left(\cos \frac{b^2}{a} + \sin \frac{b^2}{a} \right).$$

(e)

$$f(x) = \begin{cases} \sin \omega_0 x, & |x| \leq \frac{N\pi}{\omega_0}, \\ 0, & |x| > \frac{N\pi}{\omega_0}. \end{cases}$$

ω_0 est ici une constante et N un entier non-nul. Quelles sont les valeurs de $\hat{f}$ en $\pm\omega_0$?

(f)

$$f(x) = \begin{cases} e^{-x}, & x > 0, \\ 0, & x \leq 0. \end{cases}$$

(g)

$$f(x) = \frac{1}{\cosh x}.$$

Note : A nouveau, voici une identité utile :

$$\int_0^\infty \frac{\cos ax}{\cosh bx} dx = \frac{\pi}{2b} \operatorname{sech} \frac{a\pi}{2b}.$$

(h) La fonction f_ν suivante pour $\nu + \frac{1}{2} > 0$:

$$f_\nu = \begin{cases} (1 - t^2)^{\nu - \frac{1}{2}}, & |t| \leq 1, \\ 0, & \text{autrement.} \end{cases}$$

2*. (a) Pour chacune des fonctions du problème précédent, vérifier que la fonction $\hat{f}$ est continue sur $\mathbb{R}$.(b) Quelles sont les fonctions $\hat{f}$ qui ne sont pas intégrables en valeur absolue, c'est-à-dire celles pour lesquelles $\int_{-\infty}^\infty |\hat{f}(\omega)| d\omega$ n'existe pas.Note : Il est utile de savoir que $\int_0^\infty \frac{\sin x}{x} dx$ et $\int_1^\infty \frac{\cos x}{x} dx$ convergent.**3.** (a) Montrer que, si $f : \mathbb{R} \rightarrow \mathbb{R}$ est une fonction paire, alors

$$\hat{f}(\omega) = 2 \int_0^\infty f(x) \cos \omega x dx$$

(b) et si elle est impaire, alors

$$\hat{f}(\omega) = -2i \int_0^\infty f(x) \sin \omega x dx.$$

(c) Quelle condition doit satisfaire f pour que $\hat{f}$ soit paire? Pour qu'elle soit impaire?**4*.** L'étude de l'oscillateur harmonique en mécanique quantique mène au problème de Sturm-Liouville

$$\left(-\frac{1}{2} \frac{d^2}{dx^2} + \frac{1}{2} x^2 \right) \psi = \lambda \psi.$$

Il est possible de montrer que si les conditions aux limites

$$\lim_{x \rightarrow \infty} \psi(x) = \lim_{x \rightarrow -\infty} \psi(x) = 0$$

sont choisies, les valeurs propres sont $\lambda_n = n + \frac{1}{2}$, $n \in \mathbb{N} \cup \{0\}$, et non-dégénérées.

(a) Montrer que $\psi_n(x) = a_n e^{-x^2/2} H_n(x)$ est une fonction propre de valeur propre $(n + \frac{1}{2})$. (H_n est le n -ième polynôme de Hermite. Voir les exercices de la section 3.3.)

(b) Trouver a_n tel que

$$\int_{-\infty}^{\infty} |\psi_n(x)|^2 dx = 1.$$

(c) Montrer que $\psi_{n+1}(x)$ et $(x - \frac{d}{dx})\psi_n$ sont proportionnels. Déterminer la constante de proportionnalité.

(d) Calculer les transformées de Fourier $\hat{\psi}_n(\omega)$ de $\psi_n(x)$.

Note : Trouver le lien entre $\hat{g}(\omega)$ et $\hat{f}(\omega)$ où g est la fonction $g(x) = xf(x)$.

(e) Vérifier que

$$\frac{1}{2\pi} \int_{-\infty}^{\infty} |\hat{\psi}_n(\omega)|^2 d\omega = \int_{-\infty}^{\infty} |\psi_n(x)|^2 dx.$$

5. Soit f une fonction dont la transformée de Fourier $\hat{f}$ existe. Exprimer la transformée de Fourier des fonctions suivantes en fonction de $\hat{f}$.

(a) La fonction g_d définie par $g_d(x) = f(dx)$ où d est un nombre réel non-nul.

(b) La fonction h_a définie par $h_a(x) = f(x - a)$ où a est un nombre réel.

6*. A la section 1.1, nous avons obtenu la série de Fourier de la fonction en créneau :

$$f(x) = \begin{cases} 1, & 0 \leq x \leq \pi, \\ -1, & -\pi < x < 0 \end{cases}$$

qui est

$$f(x) = \sum_{n=0}^{\infty} \frac{4}{(2n+1)\pi} \sin(2n+1)x,$$

l'égalité étant satisfaite en tout x qui n'est pas un multiple de π . Nous avons vu cependant que les graphiques 1.1 des sommes partielles

$$s_N(x) = \sum_{n=0}^N \frac{4}{(2n+1)\pi} \sin(2n+1)x$$

possèdent un maximum proche de zéro qui dépasse clairement la valeur $+1$. De plus ce maximum semble être à peu près constant pour N grand. Ce pic est connu sous le nom de *phénomène de Gibbs* et le but du présent exercice est d'évaluer sa hauteur. En faisant ce calcul, vous pratiquerez également le passage d'une somme à une intégrale que nous avons utilisé dans cette section pour introduire la transformation de Fourier.

phénomène de Gibbs
Willard Gibbs
(1839-1903)

(a) Montrer que le premier maximum $x_{\max}$ strictement positif de la somme partielle s_N est $\pi/(2(N+1))$.

(b) Montrer que

$$s_N \left(\frac{\pi}{2(N+1)} \right) \xrightarrow{N \rightarrow \infty} \frac{2}{\pi} \int_0^\pi \frac{\sin x}{x} dx.$$

Pour cela, introduire la variable

$$x = \frac{(2n+1)\pi}{2(N+1)}$$

et noter que les valeurs de x sont $\in [0, \pi]$. En se rappelant que la définition naïve de l'intégrale d'une fonction est la somme des aires de rectangles, écrivez la somme sous cette forme et prenez la limite. (Malgré que ce passage à la limite soit un peu « intuitif », il est possible de le rendre rigoureux.)

(c) Obtenir une valeur numérique pour la hauteur du pic.

4.2 Quelques théorèmes²

Les résultats rigoureux sur la transformation de Fourier ne sont guère plus difficiles à démontrer que ceux sur les séries de Fourier. En fait, beaucoup sont très semblables. La seule différence réside dans la nécessité d'intervertir deux intégrations $\int_{-\infty}^{\infty} d\lambda \int_{-\infty}^{\infty} dx \cdots \rightarrow \int_{-\infty}^{\infty} dx \int_{-\infty}^{\infty} d\lambda \dots$, ce qui nécessite des conditions particulières sur le « ... ». Nous énoncerons ces conditions sans preuve.

Pour le restant de cette section, nous nous restreindrons à des fonctions particulières.

Définition 18. $L^1 \cap C$ est l'ensemble des fonctions $\mathbb{R} \rightarrow \mathbb{C}$ continues et telles que $\int_{-\infty}^{\infty} |f(x)| dx$ existe (et soit $< \infty$).

Le théorème de Fejér pour la transformation de Fourier s'énonce comme suit.

Théorème 27 (Fejér). (i) Si $f \in L^1 \cap C$, alors

$$\frac{1}{2\pi} \int_{-R}^R \left(1 - \frac{|\lambda|}{R}\right) \hat{f}(\lambda) e^{i\lambda x} d\lambda \xrightarrow{R \rightarrow \infty} f(x), \quad \forall x \in \mathbb{R}.$$

(ii) Si f est uniformément continue et bornée, alors la convergence est uniforme.

2. Le traitement de cette section ainsi que les exercices 2 et 6 suivent Körner.

La preuve de ce théorème est similaire à celle pour les séries de Fourier. Rappelons que les moyennes σ_n des sommes partielles définies à la section 2.3 peuvent être réécrites sous la forme

$$\sigma_n(f, t) = \sum_{r=-n}^n \frac{n+1-|r|}{n+1} c_r e^{irt} = \sum_{r=-n}^n \left(1 - \frac{|r|}{n+1}\right) c_r e^{irt}.$$

Le théorème de Fejér 5 énonce que, si f est continue sur le tore, alors $\sigma_n(f, \cdot) \rightarrow f$ uniformément quand $n \rightarrow \infty$. Pour les transformations de Fourier, la somme dans σ_n est remplacée par l'intégrale de $-R$ à R et le facteur $(1 - \frac{|r|}{n+1})$ est remplacé par $(1 - \frac{|\lambda|}{R})$. L'énoncé ci-dessus semble donc fort raisonnable.

Mais quel est l'objet jouant ici le rôle du noyau de Fejér des séries de Fourier. Il s'agit de faire un calcul semblable à celui fait pour les séries de Fourier pour l'identifier :

$$\begin{aligned} \frac{1}{2\pi} \int_{-R}^R \left(1 - \frac{|\lambda|}{R}\right) \hat{f}(\lambda) e^{i\lambda x} d\lambda &= \\ &= \frac{1}{2\pi} \int_{-R}^R \left(1 - \frac{|\lambda|}{R}\right) \left(\int_{-\infty}^{\infty} f(t) e^{-i\lambda t} dt \right) e^{i\lambda x} d\lambda \\ &= \frac{1}{2\pi} \int_{-R}^R \int_{-\infty}^{\infty} \left(1 - \frac{|\lambda|}{R}\right) f(t) e^{i\lambda(x-t)} dt d\lambda \\ &\stackrel{?}{=} \frac{1}{2\pi} \int_{-\infty}^{\infty} dt f(t) \left(\int_{-R}^R \left(1 - \frac{|\lambda|}{R}\right) e^{i\lambda(x-t)} d\lambda \right) \\ &= \int_{-\infty}^{\infty} dt f(t) \tilde{K}_R(x-t) \\ &= \int_{-\infty}^{\infty} dy \tilde{K}_R(y) f(x-y) \end{aligned}$$

où, comme précédemment, la preuve reposera sur

$$\tilde{K}_R(y) = \frac{1}{2\pi} \int_{-R}^R \left(1 - \frac{|\lambda|}{R}\right) e^{i\lambda y} d\lambda.$$

Une des étapes ci-dessus a été marquée par « $\stackrel{?}{=}$ ». C'est précisément l'étape où deux intégrales ont été interverties et qui devra donc être justifiée. Pour l'instant, notons que

$$\begin{aligned} \tilde{K}_R(y) &= \frac{1}{2\pi} \int_{-R}^R \left(1 - \frac{|\lambda|}{R}\right) e^{i\lambda y} d\lambda \\ &= \frac{2}{2\pi} \operatorname{Re} \int_0^R \left(1 - \frac{\lambda}{R}\right) e^{i\lambda y} d\lambda \\ &= \frac{1}{\pi} \int_0^R \left(1 - \frac{\lambda}{R}\right) \cos \lambda y d\lambda \end{aligned}$$

et intégrant par parties pour le cas $y \neq 0$:

$$\begin{aligned}
 &= \frac{1}{\pi} \left(\left(1 - \frac{\lambda}{R} \right) \frac{\sin \lambda y}{y} \Big|_0^R + \frac{1}{yR} \int_0^R \sin \lambda y \, d\lambda \right) \\
 &= -\frac{1}{\pi} \frac{1}{y^2 R} \cos \lambda y \Big|_0^R \\
 &= \frac{1 - \cos yR}{y^2 \pi R} \\
 &= \frac{2R \sin^2 \frac{yR}{2}}{4\pi \frac{y^2 R^2}{4}} \\
 &= \frac{R}{2\pi} \left(\frac{\sin \frac{yR}{2}}{\frac{yR}{2}} \right)^2.
 \end{aligned}$$

Pour le cas $y = 0$

$$\begin{aligned}
 \tilde{K}_R(0) &= \frac{1}{2\pi} \int_{-R}^R \left(1 - \frac{|\lambda|}{R} \right) d\lambda \\
 &= \frac{2}{2\pi} \left(\lambda \Big|_0^R - \frac{\lambda^2}{2R} \Big|_0^R \right) \\
 &= \frac{1}{\pi} \left(R - \frac{R^2}{2R} \right) \\
 &= \frac{R}{2\pi}
 \end{aligned}$$

qui est la limite du cas $y \neq 0$ quand $y \rightarrow 0$. (Il est possible que le calcul de $\tilde{K}_R$ vous rappelle l'exemple 3 de la section précédente. Exercice : quel est le lien entre les deux calculs ?)

Revenons maintenant sur cette étape « ? » déclarée dangereuse. Et pourquoi est-il si dangereux d'intervertir deux intégrales ? La meilleure réponse est donnée par l'exemple simple suivant. Faisons l'intégrale double suivante :

$$\begin{aligned}
 \int_1^\infty dx \left(\int_1^\infty \frac{x^2 - y^2}{(x^2 + y^2)^2} dy \right) &= \int_1^\infty \frac{y}{x^2 + y^2} \Big|_{y=1}^\infty dx, && \text{voir plus bas} \\
 &= -\int_1^\infty \frac{1}{1 + x^2} dx \\
 &= -\arctan x \Big|_1^\infty \\
 &= -\left(\frac{\pi}{2} - \frac{\pi}{4} \right) \\
 &= -\frac{\pi}{4}
 \end{aligned}$$

où nous avons utilisé le fait que

$$\begin{aligned}\frac{\partial}{\partial y} \frac{y}{x^2 + y^2} &= \frac{1}{x^2 + y^2} - \frac{2y^2}{(x^2 + y^2)^2} \\ &= \frac{x^2 + y^2 - 2y^2}{(x^2 + y^2)^2} \\ &= \frac{x^2 - y^2}{(x^2 + y^2)^2}.\end{aligned}$$

Refaisons maintenant la même double intégrale, mais dans l'autre ordre.

$$\int_1^\infty dy \left(\int_1^\infty \frac{x^2 - y^2}{(x^2 + y^2)^2} dx \right) = - \int_1^\infty dy \left(\int_1^\infty \frac{y^2 - x^2}{(x^2 + y^2)^2} dx \right) = \frac{\pi}{4}$$

puisque la seconde forme de la double intégrale est identique à celle ci-dessus où le rôle de x et de y a été échangé. Vous pouvez chercher l'erreur dans ce raisonnement. Le fait est qu'il n'y en a pas. Ce que cet exemple nous rappelle est qu'intervertir deux processus limites (dans ce cas-ci l'intégration sur un intervalle infini) ne peut pas toujours être fait. Alors quand peut-on intervertir les intégrales ? ! Le cas le plus simple est quand les domaines des deux intégrales sont des intervalles finis. Alors :

Théorème 28 (Fubini). *Si $f : [a, b] \times [c, d] \rightarrow \mathbb{C}$ est continue, alors :*

$$\int_a^b dx \int_c^d dy f(x, y) = \int_c^d dy \int_a^b dx f(x, y).$$

Ceci est rassurant mais nous avons besoin d'un résultat similaire pour les intégrales sur des intervalles infinis.

Théorème 29. *Soit $f : \mathbb{R}^2 \rightarrow \mathbb{C}$ une fonction continue telle que*

(A1) $\int_{-\infty}^\infty |f(x, y)| dy$ soit une fonction continue de x ,

(A2) $\int_{-\infty}^\infty |f(x, y)| dx$ soit une fonction continue de y

(B) et telle que $\int_{-\infty}^\infty \int_{-\infty}^\infty |f(x, y)| dx dy$ converge.

Alors les intégrales

$$\int_{-\infty}^\infty dx \int_{-\infty}^\infty dy f(x, y) \quad \text{et} \quad \int_{-\infty}^\infty dy \int_{-\infty}^\infty dx f(x, y)$$

existent et sont égales.

Nous accepterons ce résultat sans preuve. Dans notre cas

$$g(t, \lambda) = \begin{cases} \left(1 - \frac{|\lambda|}{R}\right) f(t) e^{i\lambda(x-t)}, & |\lambda| < R, \\ 0, & \text{autrement.} \end{cases}$$

Remarquons que

$$\int_{-\infty}^{\infty} |g(t, \lambda)| dt = \begin{cases} \left(1 - \frac{|\lambda|}{R}\right) \int_{-\infty}^{\infty} |f(t)| dt, & |\lambda| < R, \\ 0, & \text{autrement,} \end{cases}$$

qui est une fonction continue de λ et

$$\int_{-\infty}^{\infty} |g(t, \lambda)| d\lambda = |f(t)| \int_{-R}^R \left(1 - \frac{|\lambda|}{R}\right) d\lambda = R|f(t)|$$

qui est une fonction continue de t . Enfin

$$\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} |g(t, \lambda)| d\lambda dt = R \int_{-\infty}^{\infty} |f(t)| dt < \infty$$

si $f \in L^1 \cap C$. Donc, l'ordre d'intégration peut être changé.

Nous pouvons maintenant revenir à la preuve du théorème de Fejér pour la transformation de Fourier.

Lemme 30. (i) $\tilde{K}_R(s) \geq 0$, $\forall s \in \mathbb{R}$.
(ii) $\tilde{K}_R(s) \rightarrow 0$ uniformément sur $\mathbb{R} \setminus [-\delta, \delta]$ et

$$\int_{|s| \geq \delta} \tilde{K}_R(s) ds \xrightarrow{R \rightarrow \infty} 0, \quad \forall \delta > 0.$$

(iii) $\int_{-\infty}^{\infty} \tilde{K}_R(s) ds = 1$.

Ce lemme peut être évidemment mis en parallèle avec le lemme 4. Pour la partie (iii), nous ne montrerons que $\int_{-\infty}^{\infty} \tilde{K}_R(s) ds < \infty$.

Démonstration. (i) Évident.

(ii)

$$\left| \frac{R}{2\pi} \left(\frac{\sin \frac{sR}{2}}{\frac{sR}{2}} \right)^2 \right| \leq \frac{R}{2\pi} \frac{4}{s^2 R^2} \leq \frac{2}{\pi R \delta^2}$$

pour tout $s \in \mathbb{R} \setminus [-\delta, \delta]$ et

$$\begin{aligned} \int_{|s| \geq \delta} \tilde{K}_R(s) ds &< \int_{|s| \geq \delta} \frac{2}{\pi R s^2} ds = -\frac{4}{\pi R} \frac{1}{s} \Big|_{\delta}^{\infty} \\ &= \frac{4}{\pi R \delta} \xrightarrow{R \rightarrow \infty} 0 \end{aligned}$$

pour tout $\delta > 0$.

(iii) Enfin :

$$\begin{aligned}
 \int_{-\infty}^{\infty} K_R(s) ds &= \frac{R}{2\pi} \int_{-\infty}^{\infty} \left(\frac{\sin \frac{sR}{2}}{\frac{sR}{2}} \right)^2 ds \\
 &= \frac{1}{2\pi} \int_{-\infty}^{\infty} \left(\frac{\sin \frac{u}{2}}{\frac{u}{2}} \right)^2 du, \quad \text{en posant } u = sR \\
 &= \int_{-\infty}^{\infty} \tilde{K}_1(s) ds \\
 &< \infty
 \end{aligned}$$

car, sur $[-\delta, \delta]$, la fonction $\tilde{K}_1$ est continue et donc bornée et alors $\int_{-\delta}^{\delta} \tilde{K}_1 ds < \infty$ et $\int_{\mathbb{R} \setminus [-\delta, \delta]} \tilde{K}_1 ds < 4/\pi\delta$ par (ii). $\square$

La preuve du théorème de Fejér est alors aisée.

Preuve du théorème 27. Comme pour le théorème de Fejér pour les séries de Fourier, la preuve se résume à calculer

$$\begin{aligned}
 &\left| \frac{1}{2\pi} \int_{-R}^R \left(1 - \frac{|\lambda|}{R} \right) \hat{f}(\lambda) e^{i\lambda t} d\lambda - f(t) \right| \\
 &= \left| \int_{-\infty}^{\infty} \tilde{K}_R(x) f(t-x) dx - f(t) \right| \\
 &= \left| \int_{-\infty}^{\infty} \tilde{K}_R(x) (f(t-x) - f(t)) dx \right|, \quad \text{par (iii) ci-dessus} \\
 &\leq \left| \int_{|x| \geq \delta} \tilde{K}_R(x) (f(t-x) - f(t)) dx \right| \\
 &\quad + \left| \int_{|x| \leq \delta} \tilde{K}_R(x) (f(t-x) - f(t)) dx \right| \\
 &\leq \left| \int_{|x| \geq \delta} \tilde{K}_R(x) f(t-x) dx \right| + \left| \int_{|x| \geq \delta} \tilde{K}_R(x) f(t) dx \right| \\
 &\quad + \left| \int_{|x| \leq \delta} \tilde{K}_R(x) (f(t-x) - f(t)) dx \right| \\
 &\leq \left(\sup_{|x| \geq \delta} \tilde{K}_R(x) \right) \int_{-\infty}^{\infty} |f(t-x)| dx + |f(t)| \int_{|x| \geq \delta} \tilde{K}_R(x) dx \\
 &\quad + \left(\sup_{|x| \leq \delta} |f(t-x) - f(t)| \right) \int_{|x| \leq \delta} \tilde{K}_R(x) dx.
 \end{aligned}$$

Supposons que $\epsilon > 0$ soit donné. Puisque f est continue, on peut trouver $\delta(\epsilon) > 0$ tel que

$$\text{si } |x| < \delta(\epsilon) \quad \text{alors} \quad |f(t-x) - f(t)| < \frac{\epsilon}{3}.$$

Donné ce $\delta(\epsilon)$, le lemme précédent assure qu'il existe $R_0(\epsilon)$ tel que

$$\sup_{|x| \geq \delta} \tilde{K}_R(x) < \frac{\epsilon}{3} \frac{1}{1 + \int_{-\infty}^{\infty} |f(t-x)| dx}, \quad \text{si } R > R_0(\epsilon)$$

et

$$\int_{|x| > \delta} \tilde{K}_R(x) dx < \frac{\epsilon}{3} \frac{1}{1 + |f(t)|}.$$

(Les nombres « 1 » ont été ajoutés aux dénominateurs pour assurer que ceux-ci soient non nuls.) Ainsi

$$\left| \frac{1}{2\pi} \int_{-R}^R \left(1 - \frac{|\lambda|}{R}\right) \hat{f}(\lambda) e^{i\lambda t} d\lambda - f(t) \right| < \frac{\epsilon}{3} + \frac{\epsilon}{3} + \frac{\epsilon}{3} = \epsilon.$$

□

Nous sommes maintenant prêts à établir la relation $(\mathcal{F})$ introduite à la section précédente et qui assure que f et $\hat{f}$ contiennent la « même information ». Plus précisément :

Théorème 31 (Théorème d'inversion). *Si $f \in L^1 \cap C$ et $\hat{f} \in L^1 \cap C$, alors* théorème d'inversion

$$f(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\lambda) e^{i\lambda t} d\lambda,$$

l'intégrale étant convergente absolument uniformément.

Démonstration. Rappelons tout d'abord que, par définition

$$\int_{-\infty}^{\infty} \hat{f}(\lambda) e^{i\lambda t} d\lambda = \lim_{R, S \rightarrow \infty} \int_{-R}^S \hat{f}(\lambda) e^{i\lambda t} d\lambda.$$

La convergence uniforme de l'intégrale signifie simplement que le choix de R_0 et S_0 tel que

$$\text{si } R > R_0 \text{ et } S > S_0, \quad \text{alors} \quad \left| \int_{-\infty}^{\infty} \hat{f}(\lambda) e^{i\lambda t} d\lambda - \int_{-R}^S \hat{f}(\lambda) e^{i\lambda t} d\lambda \right| < \epsilon,$$

ne dépend pas de t . Mais ceci découle rapidement du fait que $\hat{f} \in L^1 \cap C$:

$$\begin{aligned} \left| \int_{-\infty}^{\infty} \hat{f}(\lambda) e^{i\lambda t} d\lambda - \int_{-R}^S \hat{f}(\lambda) e^{i\lambda t} d\lambda \right| &= \left| \int_{\lambda > S \text{ ou } \lambda < -R} \hat{f}(\lambda) e^{i\lambda t} d\lambda \right| \\ &\leq \int_{\lambda > S \text{ ou } \lambda < -R} |\hat{f}(\lambda) e^{i\lambda t}| d\lambda \\ &= \int_{\lambda > S \text{ ou } \lambda < -R} |\hat{f}(\lambda)| d\lambda \end{aligned}$$

qui est une quantité (indépendante de t) tendant vers 0 lorsque S et R tendent vers ∞ .

Reste à montrer que l'intégrale converge bien vers $f(t)$. Le théorème 27 a montré que

$$\frac{1}{2\pi} \int_{-\infty}^{\infty} \Delta_R(\lambda) \hat{f}(\lambda) e^{i\lambda t} d\lambda \xrightarrow{R \rightarrow \infty} f(t)$$

où

$$\Delta_R(\lambda) = \begin{cases} 1 - \frac{|\lambda|}{R}, & |\lambda| \leq R, \\ 0, & \text{autrement.} \end{cases}$$

Ainsi il suffit donc de montrer que

$$\frac{1}{2\pi} \int_{-\infty}^{\infty} \Delta_R(\lambda) \hat{f}(\lambda) e^{i\lambda t} d\lambda \xrightarrow{R \rightarrow \infty} \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\lambda) e^{i\lambda t} d\lambda.$$

Puisque $\hat{f} \in L^1 \cap C$, pour tout $\epsilon > 0$, il existe $S > 0$ tel $\int_{|\lambda| > S} |\hat{f}(\lambda)| d\lambda < \epsilon$. Alors, pour $R > \frac{S}{\epsilon}$:

$$\begin{aligned} \left| \int_{-\infty}^{\infty} \hat{f}(\lambda) e^{i\lambda t} d\lambda - \int_{-\infty}^{\infty} \Delta_R(\lambda) \hat{f}(\lambda) e^{i\lambda t} d\lambda \right| \\ = \left| \int_{-\infty}^{\infty} (1 - \Delta_R(\lambda)) \hat{f}(\lambda) e^{i\lambda t} d\lambda \right| \end{aligned}$$

et puisque $(1 - \Delta_R)$ est toujours positif :

$$\begin{aligned} &\leq \int_{-\infty}^{\infty} (1 - \Delta_R(\lambda)) |\hat{f}(\lambda)| d\lambda \\ &= \int_{|\lambda| < S} (1 - \Delta_R(\lambda)) |\hat{f}(\lambda)| d\lambda + \int_{|\lambda| \geq S} (1 - \Delta_R(\lambda)) |\hat{f}(\lambda)| d\lambda \\ &\leq \int_{|\lambda| < S} \frac{S}{R} |\hat{f}(\lambda)| d\lambda + \int_{|\lambda| \geq S} |\hat{f}(\lambda)| d\lambda \\ &\leq \epsilon \left(\int_{-\infty}^{\infty} |\hat{f}(\lambda)| d\lambda + 1 \right) \end{aligned}$$

ce qui termine la preuve. (La parenthèse ci-dessus peut être éliminée par un nouveau choix de R et S correspondant à $\epsilon' = \epsilon / (\int |\hat{f}| d\lambda + 1)$.) $\square$

Le premier exercice ci-dessous établira une autre conséquence du théorème 29 sur l'ordre d'intégration. Il s'agit de l'analogue de l'identité de Parseval (théorème 23) pour les séries de Fourier.

Corollaire 32 (Identité de Parseval). *Si $f \in L^1 \cap C$ et $\hat{f} \in L^1 \cap C$, alors*

$$\int_{-\infty}^{\infty} |f(t)|^2 dt = \frac{1}{2\pi} \int_{-\infty}^{\infty} |\hat{f}(\lambda)|^2 d\lambda.$$

identité
de Parseval

Cette identité permet de calculer quelques intégrales impropres comme l'indique l'exemple qui suit.

Exemple. L'exemple 2 de la section précédente a établi que pour

$$f(x) = e^{-a|x|}, \quad a > 0,$$

la fonction $\hat{f}$ est

$$\hat{f}(\lambda) = \frac{2a}{a^2 + \lambda^2}.$$

Les deux fonctions f et $\hat{f}$ sont clairement continues. La fonction f est aussi intégrable et donc $f \in L^1 \cap C$. Quant à $\hat{f}$, notons que

$$|\hat{f}(\lambda)| = \frac{1}{a} \frac{2}{1 + \frac{\lambda^2}{a^2}} \leq \begin{cases} \frac{2}{a}, & |\lambda| < a \\ \frac{2a}{\lambda^2}, & |\lambda| \geq a. \end{cases}$$

et donc que $\hat{f} \in L^1 \cap C$ également.

L'identité de Parseval affirme dans ce cas que

$$\int_{-\infty}^{\infty} e^{-2a|x|} dx = \frac{1}{2\pi} \int_{-\infty}^{\infty} \frac{4a^2}{(a^2 + \lambda^2)^2} d\lambda.$$

Pour cet exemple, il est aisé de vérifier cette relation par intégration explicite. Mais dans bien des cas, au moins une des deux intégrales est très difficile à faire par les méthodes usuelles.

Vérifions l'identité dans le cas présent.

$$\begin{aligned} \int_{-\infty}^{\infty} e^{-2a|x|} dx &= 2 \int_0^{\infty} e^{-2ax} dx \\ &= -2 \left. \frac{e^{-2ax}}{2a} \right|_0^{\infty} \\ &= \frac{1}{a}. \end{aligned}$$

Pour le membre de droite :

$$\begin{aligned} \int_{-\infty}^{\infty} \frac{4a^2}{(a^2 + \lambda^2)^2} d\lambda &= \int_{-\infty}^{\infty} \frac{4a^2}{a^4 (1 + \frac{\lambda^2}{a^2})^2} d\lambda \\ &= \frac{4}{a} \int_{-\infty}^{\infty} \frac{1}{(1 + \omega^2)^2} d\omega \end{aligned}$$

et la primitive est

$$\begin{aligned} &= \frac{2}{a} \left(\frac{\omega}{1 + \omega^2} + \arctan \omega \right) \Big|_{-\infty}^{\infty} \\ &= \frac{2}{a} (0 + \pi) = \frac{2\pi}{a} \end{aligned}$$

et l'égalité est donc vérifiée.

Exercices

7*. Quelles hypothèses du théorème 29 ne sont pas vérifiées par la fonction

$$f(x, y) = \begin{cases} \frac{x^2 - y^2}{(x^2 + y^2)^2}, & x, y \geq 1, \\ 0, & \text{autrement?} \end{cases}$$

8. Le but de cet exercice est d'établir l'identité de Parseval.

(a) Si $f, g, \hat{f}, \hat{g} \in L^1 \cap C$, alors utiliser le théorème 29 pour la fonction $F(x, y) = f(x)g(y)e^{-ixy}$ pour montrer que

$$\int_{-\infty}^{\infty} f(t)\hat{g}(t)dt = \int_{-\infty}^{\infty} \hat{f}(t)g(t)dt.$$

(b) Montrer que, si $f, g, \hat{f}, \hat{g} \in L^1 \cap C$, alors

$$\int_{-\infty}^{\infty} f(t)g(t)^*dt = \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\lambda)\hat{g}(\lambda)^*d\lambda.$$

identité de Parseval

(c) En déduire

$$\int_{-\infty}^{\infty} |f(t)|^2 dt = \frac{1}{2\pi} \int_{-\infty}^{\infty} |\hat{f}(\lambda)|^2 d\lambda$$

pour $f, \hat{f} \in L^1 \cap C$. Ceci est l'analogue de l'identité de Parseval pour les séries de Fourier.

9*. (a) Parmi les exemples et les fonctions de l'exercice **1.** de la section précédente, quels sont ceux dont la paire f et $\hat{f}$ appartient à $L^1 \cap C$.

(b) Appliquer l'identité de Parseval à ces paires. (Parfois vous pourrez faire explicitement les deux intégrales, mais pas dans tous les cas.)

10. (a) Soit $f \in L^1 \cap C$. Montrer que

$$f_n(x) = \sum_{j=0}^n f(x+j)$$

est également dans $L^1 \cap C$ si $n \in \mathbb{N}$.

(b) Montrer que

$$\hat{f}_n(\omega) = e^{i\omega n/2} \frac{\sin\left(\frac{n+1}{2}\omega\right)}{\sin\frac{\omega}{2}} \hat{f}(\omega).$$

11. Cet exercice poursuit le précédent.

(a) Supposons maintenant que $f \in L^1 \cap C$ décroisse à l'infini ($|x| \rightarrow \infty$) au moins comme $\frac{1}{x^2}$. (En d'autres mots, il existe $M, R > 0$ tel que, si $|x| > R$, alors $|f(x)| \leq \frac{M}{x^2}$.) Montrer que la fonction

$$\sum_{j \in \mathbb{Z}} f(x+j)$$

existe et est uniformément continue. Nous appellerons f_∞ cette fonction.

(b) Peut-on calculer $\widehat{f_\infty}$, par exemple, comme la limite du $\widehat{f_n}$ obtenu à l'exercice précédent ? Pourquoi ?

(c) La fonction f_∞ étant périodique de période 1, il est possible de calculer son développement en série de Fourier. En évaluant cette série en $x = 0$, montrer que

$$\sum_{n \in \mathbb{Z}} f(n) = \sum_{m \in \mathbb{Z}} \widehat{f}(2\pi m).$$

formule de somme
de Poisson
Siméon Denis Poisson
(1781-1840)

Cette identité est la formule de somme de Poisson.

(d) Montrer que $f(y) = e^{-\pi xy^2}$ satisfait aux exigences de l'exercice (a) (ou, plus précisément, pour quelles valeurs du paramètre x la fonction $f(y)$ satisfait aux condition de (a)).

(e) En déduire l'identité :

$$\sum_{n \in \mathbb{Z}} e^{-\pi n^2 x} = x^{-\frac{1}{2}} \sum_{m \in \mathbb{Z}} e^{-\pi m^2 / x}.$$

Note : La fonction elliptique $\theta(z) = \sum_{n \in \mathbb{Z}} e^{i\pi n^2 z}$ joue un rôle important en analyse complexe, en théorie des nombres et même en physique théorique. L'identité ci-dessus peut être réécrite sous la forme :

$$\theta(z) = e^{\frac{i\pi}{4} z - \frac{1}{2} z} \theta\left(-\frac{1}{z}\right).$$

Elle a ceci de remarquable qu'elle relie, par exemple, le comportement de la même fonction autour de 0 et de l'infini.

12. L'inégalité de Heisenberg. Cet exercice a pour but de démontrer l'inégalité de Heisenberg qui est la formulation mathématique du principe d'incertitude de Heisenberg. Cette relation est probablement l'inégalité la plus spectaculaire du XXIème siècle car elle incarne la différence profonde entre les descriptions classique et quantique de la nature. La discussion de son contenu physique occupe plusieurs semaines des cours de mécanique quantique. Nous nous en tiendrons ici qu'à la preuve mathématique. (Körner)

inégalité
de Heisenberg
Werner Heisenberg
(1901-1976)

(a) Soit $f \in L^1 \cap C$ telle que $f(x) \geq 0$ pour tout $x \in \mathbb{R}$ et $\int_{-\infty}^{\infty} f(x) dx = 1$. Montrer que si $f(x) = 0$ pour $|x| \geq \delta > 0$, alors $|\widehat{f}(\lambda)| \geq \frac{1}{2}$ pour $|\lambda| \leq \frac{1}{100\delta}$. (La constante

$\frac{1}{100}$ n'est sûrement pas la meilleure borne supérieure.) Ce premier exercice montre que, si f est concentrée dans un intervalle $[-\delta, \delta]$, sa transformée $\hat{f}$ doit être étalée (c'est-à-dire non négligeable) sur un énorme intervalle $[-\frac{1}{100\delta}, \frac{1}{100\delta}]$. C'est ce « compromis » entre les plages où f et $\hat{f}$ prennent des valeurs importantes que l'inégalité de Heisenberg régit.

(b) En supposant que f se comporte bien, justifier chacune des étapes suivantes. (Vous devrez utiliser le résultat prouvé à l'exercice 8 (a) ci-dessus.) Préciser les conditions sur f .

$$\begin{aligned}
 & \frac{1}{2\pi} \int_{-\infty}^{\infty} x^2 |f(x)|^2 dx \int_{-\infty}^{\infty} \lambda^2 |\hat{f}(\lambda)|^2 d\lambda \\
 &= \frac{1}{2\pi} \int_{-\infty}^{\infty} |xf(x)|^2 dx \int_{-\infty}^{\infty} |\lambda \hat{f}(\lambda)|^2 d\lambda \\
 &= \int_{-\infty}^{\infty} |xf(x)|^2 dx \int_{-\infty}^{\infty} |f'(x)|^2 dx \\
 &\geq \left(\int_{-\infty}^{\infty} |xf'(x)f(x)| dx \right)^2 \\
 &\geq \left(\int_{-\infty}^{\infty} \frac{x}{2} (f'(x)f^*(x) + f(x)f^{*'}(x)) dx \right)^2 \\
 &= \frac{1}{4} \left(\int_{-\infty}^{\infty} x \left(\frac{d}{dx} |f(x)|^2 \right) dx \right)^2 \\
 &= \frac{1}{4} \left(\int_{-\infty}^{\infty} |f(x)|^2 dx \right)^2 \\
 &= \frac{1}{8\pi} \int_{-\infty}^{\infty} |f(x)|^2 dx \int_{-\infty}^{\infty} |\hat{f}(\lambda)|^2 d\lambda.
 \end{aligned}$$

(c) Expliquer en mots pourquoi l'inégalité démontrée ci-dessus :

$$\frac{\int_{-\infty}^{\infty} x^2 |f(x)|^2 dx}{\int_{-\infty}^{\infty} |f(x)|^2 dx} \cdot \frac{\int_{-\infty}^{\infty} \lambda^2 |\hat{f}(\lambda)|^2 d\lambda}{\int_{-\infty}^{\infty} |\hat{f}(\lambda)|^2 d\lambda} \geq \frac{1}{4}$$

affirme que, si f est concentrée en l'origine, $\hat{f}$ ne peut pas l'être et vice versa.

(d) Montrer que l'inégalité « $\geq$ » ci-dessus est une égalité « $=$ » si

$$f'(x) = kxf(x)$$

pour tout x et une certaine constante k . En conclure que les seules fonctions réelles (lisses) vérifiant l'égalité « $=$ » sont celles de la forme

$$f(x) = Ae^{-|k|x^2/2}, \quad A, k \in \mathbb{R}.$$

Vérifier qu'en effet, l'égalité est satisfaite pour ces fonctions.

4.3 La convolution

Avant d'appliquer la transformation de Fourier à la solution d'équations aux dérivées partielles, nous devons encore développer deux outils :

- la dérivation de $\hat{f}$;
- la convolution.

La première est une simple propriété technique dont nous aurons besoin. La seconde est cependant un outil fort étrange lorsqu'on le voit pour la première fois.

Commençons donc par citer les théorèmes qui nous permettent de dériver sous le signe de l'intégrale. Le résultat le plus simple décrit le cas d'une intégrale sur un intervalle fermé borné.

Théorème 33. Soit $g : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{C}$ une fonction continue : $g = g(x, t)$. Si la dérivée partielle $\frac{\partial g}{\partial t}$ existe et est continue, alors $\int_a^b g(x, t) dx$ sur l'intervalle borné $[a, b]$ est différentiable et

$$\frac{d}{dt} \int_a^b g(x, t) dx = \int_a^b \frac{\partial g}{\partial t}(x, t) dx.$$

Ce résultat est simple et direct. Mais le cas dont nous aurons besoin nécessite plus de conditions sur la fonction g . En effet, pour obtenir $\hat{f}'$ de façon similaire, il faut passer le signe de dérivée sous le signe d'intégrale qui définit la transformée de Fourier, c'est-à-dire d'une intégrale de $-\infty$ à $+\infty$. Dans ce cas, le théorème s'énonce comme suit.

Théorème 34. Soit $g : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{C}$ une fonction continue telle que la dérivée $\frac{\partial g}{\partial t}$ existe et soit continue. Si, de plus, $\int_{-\infty}^{\infty} |g(x, t)| dx$ et $\int_{-\infty}^{\infty} |\frac{\partial g}{\partial t}| dx$ existent en tout t et $\int_{|x|>R} |\frac{\partial g}{\partial t}| dx \rightarrow 0$ quand $R \rightarrow \infty$ uniformément en t sur tout intervalle fermé $[a, b]$, alors $\int_{-\infty}^{\infty} g(x, t) dx$ est différentiable avec

$$\frac{d}{dt} \int_{-\infty}^{\infty} g(x, t) dx = \int_{-\infty}^{\infty} \frac{\partial g}{\partial t}(x, t) dx.$$

Nous tiendrons pour acquis ces résultats. (Voir le chapitre 53 de Körner où ces théorèmes sont démontrés.)

Nous nous tournons maintenant vers l'autre propriété étudiée dans cette section.

Définition 19. Soit $f, g : \mathbb{R} \rightarrow \mathbb{C}$ deux fonctions intégrables au sens de Riemann convolution sur tous les intervalles $[a, b]$ et telles que $\int_{-\infty}^{\infty} |f(t-x)g(x)| dx$ converge pour tout t . Alors la convolution de f et g , notée $f * g$, est :

$$(f * g)(t) = \int_{-\infty}^{\infty} f(t-x)g(x) dx.$$

On notera $\mathcal{C}$ l'ensemble des paires $(f; g)$ satisfaisant les conditions ci-dessus.

Notons que, puisque f est intégrable sur $[t - R, t + S]$ et g sur $[-S, R]$, alors $h(t) = f(t - x)g(x)$ est intégrable sur $[-S, R]$. Puisque $\int_{-\infty}^{\infty} |h(x)| dx$ converge, alors $\int_{-\infty}^{\infty} h(x) dx$ converge aussi.

Une seconde remarque peut également être utile. Elle concerne la notation. Le symbole $f * g$ désigne une fonction qui est obtenue à partir des deux fonctions f et g . Le nom de cette fonction « $f * g$ » contient trois symboles et nous permet de se rappeler l'origine de cette nouvelle fonction. Ainsi $f * g(t)$ est l'évaluation de la convolution de f et de g au point t . Ce n'est pas le produit de f et g évaluées en t !

La propriété la plus importante de la convolution (à part la linéarité) est

$$\widehat{f * g}(\omega) = \hat{f}(\omega)\hat{g}(\omega) \quad (C)$$

où, maintenant, le membre de droite est le produit de $\hat{f}$ et de $\hat{g}$ évaluées en ω . Le but principal de cette section est de démontrer cette propriété. La description du rôle et de la signification de la convolution sera donnée à la prochaine section. Pour l'instant nous obtiendrons les propriétés élémentaires de la convolution qui mèneront à la propriété (C).

la convolution est
symétrique,
linéaire et
et bornée si f l'est.

Lemme 35. Soit $(f; g) \in \mathcal{C}$ et $(f; h) \in \mathcal{C}$.

(i) Alors $(g; f) \in \mathcal{C}$ et $f * g = g * f$.

(ii) Si $\lambda, \mu \in \mathbb{C}$, alors $f * (\lambda g + \mu h) = \lambda(f * g) + \mu(f * h)$.

(iii) si $f, g : \mathbb{R} \rightarrow \mathbb{C}$ sont intégrables sur tout $[a, b]$, f est bornée et $\int_{-\infty}^{\infty} |g(x)| dx$ converge, alors $(f; g) \in \mathcal{C}$ et $f * g$ est bornée.

Démonstration. (i) En changeant la variable d'intégration $y = t - x$:

$$\int_{-S}^R g(t - x)f(x) dx = - \int_{t+S}^{t-R} g(y)f(t - y) dy = \int_{t-R}^{t+S} f(t - y)g(y) dy;$$

ainsi, si $R, S \rightarrow \infty$ et t est tenu fixe, alors l'intégrale ci-dessus est simplement $= f * g(t)$. Donc $(g; f) \in \mathcal{C}$ et $g * f = f * g$.

(ii) Trivial puisque l'intégration est linéaire.

(iii) Si f est bornée, c'est qu'il existe M tel que $|f(x)| \leq M$, pour tout $x \in \mathbb{R}$. Alors $|f(t - x)g(x)| \leq M|g(x)|$ et

$$\left| \int_{-\infty}^{\infty} f(t - x)g(x) dx \right| \leq M \int_{-\infty}^{\infty} |g(x)| dx.$$

Cette dernière intégrale converge par hypothèse et elle ne dépend pas de t . Donc $|f * g|$ est bornée pour tout t . $\square$

continuité de $f * g$

Lemme 36. (i) Soit $f : \mathbb{R} \rightarrow \mathbb{C}$ une fonction continue et bornée, $g : \mathbb{R} \rightarrow \mathbb{C}$ une fonction intégrable sur tout $[a, b]$ telle que $\int_{-\infty}^{\infty} |g(x)| dx$ converge. Alors $f * g$ est bornée et continue.

(ii) Si f est uniformément continue, $f * g$ l'est aussi.

Preuve de (i). Par l'énoncé (iii) du lemme précédent, $f * g$ est bornée. Reste la continuité. Réécrivons la différence de $f * g$ évaluée en deux points voisins.

$$\begin{aligned} |f * g(x) - f * g(y)| &= \left| \int_{-\infty}^{\infty} (f(x-t)g(t) - f(y-t)g(t)) dt \right| \\ &\leq \int_{-\infty}^{\infty} |f(x-t)g(t) - f(y-t)g(t)| dt \end{aligned}$$

et si R est un nombre réel positif :

$$\begin{aligned} &\leq \int_{|t| \leq R} |f(x-t) - f(y-t)| |g(t)| dt \\ &\quad + \int_{|t| > R} |f(x-t) - f(y-t)| |g(t)| dt \\ &\leq \sup_{|t| \leq R} |f(x-t) - f(y-t)| \int_{-\infty}^{\infty} |g(t)| dt + 2M \int_{|t| > R} |g(t)| dt. \end{aligned}$$

Pour tout $\epsilon > 0$, il faut trouver $\delta > 0$ tel que si $|x - y| < \delta$, alors l'expression ci-dessus est $< \epsilon$. Puisque $\int_{-\infty}^{\infty} |g(t)| dt$ converge, il est possible de choisir le nombre $R > 0$ tel que

$$\int_{|t| > R} |g(t)| dt < \frac{\epsilon}{4M + 1}.$$

(Le « 1 » a été ajouté au cas où M serait nul.) Donnée ce R , puisque f est continue, elle l'est uniformément sur l'intervalle fermé $[x - R - 1, x + R + 1]$, et on peut donc trouver $0 < \delta < 1$ tel que

$$|f(u) - f(v)| \leq \frac{\epsilon}{2 \int_{-\infty}^{\infty} |g(t)| dt + 1}$$

si $|u - v| < \delta$ et $u, v \in [x - R - 1, x + R + 1]$. Donc

$$|f(x-t) - f(y-t)| \leq \frac{\epsilon}{2 \int_{-\infty}^{\infty} |g(t)| dt + 1}$$

si $|x - y| < \delta$ et $t \in [-R, R]$. Alors, si $|x - y| < \delta$:

$$|f * g(x) - f * g(y)| < \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon.$$

□

Nous sommes maintenant prêts à montrer la propriété fondamentale de la convolution.

Théorème 37. Soient f et g des fonctions bornées $\in L^1 \cap C$. Alors :

(i) $f * g \in L^1 \cap C$ et $f * g$ est bornée.

(ii)

$$\int_{-\infty}^{\infty} f * g(y) dy = \left(\int_{-\infty}^{\infty} f(y) dy \right) \left(\int_{-\infty}^{\infty} g(y) dy \right)$$

(iii)

$$\int_{-\infty}^{\infty} |f * g(y)| dy \leq \int_{-\infty}^{\infty} |f(y)| dy \int_{-\infty}^{\infty} |g(y)| dy$$

propriété fondamentale
de la convolution

(iv)

$$\widehat{f * g}(\omega) = \hat{f}(\omega) \hat{g}(\omega), \quad \text{pour tout } \omega \in \mathbb{R}. \quad (\text{C})$$

Démonstration. (i) et (ii) : Par le lemme ci-dessus, $f * g$ est bornée et continue. Appliquons le même lemme à $|f| * |g|$:

$$\int_{-\infty}^{\infty} |f(x-t)| |g(t)| dt = |f| * |g|(x) \quad (\text{A1})$$

est donc continue en x . De plus, si on intègre par rapport à x plutôt :

$$\int_{-\infty}^{\infty} |f(x-t)| |g(t)| dx = |g(t)| \underbrace{\int_{-\infty}^{\infty} |f(x)| dx}_{< \infty} \quad (\text{A2})$$

est donc continue en t . Finalement

$$\begin{aligned} \int_{-\infty}^{\infty} \left(\int_{-\infty}^{\infty} |f(x-t)| |g(t)| dx \right) dt &= \int_{-\infty}^{\infty} \left(|g(t)| \int_{-\infty}^{\infty} |f(x-t)| dx \right) dt \\ &= \left(\int_{-\infty}^{\infty} |g(t)| dt \right) \left(\int_{-\infty}^{\infty} |f(x)| dx \right) \\ &< \infty \end{aligned} \quad (\text{B})$$

et $f * g$ est donc élément de $L^1 \cap C$. Ces trois propriétés (A1), (A2) et (B) sont précisément les trois conditions du théorème 29 pour intervertir l'ordre d'intégration :

$$\begin{aligned} \int_{-\infty}^{\infty} f * g(x) dx &= \int_{-\infty}^{\infty} dx \int_{-\infty}^{\infty} f(x-t) g(t) dt \\ &= \int_{-\infty}^{\infty} dt \int_{-\infty}^{\infty} f(x-t) g(t) dx \\ &= \left(\int_{-\infty}^{\infty} g(t) dt \right) \left(\int_{-\infty}^{\infty} f(x) dx \right). \end{aligned}$$

(iii) Puisque

$$|f * g(x)| = \left| \int_{-\infty}^{\infty} f(x-t) g(t) dt \right| \leq \int_{-\infty}^{\infty} |f(x-t)| |g(t)| dt = |f| * |g|(x)$$

pour tout $x \in \mathbb{R}$, alors

$$\int_{-\infty}^{\infty} |f * g(x)| dx \leq \int_{-\infty}^{\infty} |f| * |g|(x) dx = \left(\int_{-\infty}^{\infty} |f(x)| dx \right) \left(\int_{-\infty}^{\infty} |g(x)| dx \right)$$

par (ii) appliqué aux fonctions $|f|$ et $|g|$.

(iv)

$$\begin{aligned}\widehat{f * g}(\omega) &= \int_{-\infty}^{\infty} \left(\int_{-\infty}^{\infty} f(x-t)g(t) dt \right) e^{-i\omega x} dx \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (f(x-t)e^{-i\omega x} e^{i\omega t}) (e^{-i\omega t} g(t)) dt dx \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (f(x-t)e^{-i\omega(x-t)}) (g(t)e^{-i\omega t}) dt dx.\end{aligned}$$

Notons que l'intégrale intérieure est la convolution de $f(x-t)e^{-i\omega(x-t)}$ et de $g(t)e^{-i\omega t}$ et donc par (ii) :

$$\begin{aligned}&= \int_{-\infty}^{\infty} f(x-t)e^{-i\omega(x-t)} dx \int_{-\infty}^{\infty} g(t)e^{-i\omega t} dt \\ &= \hat{f}(\omega)\hat{g}(\omega).\end{aligned}$$

□

Exemple. Le théorème ci-dessus semble joli, sans plus. Pourtant, calculer un exemple, aussi simple soit-il, indique combien ce résultat est non-trivial. Reprenons notre premier exemple d'une transformation de Fourier fait à la section 4.1. Rappelons que

$$f(x) = \begin{cases} 1, & |x| \leq 1 \\ 0, & \text{autrement.} \end{cases}$$

Convolons cette fonction avec elle-même, c'est-à-dire calculons $f * f$.

$$\begin{aligned}f * f(t) &= \int_{-\infty}^{\infty} f(t-x)f(x) dx \\ &= \int_{-1}^1 f(t-x) dx\end{aligned}$$

puisque $f(x)$ est égale à 1 si $|x| \leq 1$ et 0 autrement. En faisant le changement de variable $y = t - x$, on obtient :

$$= \int_{t-1}^{t+1} f(y) dy.$$

A partir de ce point, un peu de minutie s'impose.

- Si $-1 < t+1 < 1$ (et donc $-2 < t < 0$), la borne supérieure de l'intégrale se retrouve entre les deux valeurs extrêmes -1 et 1 entre lesquelles f est non-nulle. Alors $f * f$ dans ce cas devient :

$$\int_{t-1}^{t+1} f(y) dy = \int_{-1}^{t+1} dy = t + 2.$$

- Si $-1 < t-1 < 1$ (et donc $0 < t < 2$), la borne inférieure se trouve entre 1 et -1 et $f * f$ est alors

$$\int_{t-1}^{t+1} f(y) dy = \int_{t-1}^1 dy = 2 - t.$$

- Si $t-1 > 1$ (et donc $t > 2$), la fonction f à intégrer est nulle sur l'intervalle d'intégration et $f * f$ est

$$\int_{t-1}^{t+1} f(y) dy = 0.$$

- Si $t+1 < -1$ (et donc $t < -2$), même conclusion :

$$\int_{t-1}^{t+1} f(y) dy = 0.$$

Ainsi

$$f * f(t) = \begin{cases} 0, & t \geq 2, \\ 2 - t, & 0 \leq t \leq 2, \\ 2 + t, & -2 \leq t \leq 0, \\ 0, & t \leq -2. \end{cases}$$

Tout ceci peut être réécrit plus simplement sous la forme :

$$f(t) = \begin{cases} 2 - |t|, & \text{si } |t| \leq 2, \\ 0, & \text{autrement.} \end{cases}$$

Comme on pourra le constater, le calcul de la convolution est, dans ce cas, relativement simple, mais nécessite quand même un peu d'attention.

Quelle est la transformée de Fourier de $f * f$? Nous l'avons déjà calculée à l'exemple 3. de la section 4.1. En prenant $h = 2$ et $a = \frac{1}{2}$ dans cet exemple, on trouve immédiatement :

$$\widehat{f * f}(\omega) = \frac{2}{\omega^2} (1 - \cos 2\omega).$$

D'autre part, le carré de la transformée de f donne

$$(\hat{f}(\omega))^2 = \frac{4}{\omega^2} \sin^2 \omega$$

et, en utilisant les identités trigonométriques usuelles, on vérifie le théorème de convolution.

Exercices

13. Vérifier que les exigences du théorème 34 sont vérifiées dans le cas de la fonction $g(x, \omega) = e^{-x^2/2} e^{-i\omega x}$. Ceci rend donc rigoureux le calcul de la transformée de la gaussienne. (Voir l'exemple 4. de la section 4.1.)

14*. Soient les fonctions suivantes :

$$f(x) = \begin{cases} 1, & |x| \leq 1 \\ 0, & \text{autrement.} \end{cases}$$

et

$$g(x) = e^{-a|x|}.$$

- (a) Calculer la convolution $f * g$.
- (b) Vérifier le théorème de convolution par calcul direct.

15*. Soit

$$f(x) = \begin{cases} e^{-x}, & x \geq 0, \\ -e^{+x}, & x < 0. \end{cases}$$

- (a) Obtenir $\hat{f}(\omega)$.
- (b) Est-ce que $\hat{f} \in L^1 \cap C$?
- (c) Vérifier par calcul direct que, même si $f \notin L^1 \cap C$ (pourquoi ?), l'identité de Parseval pour la paire $(f, \hat{f})$ est satisfaite.

16*. Supposons que toutes les paires de fonctions intervenant dans cet exercice soient convoluables.

- (a) La convolution de deux fonctions paires possède-t-elle une parité ? Si oui, laquelle ?
- (b) Et la convolution de deux fonctions impaires ?
- (c) Et d'une paire et d'une impaire ?

17*. Soit la gaussienne

$$G_\sigma = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{x^2}{2\sigma^2}}.$$

Montrer que

$$G_\sigma * G_\tau = G_{\sqrt{\sigma^2 + \tau^2}}.$$

18. Soit $\theta : \mathbb{R} \rightarrow \mathbb{R}$ la fonction de Heaviside

$$\theta(x) = \begin{cases} 1, & x \geq 0, \\ 0, & \text{autrement.} \end{cases}$$

Soit maintenant les fonctions $\theta_\lambda^\alpha : \mathbb{R} \rightarrow \mathbb{R}$ définies par

$$\theta_\lambda^\alpha(x) = \theta(x) \frac{x^{\alpha-1}}{\Gamma(\alpha)} e^{-\lambda x},$$

si $\alpha > 0$ et $\lambda > 0$. Montrer que

$$\theta_\lambda^\alpha * \theta_\lambda^\beta = \theta_\lambda^{\alpha+\beta}.$$

19*. Soit

$$f(x) = \begin{cases} 1, & |x| \leq 1 \\ 0, & \text{autrement.} \end{cases}$$

(a) Nous avons calculé en exemple $f * f$ dans la présente section. Obtenir $f * f * f$ et $f * f * f * f$.

(b) Vérifier que

$$\underbrace{f * f * \cdots * f}_{n \text{ fois}}, \quad n = 2, 3, 4$$

est élément de l'ensemble C^{n-1} des fonctions dont les $(n-1)$ dérivées existent et sont continues.

(c) Calculer

$$\int_{-\infty}^{\infty} \left(\frac{\sin \omega}{\omega} \right)^{2n} d\omega$$

pour $n = 1, 2, 3, 4$.

(d) Montrer que

$$\underbrace{f * f * \cdots * f}_{n \text{ fois}}$$

est élément de C^{n-1} , $n \geq 2$.

20. Soit

$$f(x) = \begin{cases} e^{-x}, & x \geq 0, \\ 0, & x < 0. \end{cases}$$

Soit $f_1 = f$ et $f_n = f_{n-1} * f$, $n \geq 2$.

(a) Montrer que

$$f_{n+1} = \begin{cases} \frac{x^n}{n!} e^{-x}, & x \geq 0, \\ 0, & x < 0. \end{cases}$$

(b) Vérifier par calcul direct que $\hat{f}_n(\zeta) = (\hat{f}(\zeta))^n$, pour tout $\zeta \in \mathbb{R}$ et $n \geq 1$.

4.4 Une application : l'âge de la terre selon Lord Kelvin

*Lorsqu'ils parlaient de cette région de la fosse,
les mineurs du pays pâlessaient et baissaient la voix,
comme s'ils avaient parlé de l'enfer ...
Les tailles, au point où l'on en était arrivé,
avaient une température moyenne
de quarante-cinq degrés.*

Germinal, Zola

Cette section a pour but de donner une application de la transformation de Fourier. Cet exemple consistera à résoudre l'équation de la chaleur pour une barre métallique de longueur infinie. Cette solution obtenue, nous relaterons, suivant Körner, le calcul de l'âge de la terre par Lord Kelvin. Ce calcul est un des faits saillants de l'histoire des sciences puisqu'il fut fait à un moment où deux évaluations sérieuses, mais irréconciliables, de cet âge étaient proposées, celle de Kelvin et celle de Darwin.

Commençons par trois petits lemmes.

Lemme 38. Si $f : \mathbb{R} \rightarrow \mathbb{C}$ est une fonction différentiable et f, f' et g sont bornées avec $f, f', g \in L^1 \cap C$, alors $f * g$ est différentiable avec $(f * g)' = f' * g$.

Démonstration. En différentiant directement la convolution de f et g :

$$\begin{aligned} (f * g)' &= \frac{d}{dx} \int_{-\infty}^{\infty} f(x-t)g(t)dt \\ &= \int_{-\infty}^{\infty} \frac{d}{dx} f(x-t)g(t)dt \\ &= \int_{-\infty}^{\infty} f'(x-t)g(t)dt \\ &= (f' * g)(x). \end{aligned}$$

Le lecteur pourra vérifier les hypothèses du théorème 34 que nous avons utilisé à la première étape. $\square$

Lemme 39. Soit $f : \mathbb{R} \rightarrow \mathbb{C}$. Notons par $\mathcal{T}f$ la fonction $(\mathcal{T}f)(t) = tf(t)$. Si f et $(\mathcal{T}f) \in L^1 \cap C$, alors $\hat{f}$ est différentiable avec $\frac{d\hat{f}}{d\zeta}(\zeta) = -i(\widehat{\mathcal{T}f})(\zeta)$.

Démonstration. A nouveau, par un calcul direct :

$$\begin{aligned} \frac{d}{d\zeta} \int_{-\infty}^{\infty} f(t)e^{-i\zeta t} dt &= \int_{-\infty}^{\infty} f(t)(-it)e^{-i\zeta t} dt \\ &= -i \int_{-\infty}^{\infty} (\mathcal{T}f)(t)e^{-i\zeta t} dt \\ &= -i(\widehat{\mathcal{T}f})(\zeta). \end{aligned}$$

La même vérification des conditions du théorème 34 s'impose. $\square$

Lemme 40. (i) Si $f \in L^1 \cap C$ et $R > 0$, alors $f_R(x) = Rf(Rx)$ est $\in L^1 \cap C$ et $\hat{f}_R(\zeta) = \hat{f}(\zeta/R)$.

(ii) Si $f : \mathbb{R} \rightarrow \mathbb{C}$ est différentiable une fois avec $f, f' \in L^1 \cap C$ et $f(x) \rightarrow 0$ quand $|x| \rightarrow \infty$, alors $\hat{f}'(\zeta) = i\zeta \hat{f}(\zeta)$.

Attention : il faut distinguer les deux fonctions $\hat{f}'$ (décrite dans ce lemme) et $(\hat{f})'$ (décrite dans le lemme précédent).

Démonstration. (i)

$$\hat{f}_R(\zeta) = \int_{-\infty}^{\infty} Rf(Rx)e^{-i\zeta x} dx$$

et avec le changement de variable $u = Rx$:

$$\begin{aligned} &= \frac{1}{R} \int_{-\infty}^{\infty} Rf(u)e^{-i\zeta u/R} du \\ &= \hat{f}(\zeta/R). \end{aligned}$$

(ii) En intégrant par parties :

$$\begin{aligned} \int_{-\infty}^{\infty} f'(x)e^{-i\zeta x} dx &= f(x)e^{-i\zeta x} \Big|_{-\infty}^{\infty} + i\zeta \int_{-\infty}^{\infty} f(x)e^{-i\zeta x} dx \\ &= 0 + i\zeta \hat{f}(\zeta). \end{aligned}$$

$\square$

Après ces lemmes que nous utiliserons ci-dessous, nous sommes prêts à poser le problème : quelle est la température le long d'une barre métallique de longueur infinie si, au temps $t = 0$, la température est donné par $G : \mathbb{R} \rightarrow \mathbb{R}$. Mathématiquement la solution de ce problème consiste en la résolution de l'équation de la chaleur

$$\frac{\partial \phi}{\partial t}(x, t) = \kappa \frac{\partial^2 \phi}{\partial x^2}(x, t), \quad x \in \mathbb{R} \text{ et } t \in \mathbb{R}^+$$

avec conditions initiales :

$$\phi(x, t) \rightarrow G(x), \quad \text{lorsque } t \rightarrow 0^+.$$

Nous avons déjà résolu des équations aux dérivées partielles au chapitre 3 : équation de la corde vibrante, équation du tambour (carré et rond), équation de la chaleur pour une barre de longueur L , etc. Toutes ces équations avaient comme point commun que la ou les variables d'espace étaient limitées à un domaine fermé borné, c'est-à-dire à un ensemble compact. Maintenant nous avons affaire à une

barre de longueur infinie et les séries de Fourier ne pourront pas s'appliquer. Notre stratégie de résolution est représentée au diagramme suivant.

$$\begin{array}{ccc}
 D_2\phi = \kappa D_1 D_1\phi & \xrightarrow{\quad ? \quad} & \phi \\
 \text{transformation de Fourier} \downarrow & & \uparrow \text{transformation inverse} \\
 \mathcal{F}_1(D_2\phi) = -\kappa\zeta^2(\mathcal{F}_1\phi) & \xrightarrow{\text{résolution d'une édo}} & \mathcal{F}_1\phi
 \end{array}$$

Expliquons d'abord la notation. Puisque ϕ est une fonction de deux variables, la notation $\hat{\phi}$ est confuse. Nous noterons donc

$$(\mathcal{F}_1\phi)(\zeta, t) = \int_{-\infty}^{\infty} \phi(x, t) e^{-i\zeta x} dx$$

pour tout t dans le domaine de ϕ . De plus nous noterons

$$\begin{aligned}
 D_1\phi &= \frac{\partial}{\partial x}\phi \\
 D_2\phi &= \frac{\partial}{\partial t}\phi.
 \end{aligned}$$

Avec ces nouvelles notations, l'équation de la chaleur se lit simplement

$$D_2\phi = \kappa D_1 D_1\phi.$$

La flèche « $\rightarrow$ » est le problème que nous désirons résoudre : trouver la solution ϕ à partir de l'équation aux dérivées partielles $D_2\phi = \kappa D_1 D_1\phi$. Pour cela nous allons suivre le chemin indiqué par les trois autres flèches. Nous ferons la transformation de Fourier de ϕ sur la variable spatiale x ; ceci correspond à la flèche « $\downarrow$ ». L'équation résultante sera une simple équation différentielle ordinaire en t . Sa solution « $\rightarrow$ » sera aisée. La variable ζ , introduite lors de la transformation de Fourier en la variable x , y apparaîtra comme un simple paramètre. Ainsi, la transformation de Fourier aura *séparé* les coordonnées spatiale et temporelle. En faisant la transformation inverse « $\uparrow$ », nous obtiendrons la solution générale ϕ recherchée.

La première étape de notre solution consiste donc à faire la transformation de Fourier sur la variable spatiale de chacun des membres de l'équation

$$\mathcal{F}_1(D_2\phi)(\zeta, t) = \kappa \mathcal{F}_1(D_1 D_1\phi)(\zeta, t)$$

qui, par le l'énoncé (ii) du lemme précédent, peut être écrite :

$$= (i\zeta)^2 \kappa (\mathcal{F}_1\phi)(\zeta, t).$$

Supposons que l'on puisse sortir le $\frac{\partial}{\partial t}$ de sous le signe d'intégration. Alors

$$\mathcal{F}_1(D_2\phi) = \int_{-\infty}^{\infty} \frac{\partial \phi}{\partial t}(x, t) e^{-i\zeta x} dx = \frac{\partial}{\partial t} \int_{-\infty}^{\infty} \phi(x, t) e^{-i\zeta x} dx = D_2(\mathcal{F}_1\phi)(\zeta, t)$$

équation différentielle
pour $\mathcal{F}_1\phi$

et donc

$$D_2(\mathcal{F}_1\phi)(\zeta, t) = -\kappa\zeta^2(\mathcal{F}_1\phi)(\zeta, t).$$

Ainsi, si toutes les hypothèses précédentes sont satisfaites, nous sommes amenés à résoudre

$$u'(t) = -\kappa\zeta^2 u(t)$$

où $u(t) = (\mathcal{F}_1\phi)(\zeta, t)$ pour un ζ fixé. La solution est immédiate :

$$(\mathcal{F}_1\phi)(\zeta, t) = A(\zeta)e^{-\kappa\zeta^2 t}$$

où $A(\zeta)$ est la constante d'intégration. Elle dépend de ζ puisque le problème pour $u(t)$ est pour chaque valeur de ζ . Ainsi A est une fonction de $\mathbb{R} \rightarrow \mathbb{C}$.

conditions initiales

Si ϕ se comporte suffisamment bien

$$\int_{-\infty}^{\infty} |\phi(x, t) - \phi(x, 0)| dx \rightarrow 0, \quad \text{lorsque } t \rightarrow 0^+$$

et donc

$$\begin{aligned} |\mathcal{F}_1\phi(\zeta, t) - \mathcal{F}_1\phi(\zeta, 0)| &= \left| \int_{-\infty}^{\infty} (\phi(x, t) - \phi(x, 0))e^{-i\zeta x} dx \right| \\ &\leq \int_{-\infty}^{\infty} |\phi(x, t) - \phi(x, 0)| |e^{-i\zeta x}| dx \\ &\rightarrow 0 \end{aligned}$$

pour $t \rightarrow 0^+$. Alors, puisque $e^{-\kappa\zeta^2 t} \rightarrow 1$ pour $t \rightarrow 0^+$:

$$(\mathcal{F}_1\phi)(\zeta, 0) = \int_{-\infty}^{\infty} G(x)e^{-i\zeta x} dx = \hat{G}(\zeta) = A(\zeta).$$

Ainsi, pour t fixé, si $F(x) = \phi(x, t)$, la solution est

$$\hat{F}(\zeta) = \hat{G}(\zeta)e^{-\kappa\zeta^2 t}$$

ou encore $(\mathcal{F}_1\phi)(\zeta, t) = \hat{G}(\zeta)e^{-\kappa\zeta^2 t}$, ce qui conclut la seconde étape « $\rightarrow$ ».

$\hat{F}$ est le produit
de deux transformées
de Fourier

Rappelons maintenant que la transformée de Fourier d'une gaussienne est également une gaussienne. (Voir l'exemple 4. de la section 4.1.) Ainsi la paire

$$E(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}} \quad \hat{E}(\zeta) = e^{-\frac{\zeta^2}{2}}$$

est une fonction et sa transformée de Fourier. Par le lemme 40 ci-dessus, la paire

$$E_R(x) = RE(Rx) \quad \hat{E}_R(\zeta) = \hat{E}(\zeta/R)$$

est conséquemment aussi une fonction et sa transformée de Fourier. Posons

$$R = \frac{1}{\sqrt{2\kappa t}}.$$

Alors

$$\hat{E}_R(\zeta) = e^{-\frac{1}{2}\left(\frac{\zeta}{1/\sqrt{2\kappa t}}\right)^2} = e^{-\kappa\zeta^2 t}.$$

Nous pouvons donc écrire la solution trouvée ci-dessus en termes de cette notation sous la forme :

$$\hat{F}(\zeta) = \hat{G}(\zeta) \widehat{E_{1/\sqrt{2\kappa t}}}(\zeta)$$

et par le théorème de convolution 37 :

$$\hat{F}(\zeta) = (\widehat{G * E_{1/\sqrt{2\kappa t}}})(\zeta)$$

et par le théorème d'inversion 31

$$F(x) = (G * E_{1/\sqrt{2\kappa t}})(x)$$

pour tout x et t fixé, c'est-à-dire :

la solution

$$\phi(x, t) = (G * E_{1/\sqrt{2\kappa t}})(x)$$

qui ne dépend plus que d'une convolution ! Nous venons donc de démontrer :

Théorème 41. Soit $G \in L^1 \cap C$. Si $\phi : \mathbb{R} \times \mathbb{R}^+ \rightarrow \mathbb{C}$ telle que $D_1\phi, D_1D_1\phi$ et $D_2\phi$ existe avec

- (i) $\phi(\cdot, t), D_1\phi(\cdot, t), D_1D_1\phi(\cdot, t), D_2\phi(\cdot, t) \in L^1 \cap C$,
- (ii) $\phi(x, t), D_1\phi(x, t) \rightarrow 0$ quand $|x| \rightarrow \infty$ pour tout $t \in \mathbb{R}^+$,
- (iii)

$$\frac{\partial \phi}{\partial t} = \kappa \frac{\partial^2 \phi}{\partial x^2}$$

pour tout $(x, t) \in \mathbb{R} \times \mathbb{R}^+$ et

- (iv) $\int_{-\infty}^{\infty} |\phi(x, t) - G(x)| dx \rightarrow 0$ quand $t \rightarrow 0^+$,
- alors

$$\begin{aligned} \phi(x, t) &= (G * E_{1/\sqrt{2\kappa t}})(x) \\ &= \frac{1}{\sqrt{2\kappa t}} \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} G(x-y) e^{-y^2/4\kappa t} dy \end{aligned}$$

pour tout $(x, t) \in \mathbb{R} \times \mathbb{R}^+$.

Avant de relater la petite histoire du calcul de l'âge de la terre par Lord Kelvin à partir de ce théorème, il est utile de s'arrêter pour comprendre, sur cette solution, la signification de l'opération « * » de convolution. Ce que dit le théorème est que, si en $t = 0$, le profil de la température sur une barre infinie est $G(x)$, alors, à un temps t ultérieur, la distribution sera

$$\phi(x, t) = \left(G * E_{1/\sqrt{2\kappa t}} \right) (x)$$

et cette convolution se lit

$$= \frac{1}{\sqrt{2\kappa t}} \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} G(x-y) e^{-y^2/4\kappa t} dy$$

Cette solution peut être lue en mots comme suit : au temps t , la température mesurée au point x sera une somme des contributions en tout point de la barre, mais plus les points seront distants, moins leur influence sera grande. En effet, la variable d'intégration y représente la distance du point x . (G est évaluée en $(x-y)$.) La contribution des points à distance $\pm y$ de x est pondérée par $e^{-y^2/4\kappa t}$ qui est d'autant plus petite que y est grand. Ainsi la fonction $e^{-y^2/4\kappa t}$ dit comment s'est « propagée » la chaleur après un temps t . Pour cette raison $E_{1/\sqrt{2\kappa t}}$ est appelé le *noyau de la chaleur*. Lorsque le temps passe, la gaussienne s'aplatit et des portions toujours plus larges des températures initiales auront un effet non-négligeable sur le point x . La convolution est donc une opération naturelle ici ; les deux fonctions G et $E_{1/\sqrt{2\kappa t}}$ sont évaluées en des points différents dans l'intégrand car leur rôle est différent : G représente la contribution de la distribution originale de la température et $E_{1/\sqrt{2\kappa t}}$ mesure la pondération à accorder à cette contribution en fonction de la distance.

noyau de la chaleur

L'étalement (ou la diffusion) de la chaleur peut être facilement vue sur un graphique. Supposons qu'à l'instant $t = 0$, la barre de longueur infinie soit maintenue à température nulle à l'exception d'un segment de longueur 2 qui est à température d'un degré. La figure 4.3 montre le profil de la température à des instants suivants. Le temps est mesuré en unité de κ . Plus le temps croît, plus le profil s'aplatit et la chaleur diffuse, ce qui est évidemment conforme à l'expérience de tous les jours. Au temps $t = 10$ (le plus grand temps tracé), le saut initial dans la température est pratiquement nivelé.

Comment peut-on utiliser ce calcul pour évaluer l'âge de la terre ? Un peu d'histoire sera nécessaire pour révéler la difficulté des arguments mis de l'avant. Trois sciences sont ici en jeu : la géologie, la physique et la biologie.

Les premières évaluations de l'âge de la terre reposent sur des hypothèses assez grossières qui seront raffinées par la suite. Si la terre et tout le règne vivant sont créés simultanément (ou en l'espace d'une semaine), l'âge de la terre peut être aussi petit que la durée de l'histoire écrite, à savoir plus petit que 10^4 années. Ceci est donc une borne inférieure. Les géologues du début du 19^{ième} siècle savent que de grandes portions de la terre étaient enfouies sous l'océan à une époque donnée de l'histoire de la planète comme l'indiquent la présence de fossiles de poisson loin des côtes des océans. Pour expliquer le changement de position des océans, les géologues de cet époque eurent recours aux changements violents et cataclysmiques. Ces événements étant de grande envergure, il en aurait fallu assez peu pour changer la face de la terre et l'âge de la terre s'en voit augmenter que d'un peu, disons jusqu'à 10^5 années.

Le géologue écossais Lyell (1797-1875) avance l'hypothèse contraire que les changements sont lents et continus et se sont produits au cours des temps à la même vitesse qu'on les observe aujourd'hui. Cette évolution *uniforme* serait

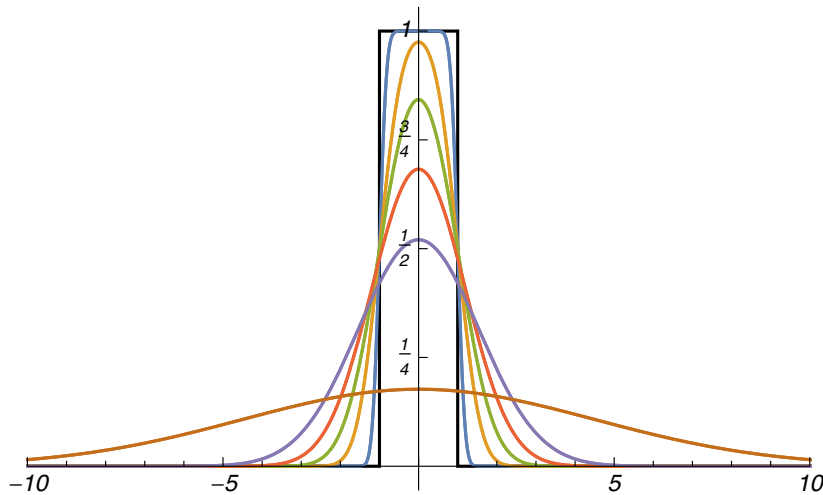


Figure 4.3 – Profils de la température aux temps $\kappa t = 0, \frac{1}{100}, \frac{1}{10}, \frac{1}{4}, \frac{1}{2}, 1$ et 10.

donc due aux volcans, aux tremblements de terre, à l'érosion. Il est clair que cette théorie nécessite un temps beaucoup plus considérable pour la formation de la planète tel qu'on la connaît : 10^6 - 10^9 années. Même si cette théorie peut expliquer la disparition d'espèces, elle ne dit rien sur l'apparition de nouvelles. Une contribution des biologistes est alors nécessaire.

Darwin, naturaliste anglais, écrit en 1859 un livre qui révolutionnera le monde scientifique : *On the Origin of Species*. Il y soutient, à l'aide d'un nombre important d'observations scientifiques, que les espèces vivantes naissent, se transforment, disparaissent. Le rythme de cette évolution est peut-être difficile à chiffrer mais son existence est pratiquement inattaquable avec les données déjà disponibles au siècle dernier. Cette évolution est lente et demande sûrement une borne *inférieure* sur l'âge de l'univers d'environ 10^8 années, probablement plus.

Charles Darwin
(1809-1882)

Entrent maintenant en scène les observations physiques. Lord Kelvin note tout d'abord que si le soleil brûlait un combustible tels ceux qu'on connaît sur terre, il serait éteint depuis longtemps. Kelvin fait donc l'hypothèse que le soleil transforme lentement son énergie potentielle gravitationnelle en chaleur. Même avec cette source d'énergie potentielle, la durée d'illumination de la terre par le soleil ne peut dépasser beaucoup $\sim 10^8$ années, et est certainement plus petite que 5×10^8 . Notons que, puisque le soleil se refroidit, la terre devait recevoir plus d'énergie à une époque lointaine et les phénomènes géologiques à cette époque devaient être plus vigoureux.

William Thomson,
Lord Kelvin (1824-1907)

Une seconde observation de Lord Kelvin concerne la terre elle-même. Puisque la croûte terrestre reste plus ou moins à la même température et que le centre de la terre est plus chaud (ceci était déjà connu par les explorations minières), il faut

admettre que la terre se refroidit. On pourrait imaginer qu'un combustible brûle au centre de la planète, mais cette hypothèse ne tient pas car aucun combustible chimique ne peut être en quantité suffisante sur terre pour soutenir une réaction pendant environ 10^8 années, période de vie approximative du soleil. Donc la terre se refroidit !

Les temps géologiques (et biologiques) n'ont pas pu commencer avant que la croûte terrestre se forme. Est-ce qu'une coquille s'est formée avant tout (comme la coquille d'un oeuf avec un centre liquide) ou est-ce que les courants de convection assurent un refroidissement uniforme ? Kelvin a montré que la première possibilité est à rejeter. Donc la terre doit avoir durci à peu près uniformément à la même époque. (C'est ce qu'on croît encore aujourd'hui avec quelques légers amendements.) Ainsi les hypothèses de Kelvin sont

$$\begin{aligned}\frac{\partial \psi}{\partial t}(\vec{x}, t) &= \kappa \nabla^2 \psi(\vec{x}, t), \\ \psi(\vec{x}, 0) &= \theta_0 = \text{température de fusion}, \quad \forall |\vec{x}| < R \\ \psi(\vec{x}, t) &= C = \text{température constante à la surface}, \quad \forall |\vec{x}| = R \text{ et } t > 0.\end{aligned}$$

Les symboles utilisés sont : R , le rayon de la terre, et $t = 0$, le moment dans l'histoire où la terre devient solide. La première équation est évidemment l'équation de la chaleur. La seconde affirme que la transition liquide à solide s'est faite (en $t = 0$) uniformément partout à l'intérieur de la terre. Enfin la troisième correspond à l'observation d'une température constante à la surface de la planète.

Kelvin savait résoudre ce problème en coordonnées sphériques mais des calculs préliminaires montrent que le refroidissement ne peut être important qu'en surface. Il est donc possible de négliger la courbure de la terre et de remplacer le système d'équations ci-dessus par

$$\begin{aligned}\frac{\partial \theta}{\partial t} &= \kappa \frac{\partial^2 \theta}{\partial y^2} \\ \theta(y, 0) &= \theta_0, \quad y > 0, \\ \theta(0, t) &= 0, \quad t > 0,\end{aligned}$$

où maintenant θ est une fonction de deux variables, y et t . La coordonnée y qui remplace la coordonnée radiale s'annule en $x = R$ et croît vers le centre de la terre. La température à la surface de la terre a été posée à 0 ce qui est possible puisque l'équation différentielle est linéaire et que « $\theta = \text{constante}$ » est une solution.

Une différence importante avec l'énoncé du théorème 41 est que la fonction ϕ du théorème est définie sur $\mathbb{R} \times \mathbb{R}^+$ et que θ l'est sur $(\mathbb{R}^+ \cup \{0\}) \times \mathbb{R}^+$, c'est-à-dire que y ne prend des valeurs que ≥ 0 . La condition supplémentaire $\theta(0, t) = 0$ pour $t > 0$ nous permet de sortir de l'impasse. Cette condition permet en effet d'étendre la solution en $y > 0$ en demandant que la solution soit impaire en y

pour tout t . Elle vérifiera donc

$$G(y) = \begin{cases} \theta_0, & y > 0, \\ 0, & y = 0, \\ -\theta_0, & y < 0. \end{cases}$$

La solution est alors

$$\theta(y, t) = \frac{\theta_0}{2\sqrt{\pi\kappa t}} \int_0^\infty \left(\exp\left(-\frac{(y-x)^2}{4\kappa t}\right) - \exp\left(-\frac{(y+x)^2}{4\kappa t}\right) \right) dx.$$

Alors Kelvin utilisa les observations sur la conductivité κ , le point de fusion de la croûte terrestre et le taux de variation de la température $\frac{\partial\theta}{\partial y}$ en fonction de la profondeur qu'on peut mesurer dans les mines. Il détermina que le $t_{\text{aujourd'hui}}$, c'est-à-dire $\text{âge}_{\text{terre}}$, est

$$\text{âge}_{\text{terre}} \approx 10^8$$

et sûrement

$$\text{âge}_{\text{terre}} < 4 \times 10^8,$$

juste un peu trop court pour les biologistes et les géologues...

Il faudra attendre le début de ce siècle pour résoudre l'énigme. La découverte de la radioactivité et les travaux de Rutherford (1871-1937) ajoutent une source d'énergie sur la terre qui n'est pas un combustible chimique tel que le concevait Kelvin et qui apporte une quantité importante de chaleur à la terre. Alors, pour refroidir la terre à la température que l'on mesure aujourd'hui, il aura fallu plus de temps...

21. Montrer que la solution θ étendue pour $y \in \mathbb{R}$ satisfait aux conditions :

Exercices

$$\begin{aligned} \frac{\partial\theta}{\partial t} &= \kappa \frac{\partial^2\theta}{\partial y^2} \\ \theta(y, 0) &= \theta_0, & y > 0, \\ \theta(0, t) &= 0, & t > 0. \end{aligned}$$

Appendice A

Fonctions eulériennes

Cet appendice ne comporte que deux courtes sections. Nous y introduisons deux fonctions : $\Gamma(z)$ et $B(p, q)$. Elles ont peu à voir avec l'analyse de Fourier ou le problème de Sturm-Liouville. Mais nous en aurons besoin pour entreprendre l'étude de certaines fonctions spéciales.

A.1 Les fonctions Γ et $B(p, q)$

La fonction factorielle $n!$ est définie sur les entiers naturels $\mathbb{N}$. La fonction gamma (Γ) que nous introduisons ici en deux étapes étend cette fonction factorielle à toute la droite réelle (en fait à tout le plan complexe).

Définition 20 (a). La fonction $\Gamma : \mathbb{R}^+ \rightarrow \mathbb{R}$ est définie par

la fonction $\Gamma(z)$

$$\Gamma(z) = \int_0^\infty e^{-x} x^{z-1} dx, \quad z > 0.$$

Remarquons tout d'abord que cette intégrale est bien définie ; pour un voisinage $(0, \delta]$ de 0, ($\delta > 0$), la fonction $e^{-x} x^{z-1} < x^{z-1}$ dont la primitive est x^z/z . Pour un voisinage de l'infini $[M, +\infty)$, l'intégrand se comporte comme

$$e^{-x} x^{z-1} = e^{-x+(z-1)\ln x} = e^{-x(1-(z-1)\frac{\ln x}{x})} \xrightarrow{x \rightarrow \infty} 0$$

puisque, par la règle de L'Hospital

$$\lim_{x \rightarrow \infty} (z-1) \frac{\ln x}{x} = (z-1) \lim_{x \rightarrow \infty} \frac{1/x}{1} = 0.$$

Donc l'intégrale existe.

Vous n'êtes peut-être pas familier avec l'idée d'une fonction définie par le résultat d'une intégrale. Il est donc utile de souligner que la variable x dans la définition n'est qu'une variable d'intégration et elle est donc muette. C'est à z qu'est évaluée la fonction Γ et z n'intervient que comme exposant de x dans l'intégrand.

continuité

La fonction Γ est continue sur $\mathbb{R}^+$. Cet énoncé repose sur un théorème dû à Lebesgue qui n'est pas démontré habituellement dans un premier cours d'analyse. Nous le citerons sans preuve.

Lebesgue, Henri-Léon
(1875-1941)

Théorème 42 (Lebesgue). *Si $f(z, t)$ est continue en z au point z_0 pour toutes les valeurs de t et si $f(z, t)$ est majorée en module par une fonction $g(t) \geq 0$ intégrable, alors $F(z) = \int_a^b f(z, t) dt$ est continue en z au point z_0 .*

Les hypothèses de ce théorème sont vérifiées pour la fonction $f(z, t) = e^{-t} t^{z-1}$. Cette fonction est en effet continue en tout point $(z, t) \in \mathbb{R}^+ \times \mathbb{R}^+$ et elle est bornée pour un z fixé par :

$$|e^{-t} t^{z-1}| \leq \begin{cases} t^{\alpha-1}, & 0 \leq t \leq 1, \\ t^{A-1} e^{-t}, & 1 \leq t < \infty, \end{cases}$$

pour un $0 < \alpha < z < A < \infty$.

positivité

La fonction Γ est positive sur $\mathbb{R}^+$ puisque l'intégrand est positif. D'autre part, si $z > 0$, alors :

$$\Gamma(z+1) = \int_0^\infty e^{-x} x^z dx = -e^{-x} x^z \Big|_0^\infty + z \int_0^\infty e^{-x} x^{z-1} dx$$

où, à la seconde étape, nous avons intégré par parties. Le terme à évaluer est nul en zéro et, à l'infini, il s'annule également puisque, comme nous venons de le voir : $e^{-x} x^z \xrightarrow{x \rightarrow \infty} 0$. Donc :

$$\Gamma(z+1) = z\Gamma(z), \quad z > 0. \quad (\Gamma_1)$$

Une valeur est relativement aisée à obtenir :

$$\Gamma(1) = \int_0^\infty e^{-x} x^{z-1} \Big|_{z=1} dx = \int_0^\infty e^{-x} dx = -e^{-x} \Big|_0^\infty = 1.$$

Donc

$$\Gamma(1) = 1.$$

$\Gamma(n) = (n-1)!$

Avec la relation fonctionnelle, ceci permet d'obtenir récursivement $\Gamma(n)$ pour tout $n \in \mathbb{N}$. La réponse est

$$\Gamma(n) = (n-1)!$$

que l'on montre par induction. Le résultat est vrai pour $n = 2$: $\Gamma(2) = 1 \cdot \Gamma(1) = 1 = (2-1)!$. Si le résultat est vrai pour n , alors :

$$\Gamma(n+1) = n\Gamma(n) = n(n-1)! = n!$$

Donc $\Gamma(z)$ est une fonction continue de z qui, pour les valeurs entières de son argument est la fonction factorielle définie sur les entiers seulement. Puisqu'on définit la fonction factorielle $n!$ par $1 \cdot 2 \cdot 3 \cdots (n-1) \cdot n$, il faut, à strictement parler, que $n \geq 1$. Il est usuel, cependant, d'étendre la notation avec le signe « ! »

au cas limite $0!$ en définissant $0! = 1$. Une autre notation est souvent utilisée dans l'étude des fonctions spéciales. Il s'agit du double factoriel « $!!$ ». Cette fonction est définie par :

$$\begin{aligned}(2s-1)!! &= (2s-1) \cdot (2s-3) \cdots 3 \cdot 1 \\ (2s)!! &= (2s) \cdot (2s-2) \cdots 4 \cdot 2\end{aligned}$$

Nous sommes maintenant prêts à prolonger Γ sur l'axe réel entier. $\Gamma(z)$ n'est présentement définie que pour $z > 0$. Cependant, il est possible d'utiliser la relation fonctionnelle pour étendre $\Gamma(z)$ à l'axe négatif. Soit $z \in \mathbb{R}$ tel que $z \neq 0, -1, -2, -3, \dots$. Alors pour n tel que $z+n > 0$, la quantité

$$\frac{\Gamma(z+n)}{(z+n-1)(z+n-2)\dots(z+1)z}$$

est bien définie et ne dépend pas du n choisi. Par ceci, je veux dire que si un autre entier $m > n$ avait été choisi, le résultat aurait été le même. Si $m = n + p$ avec $m, p \in \mathbb{N}$, alors

$$\begin{aligned}\frac{\Gamma(z+n+p)}{(z+n+p-1)(z+n+p-2)\dots(z+1)z} &= \\ &= \frac{(z+n+p-1)(z+n+p-2)\dots(z+n)\Gamma(z+n)}{(z+n+p-1)(z+n+p-2)\dots(z+1)z} \\ &= \frac{\Gamma(z+n)}{(z+n-1)(z+n-2)\dots(z+1)z}.\end{aligned}$$

Ainsi nous proposons donc la définition :

Définition 20 (b). La fonction $\Gamma : \mathbb{R} \setminus \{0, -1, -2, -3, \dots\} \rightarrow \mathbb{R}$ est définie par la fonction Γ

$$\Gamma(z) = \begin{cases} \int_0^\infty e^{-x} x^{z-1} dx, & z > 0 \\ \frac{\Gamma(z+n)}{(z+n-1)(z+n-2)\dots(z+1)z}, & z < 0, \end{cases}$$

où n est choisi tel que $z+n > 0$.

Cette définition assure que la relation fonctionnelle demeure vraie sur l'axe négatif également.

Cette définition semble indiquer cependant que la fonction $\Gamma(z)$ aura des pôles en $z = 0, -1, -2, -3, \dots$. Montrons-le. Soit $z = -n$, n un entier ≥ 0 . Alors pour u petit :

$$\begin{aligned}\Gamma(z+u) &= \Gamma(-n+u) \\ &= \frac{\Gamma(-n+u+n+1)}{(u+1-1)(u-1)(u-2)\dots(u-n)} \\ &\xrightarrow{u \rightarrow 0} \frac{\Gamma(1)}{u(-1)(-2)\dots(-n)} \\ &= \frac{(-1)^n}{un!}\end{aligned}$$

Donc $\Gamma(z)$ a des pôles simples en tous les points $-n$ entiers ($n \geq 0$) et ce sont les seuls.

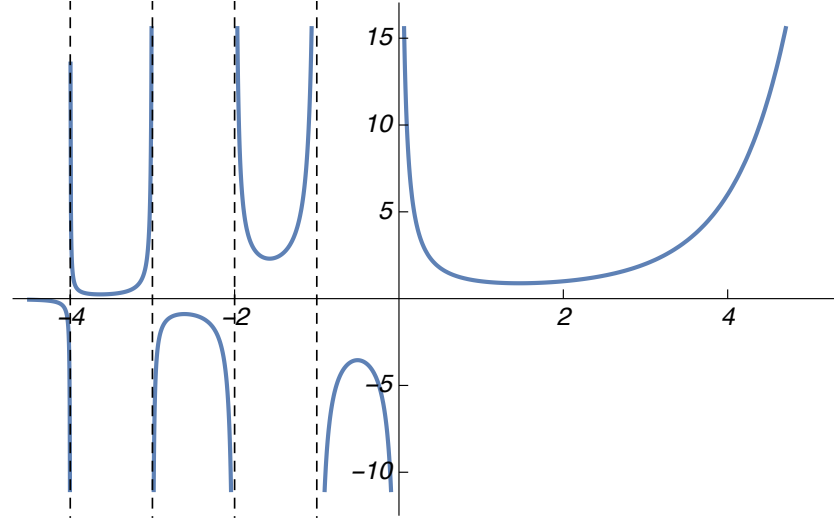


Figure A.1 – La fonction Γ

On peut réécrire $\Gamma(z)$ par changement de la variable d'intégration $x = t^2$ et donc $dx = 2t dt$:

$$\Gamma(z) = \int_0^\infty e^{-x} x^{z-1} dx = 2 \int_0^\infty e^{-t^2} t^{2z-2} t dt = 2 \int_0^\infty e^{-t^2} t^{2z-1} dt.$$

Gauss, Carl Friedrich
(1777-1855)

Cette dernière forme est appelée intégrale de Gauss. Elle fait intervenir la fonction gaussienne e^{-t^2} qui est une distribution centrale en statistique. Ainsi, ces intégrales interviennent souvent dans ce domaine. Une valeur particulière joue un rôle important :

$$\frac{1}{2}\Gamma\left(\frac{1}{2}\right) = \int_0^\infty e^{-t^2} dt$$

que nous allons évaluer dans ce qui suit.

Une façon d'évaluer $\Gamma(\frac{1}{2})$ est par la seconde fonction que cet appendice introduit. Pour cela calculons le produit de $\Gamma(p)$ et $\Gamma(q)$, $p, q > 0$:

$$\Gamma(p) = 2 \int_0^\infty e^{-u^2} u^{2p-1} du \quad \text{et} \quad \Gamma(q) = 2 \int_0^\infty e^{-v^2} v^{2q-1} dv.$$

Donc

$$\begin{aligned} \Gamma(p)\Gamma(q) &= 4 \int_0^\infty e^{-u^2} u^{2p-1} du \int_0^\infty e^{-v^2} v^{2q-1} dv \\ &= 4 \int_0^\infty \int_0^\infty e^{-u^2-v^2} u^{2p-1} v^{2q-1} dudv. \end{aligned}$$

Les coordonnées polaires sont

$$\begin{aligned} u &= r \cos \theta \\ v &= r \sin \theta \end{aligned}$$

et le jacobien du changement de variables

$$\frac{\partial(u, v)}{\partial(r, \theta)} = \begin{vmatrix} \frac{\partial u}{\partial r} & \frac{\partial u}{\partial \theta} \\ \frac{\partial v}{\partial r} & \frac{\partial v}{\partial \theta} \end{vmatrix} = \begin{vmatrix} \cos \theta & -r \sin \theta \\ \sin \theta & r \cos \theta \end{vmatrix} = r;$$

ainsi ces intégrales se réécrivent :

$$\begin{aligned} \Gamma(p)\Gamma(q) &= 4 \iint_{(u, v \geq 0)} e^{-r^2} r^{2p-1} \cos^{2p-1} \theta r^{2q-1} \sin^{2q-1} \theta r dr d\theta \\ &= 4 \iint e^{-r^2} r^{2p+2q-1} \cos^{2p-1} \theta \sin^{2q-1} \theta dr d\theta \\ &= 4 \int_0^\infty e^{-r^2} r^{2(p+q)-1} dr \int_0^{\frac{\pi}{2}} \cos^{2p-1} \theta \sin^{2q-1} \theta d\theta. \end{aligned}$$

Les bornes de l'intégrale par rapport à θ ont été choisies pour ne couvrir que le premier quadrant $u, v \geq 0$. Revenons à la fonction Γ pour l'intégrale en r ; le changement de variables

$$r^2 = s \quad 2r dr = ds$$

donne

$$\begin{aligned} &= 4 \int_0^\infty \frac{ds}{2\sqrt{s}} e^{-s} s^{p+q-\frac{1}{2}} \int_0^{\frac{\pi}{2}} \cos^{2p-1} \theta \sin^{2q-1} \theta d\theta \\ &= 2 \int_0^\infty e^{-s} s^{p+q-1} ds \int_0^{\frac{\pi}{2}} \cos^{2p-1} \theta \sin^{2q-1} \theta d\theta \\ &= 2\Gamma(p+q) \int_0^{\frac{\pi}{2}} \cos^{2p-1} \theta \sin^{2q-1} \theta d\theta. \end{aligned}$$

Définition 21. La fonction $B(p, q)$ (lue bêta) est définie par

la fonction $B(p, q)$

$$B(p, q) = \frac{\Gamma(p)\Gamma(q)}{\Gamma(p+q)} = 2 \int_0^{\frac{\pi}{2}} \cos^{2p-1} \theta \sin^{2q-1} \theta d\theta, \quad p, q > 0.$$

Une valeur de $B(p, q)$ est facile à calculer. Il s'agit de $B(\frac{1}{2}, \frac{1}{2})$:

$$B(\frac{1}{2}, \frac{1}{2})$$

$$B(\frac{1}{2}, \frac{1}{2}) = 2 \int_0^{\frac{\pi}{2}} d\theta = \pi$$

et donc

$$\frac{\Gamma(\frac{1}{2})\Gamma(\frac{1}{2})}{\Gamma(1)} = (\Gamma(\frac{1}{2}))^2 = \pi.$$

Puisque $\Gamma(z) > 0$ pour $z > 0$, on a donc :

$$\Gamma\left(\frac{1}{2}\right) = \sqrt{\pi}.$$

La forme de $B(p, q)$ est souvent présentée après le changement de variable $t = \cos^2 \theta$ avec $dt = -2 \cos \theta \sin \theta d\theta$. L'intervalle d'intégration $0 \leq \theta \leq \frac{\pi}{2}$ correspond donc à $1 \geq t \geq 0$. Alors

$$\begin{aligned} B(p, q) &= 2 \int_0^{\frac{\pi}{2}} \cos^{2p-2} \theta \sin^{2q-2} \theta \sin \theta \cos \theta d\theta \\ &= - \int_1^0 t^{p-1} (1-t)^{q-1} dt \\ &= \int_0^1 t^{p-1} (1-t)^{q-1} dt \end{aligned}$$

et donc, par exemple

$$B\left(\frac{1}{2}, \frac{1}{2}\right) = \int_0^1 \frac{dt}{\sqrt{t(1-t)}} = \pi.$$

formule
des compléments

Voici une formule reliant $\Gamma(z)$ et $\Gamma(1-z)$.

$$B(z, 1-z) = \Gamma(z)\Gamma(1-z) = \int_0^1 t^{z-1} (1-t)^{-z} dt = \frac{\pi}{\sin \pi z}.$$

La preuve de cette relation utilise une intégration dans le plan complexe et nous ne la ferons pas.

Exercices

1. Montrer que la fonction $\Gamma(z)$, $z > 0$ peut être réécrite sous la forme

$$\Gamma(z) = \int_0^1 \left(\ln \frac{1}{t} \right)^{z-1} dt.$$

2*. Réécrire l'expression suivante en termes de la fonction Γ :

$$\frac{(n+1)(n+2) \dots (n+2s)}{(2 \cdot 4 \cdot 6 \dots (2s-2)(2s)) ((2n+3)(2n+5) \dots (2n+2s+1))}$$

si $n, s \in \mathbb{N}$.

3*. (a) Montrer que, pour $z > 0$

$$\int_0^\infty e^{-xz} dx = \frac{1}{z} \Gamma\left(\frac{1}{z}\right).$$

(b) Est-ce que la limite

$$\lim_{z \rightarrow \infty} \int_0^{\infty} e^{-x^z} dx$$

existe ? Si oui, à quoi est-elle égale ?

4. Montrer pour $s \in \mathbb{N}$ et $a > 0$ que

(a)

$$\int_0^{\infty} x^{2s+1} e^{-ax^2} dx = \frac{s!}{2a^{s+1}}$$

(b)

$$\int_0^{\infty} x^{2s} e^{-ax^2} dx = \frac{(2s-1)!!}{2^{s+1} a^s} \sqrt{\frac{\pi}{a}}.$$

Note : Ces intégrales sont reliées aux moments de la distribution normale (la distribution gaussienne) en statistique. Elles jouent un rôle important en statistique et, en physique, en mécanique statistique.

5. Montrer la formule de duplication de Legendre

$$\Gamma(z)\Gamma(z - \frac{1}{2}) = 2^{-2z+2} \sqrt{\pi} \Gamma(2z - 1)$$

en intégrant la relation $(\sin 2\theta)^{2z+1} = (2 \sin \theta \cos \theta)^{2z+1}$.

formule
de duplication
de Legendre

6. Montrer que

$$\int_0^{\frac{\pi}{2}} \cos^{2s+1} \theta d\theta = \frac{(2s)!!}{(2s+1)!!}, \quad s \in \mathbb{N}.$$

7. Montrer que

$$\int_{-1}^1 (1-x^2)^{-\frac{1}{2}} x^{2n} dx = \begin{cases} \pi, & n = 0, \\ \pi \frac{(2n-1)!!}{(2n)!!}, & n \in \mathbb{N}. \end{cases}$$

8. Montrer que

$$\lim_{x \rightarrow 0} \frac{\Gamma(ax)}{\Gamma(x)} = \frac{1}{a}.$$

9*. Vérifier les propriétés suivantes de la fonction $B(a, b)$.

(a) $B(a, b) = B(a+1, b) + B(a, b+1)$

(b) $B(a, b) = \frac{a+b}{b} B(a, b+1)$

(c) $B(a, b) = \frac{b-1}{a} B(a+1, b-1)$

(d) $B(a, b)B(a+b, c) = B(b, c)B(a, b+c)$

10. La section précédente a permis de construire une fonction $\Gamma : \mathbb{R}^+ \rightarrow \mathbb{R}$ qui coïncide avec la fonction « ! » : $\mathbb{N} \rightarrow \mathbb{R}$ sur les entiers naturels. Existe-t-il d'autres fonctions $\gamma : \mathbb{R}^+ \rightarrow \mathbb{R}$ qui coïncident avec « ! » sur $\mathbb{N}$? Si oui, est-ce que l'exigence que γ soit continue implique que $\gamma = \Gamma$?

A.2 La formule de Stirling

Il est souvent utile de savoir comment se comporte une fonction lorsque son argument devient très grand. Par exemple il est utile de savoir que l'exponentielle tend vers l'infini plus vite que tout polynôme, c'est-à-dire

$$\lim_{x \rightarrow \infty} \frac{x^n}{e^x} = 0, \quad \text{pour tout } n \in \mathbb{N}.$$

comportement
asymptotique

La question qui nous occupera ici est d'identifier une fonction qui représente « bien » le comportement de $f(x)$ quand x est grand. Ce problème est celui du *comportement asymptotique*.¹ Cette question est donc différente de la question : quelle est la limite de $f(x)$ quand $x \rightarrow \infty$? Cette dernière question a comme réponse un chiffre si, évidemment, la limite existe.

Plusieurs symboles sont utilisés pour décrire les comportements asymptotiques.

Définition 22. Si $\lim_{x \rightarrow \infty} f(x)/\phi(x) = 1$, alors on écrit

$$f(x) \sim \phi(x), \quad x \rightarrow \infty.$$

Si $\lim_{x \rightarrow \infty} f(x)/\phi(x) = 0$, alors on écrit

$$f(x) = o(\phi(x)), \quad x \rightarrow \infty.$$

Si $|f(x)/\phi(x)|$ demeure borné sur un intervalle $[N, \infty)$, pour un certain $N \in \mathbb{N}$, alors on écrit

$$f(x) = \mathcal{O}(\phi(x)), \quad x \rightarrow \infty.$$

Les lettres « o » et « $\mathcal{O}$ » sont mises pour « de l'ordre de ». La forme la plus précise est la première qui donne exactement le comportement dominant quand $x \rightarrow \infty$. Puisque

$$\frac{1}{x+1} = \frac{1}{x} \frac{1}{1+\frac{1}{x}},$$

on peut écrire

$$\begin{aligned} \frac{1}{1+x} &\sim \frac{1}{x}, & x \rightarrow \infty \\ \frac{1}{1+x} &= o(1), & x \rightarrow \infty \\ \frac{1}{1+x} &= \mathcal{O}\left(\frac{1}{x}\right), & x \rightarrow \infty. \end{aligned}$$

Cependant il est également juste d'écrire

$$\begin{aligned} \frac{1}{1+x} &= o(\sqrt{x}), & x \rightarrow \infty \\ \frac{1}{1+x} &= \mathcal{O}\left(\frac{1}{\sqrt{x}}\right), & x \rightarrow \infty. \end{aligned}$$

1. Un des classiques sur le sujet est le livre

Olver, F.W.J., *Introduction to Asymptotics and Special Functions*, Academic (1974).

Après le présent cours, ce livre vous sera à toute fin pratique accessible.

La forme utilisant « $\sim$ » est donc plus précise. On utilise parfois les mêmes symboles $\sim$, o et $\mathcal{O}$ pour le comportement proche d'un point où f est singulière. Par exemple

$$\frac{1}{x + \sin x} \sim \frac{1}{2x}, \quad x \rightarrow 0.$$

Le but de cette section est de donner le comportement asymptotique de la fonction Γ . Ceci nous permettra entre autres de répondre à la question : de $n!$ et e^n , quelle est la fonction qui va le plus rapidement à l'infini ?

Nous avons vu sur le graphique que $\Gamma(x)$ augmente rapidement avec x . Son comportement asymptotique est donné par le théorème suivant.

Théorème 43 (Stirling). Pour x grand :

formule de Stirling

$$\Gamma(x+1) = \int_0^\infty e^{-t} t^x dt \sim x^x e^{-x} \sqrt{2\pi x}. \quad (\Gamma_2)$$

Cette relation est connue sous le nom de formule asymptotique de Stirling. La preuve qui est donnée ici est celle de Schwartz. Stirling, James (1692-1770)

Démonstration. Pour le montrer, nous allons faire deux changements de variable successifs, puis étudier l'intégrand résultant. Notons tout d'abord que $e^{-t} t^x$ a son maximum en $t = x$:

$$\frac{d}{dt} e^{-t} t^x = -e^{-t} t^x + x e^{-t} t^{x-1} = \left(-1 + \frac{x}{t}\right) e^{-t} t^x$$

qui s'annule en $t = x$. Faisons donc un premier changement de variable pour positionner ce maximum en 0 :

$$u = t - x \quad du = dt$$

et donc

$$\begin{aligned} \Gamma(x+1) &= \int_{-x}^\infty e^{-u} e^{-x} e^{x \ln(u+x)} du \\ &= \int_{-x}^\infty e^{-x} e^{-u} e^{x(\ln(1+\frac{u}{x})+\ln x)} du \\ &= x^x e^{-x} \int_{-x}^\infty e^{-u+x \ln(1+\frac{u}{x})} du \end{aligned}$$

Tant que u reste petit devant x , on a :

$$-u + x \ln\left(1 + \frac{u}{x}\right) \sim -u + x \left(\frac{u}{x} - \frac{u^2}{2x^2} + \mathcal{O}\left(\frac{|u|^3}{|x|^3}\right)\right) \sim -\frac{u^2}{2x} + \mathcal{O}\left(\frac{|u|^3}{|x|^2}\right).$$

Ce comportement suggère le changement de variable $u = \sqrt{x}v$, $du = \sqrt{x}dv$ et

$$\Gamma(x+1) = x^x e^{-x} \sqrt{x} \int_{-\sqrt{x}}^\infty e^{-\sqrt{x}v+x \ln(1+\frac{v}{\sqrt{x}})} dv.$$

Il reste donc à montrer que

$$I(x) = \int_{-\sqrt{x}}^{\infty} e^{-\sqrt{x}v + x \ln\left(1 + \frac{v}{\sqrt{x}}\right)} dv \xrightarrow{x \rightarrow \infty} \sqrt{2\pi}.$$

Ceci est la partie *difficile* de la démonstration car l'exposant de l'exponentielle contient deux termes dont les comportements s'opposent. Il y a cependant raison d'espérer. En effet, à la borne inférieure, c'est-à-dire pour $v = -\sqrt{x} + \epsilon$ avec $\epsilon > 0$, les deux termes de l'exposant deviennent

$$-\sqrt{x}v \sim x \quad \text{et} \quad x \ln\left(1 + \frac{v}{\sqrt{x}}\right) \sim x \ln \epsilon$$

quand $\epsilon \rightarrow 0$. Pour un x fixé, l'intégrand se comporte donc bien puisque le second terme $x \ln \epsilon \xrightarrow{\epsilon \rightarrow 0} -\infty$ et l'exponentielle s'annule. A la borne supérieure, le second terme $x \ln(1 + v/\sqrt{x})$ se comporte comme $x \ln v$ et le premier terme $-\sqrt{x}v$ sera donc dominant. Ainsi, pour un x fixé, $\lim_{v \rightarrow \infty} h(x, v) = 0$. Voyons donc comment évaluer cette intégrale ou, plutôt, sa limite.

Réécrivons cette intégrale

$$I(x) = \int_{-\infty}^{\infty} h(x, v) dv$$

où

$$h(x, v) = \begin{cases} \exp\left(-\sqrt{x}v + x \ln\left(1 + \frac{v}{\sqrt{x}}\right)\right), & v > -\sqrt{x} \\ 0, & v \leq -\sqrt{x}. \end{cases}$$

Remarquons tout d'abord que, pour v fixé, $h(x, v) \xrightarrow{x \rightarrow \infty} e^{-v^2/2}$; en effet :

$$\begin{aligned} \lim_{x \rightarrow \infty} h(x, v) &= \exp \left[\lim_{x \rightarrow \infty} -\sqrt{x}v + x \ln \left(1 + \frac{v}{\sqrt{x}} \right) \right] \\ &= \exp \left[\lim_{x \rightarrow \infty} -\sqrt{x}v + x \left(\frac{v}{\sqrt{x}} - \frac{v^2}{2x} + \mathcal{O} \left(\frac{|v|^3}{|x|^{3/2}} \right) \right) \right] \\ &= \exp \left[\lim_{x \rightarrow \infty} -\frac{v^2}{2} + \mathcal{O} \left(\frac{|v|^2}{|x|^{1/2}} \right) \right] \\ &\rightarrow \exp -\frac{v^2}{2}. \end{aligned}$$

Donc

$$\int_{-\infty}^{\infty} \lim_{x \rightarrow \infty} h(x, v) dv = \int_{-\infty}^{\infty} e^{-v^2/2} dv = \sqrt{2\pi}.$$

Peut-on dire que

$$\lim_{x \rightarrow \infty} \int_{-\infty}^{\infty} h(x, v) dv \stackrel{?}{=} \int_{-\infty}^{\infty} \lim_{x \rightarrow \infty} h(x, v) dv.$$

Si oui, la preuve est terminée. Nous devons utiliser le théorème de Lebesgue. (Ce théorème est probablement démontré dans un cours sur la théorie de la mesure ou sur l'intégrale de Lebesgue.)

Théorème 44 (Lebesgue). *Soit g une fonction intégrable sur $E \subset \mathbb{R}^n$ et soit $\{f_n\}$ une suite de fonctions intégrables telle que $|f_n| \leq g$ (presque partout) sur E et telle que $\lim_{n \rightarrow \infty} f_n(v) = f(v)$ pour (presque) tout $v \in E$. Alors*

$$\int f(v) dv = \lim_{n \rightarrow \infty} \int f_n(v) dv.$$

Même si nous ne démontrerons pas ce théorème, il est utile de donner un exemple de l'importance de l'existence de la fonction g qui plafonne les fonctions f_n . Utilisons les fonctions

$$f_n(x) = \begin{cases} n, & x \in [0, \frac{1}{n}], \\ 0, & x \in (\frac{1}{n}, 1]. \end{cases}$$

Remarquons que

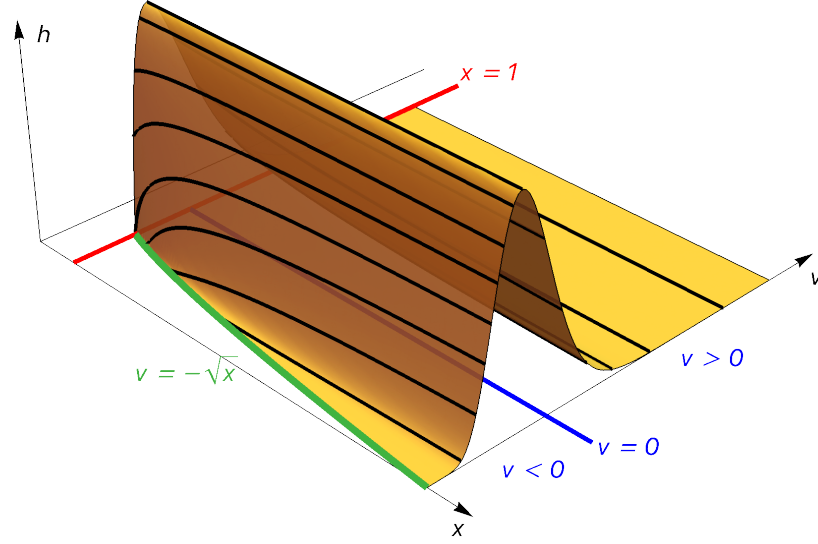
$$\int_0^1 f_n(x) dx = 1 \quad \text{pour tout } n$$

et donc $\lim \int f_n dx = 1$. Mais $\lim_{n \rightarrow \infty} f_n = 0$ sur $(0, 1]$ et n'est pas définie en $x = 0$. Elle n'est donc pas intégrable au sens de Riemann sur l'intervalle $[0, 1]$.

Dans notre cas, le rôle du n (dans le théorème) est joué par x . Pour utiliser le théorème, il faut donc trouver une fonction intégrable $g(v)$ qui majore $h(x, v)$ pour x grand. Avant de commencer cet exercice, voici le graphe de la fonction h sur le domaine où elle ne s'annule pas. Le domaine où h ne s'annule pas va, pour un x fixé, de $-\sqrt{x}$ à $+\infty$. La courbe en vert est donc la frontière inférieure de ce domaine. Puisque la propriété à démontrer vaut pour les x grands, le domaine de x à considérer est assez libre; la figure a été tracée pour les $x \geq 1$. La figure semble indiquer que le domaine de v est partitionné en trois régions selon que v est plus petit, égal ou plus grand que 0. La fonction h semble constante le long de $v = 0$. (Exercice : elle est égale à 1.) Lorsque $v < 0$, elle semble croître. Puisque nous venons de montrer que, pour v fixé, la fonction h va vers $e^{-v^2/2}$. S'il est possible de montrer qu'effectivement h croît en x pour $v < 0$, alors la fonction h sera bornée, dans la région $v < 0$, par $e^{-v^2/2}$ qui est intégrable. Finalement, il est plus difficile de deviner si, pour un $v > 0$ fixé, la fonction h est croissante ou non, mais le graphe indique au moins qu'elle semble inférieure à 1 dans la région $v > 0$.

En fait, nous allons montrer que $h(x, v)$ croît avec x pour $v < 0$ (et donc $|h(x, v)| < e^{-v^2/2}$ et $h(x, v)$ décroît avec x pour $v \geq 0$ (et donc $|h(x, v)| < |h(1, v)|$ pour $x \geq 1$). (Exercice : est-ce que $h(1, v)$ est intégrable sur $[0, \infty)$?) Puisque « exp » est croissante, il suffit de se concentrer sur $\ln h(x, v)$. Nous désirons donc montrer que

$$\begin{aligned} \ln h & \quad \text{croît avec } x & \quad \text{pour } v < 0 \\ \ln h & \quad \text{décroît avec } x & \quad \text{pour } v \geq 0, \end{aligned}$$

Figure A.2 – La fonction h

c'est-à-dire

$$\begin{aligned} \frac{\partial}{\partial \sqrt{x}} \ln h &\geq 0 && \text{pour } v < 0 \\ \frac{\partial}{\partial \sqrt{x}} \ln h &\leq 0 && \text{pour } v \geq 0. \end{aligned}$$

(La dérivée partielle par rapport à x a été remplacée par une dérivée partielle par rapport à $\sqrt{x}$ qui sera plus simple à effectuer ci-dessous. Puisque $\sqrt{x}$ croît lorsque x croît, les signes des deux dérivées sont les mêmes.) Mais $\ln h(x, 0) = 0 + x \ln(1 + 0) = 0$ et donc

$$\frac{\partial}{\partial \sqrt{x}} \ln h(x, 0) = 0.$$

Puisque la pente est nulle le long de $v = 0$, les énoncés sur le signe de $\partial \ln h / \partial \sqrt{x}$ suivront si nous montrons que cette pente augmente avec v , c'est-à-dire que la dérivée $\partial \ln h / \partial \sqrt{x}$ est plus grande que zéro lorsque $v < 0$ et plus petite lorsque $v \geq 0$. Il suffit donc de montrer que

$$\frac{\partial}{\partial v} \frac{\partial}{\partial \sqrt{x}} \ln h(x, v) < 0 \quad \text{pour } v > -\sqrt{x}.$$

Alors

$$\begin{aligned}
 \frac{\partial}{\partial \sqrt{x}} \frac{\partial}{\partial v} \left(-\sqrt{x}v + x \ln \left(1 + \frac{v}{\sqrt{x}} \right) \right) &= \frac{\partial}{\partial \sqrt{x}} \left(-\sqrt{x} + \frac{x}{1 + \frac{v}{\sqrt{x}}} \frac{1}{\sqrt{x}} \right) \\
 &= \frac{\partial}{\partial \sqrt{x}} \left(-\sqrt{x} + \frac{(\sqrt{x})^2}{v + \sqrt{x}} \right) \\
 &= -1 + \frac{2\sqrt{x}}{v + \sqrt{x}} - \frac{x}{(v + \sqrt{x})^2} \\
 &= \frac{-v^2 - 2v\sqrt{x} - x + 2v\sqrt{x} + 2x - x}{(v + \sqrt{x})^2} \\
 &= -\frac{v^2}{(v + \sqrt{x})^2} < 0
 \end{aligned}$$

pour tout $v > -\sqrt{x}$.

□

11. Montrer que

$$\lim_{n \rightarrow \infty} \frac{n}{(n!)^{\frac{1}{n}}} = e.$$

Exercices

12*. Donner le comportement asymptotique du nombre de Catalan $C_n = \frac{1}{n+1} \binom{2n}{n}$ lorsque $n \rightarrow \infty$.

13*. Quel est le comportement asymptotique des fonctions suivantes ? (Plus précisément, trouver $\phi(x)$ tel que $f(x) \sim \phi(x)$ quand $x \rightarrow \infty$.)

(a) $f(x) = \frac{\sin \frac{1}{x}}{x}$

(b) $f(x) = \cosh x$

(c) $f(x) = \tanh x - 1$

14. Soit $x = x(u)$ la fonction donnée implicitement par $x = u - \tanh x$. Montrer que

$$x = u - 1 + 2e^{-2u+2} + O(e^{-4u}).$$

15*. Calculer

$$\lim_{n \rightarrow \infty} \frac{(2n-1)!!}{(2n)!!} \sqrt{n}.$$

16. Montrer que

$$\lim_{n \rightarrow \infty} \frac{\int_0^{\frac{\pi}{2}} \sin^{2n} \theta \, d\theta}{\int_0^{\frac{\pi}{2}} \sin^{2n+1} \theta \, d\theta} = 1.$$

17. Une autre preuve de la formule de Stirling (selon Körner).

(a) Montrer que, si $x > -1$, il existe θ , $0 < \theta < 1$, tel que

$$\ln(1+x) = x - \frac{x^2}{2(1+\theta x)^2}.$$

(b) Montrer que, pour $t, a > 0$, l'équation

$$t^2 = x - a - a \ln \frac{x}{a}$$

possède deux solutions x_1 et x_2 . Nous choisirons $x_1 < x_2$.

(c) Utilisant (a) et (b), montrer qu'il existe $-1 < \theta_1 < 0 < \theta_2 < 1$ tel que

$$\frac{1}{t} \sqrt{\frac{a}{2}} = \frac{a}{a-x_1} + \theta_1 = \frac{a}{x_2-a} + \theta_2.$$

(d) En faisant le changement de variables $t^2 = x - a - a \ln \frac{x}{a}$, montrer que, pour $a > 0$:

$$\Gamma(a+1) = \int_0^\infty x^a e^{-x} dx = 2a^a e^{-a} \int_0^\infty e^{-t^2} \left(\frac{x_1}{a-x_1} + \frac{x_2}{x_2-a} \right) t dt$$

où $x_1 = x_1(t)$ et $x_2 = x_2(t)$ sont les racines identifiées en (b).

(e) Utilisant (c), montrer la formule de Stirling, c'est-à-dire que

$$\frac{\Gamma(a+1)e^a}{a^{a+\frac{1}{2}}} \longrightarrow \sqrt{2\pi}, \quad \text{lorsque } a \rightarrow \infty.$$

Bibliographie

Voici les principales sources pour ces notes de cours.

George Arfken, *Mathematical methods for physicists*, Academic Press (1985) San Diego.

Thomas William Körner, *Fourier analysis*, Cambridge University Press (1988) Cambridge. Son recueil *Exercises for Fourier analysis* (1993) qui le complète est aussi excellent.

Laurent Schwartz, *Méthodes mathématiques pour les sciences physiques*, Hermann (1997) Paris.

Georgi P. Tolstov, *Fourier series*, translated from russian by Richard A. Silverman, Dover Publications (1962) New York.

Enfin voici un autre manuel populaire qui est paru après que ces notes aient été rédigées (autour de 1995).

Elias S. Stein, Rami Sharkarshi, *Fourier analysis, an introduction*, Princeton University Press (2003) New Jersey.

Solutions et suggestions pour les exercices

Chapitre 1

Section 1.1

1.

$$A_{mn} = \begin{cases} 2, & \text{si } m = n = 0, \\ 1, & \text{si } m = n \neq 0, \\ 0, & \text{autrement,} \end{cases}$$

$$B_{mn} = 0,$$

$$C_{mn} = \begin{cases} 1, & \text{si } m = n \neq 0, \\ 0, & \text{autrement.} \end{cases}$$

2. (b) $a_n = 0$ et $b_n = 2(-)^{n+1}/n$.

(d) $b_3 = 1$ et tous les autres sont nuls.

(f) $a_n = \frac{4}{\pi(1-n^2)}$ si n est pair, $a_n = 0$ si n est impair et $b_n = 0$.

(h) $a_n = \frac{2(-)^n}{\pi} \frac{a \sinh a\pi}{n^2 + a^2}$ et $b_n = 0$.

3. (b) Non périodique si f est générique.

(d) Périodique de période 4π .

(f) Périodique de période 2π .

5. $a_m = 2$ si $m \leq n$. Tous les autres coefficients sont nuls.

Section 1.2

7. (b) Impaire.

(d) Paire.

(f) Paire.

8. (b) Les trois fonctions sont respectivement paire, impaire et paire.

9. (b) $a_0 = \pi$ et $a_m = -\frac{4}{\pi m^2}$ si m est impair. Les autres coefficients sont nuls.

11. Oui, mais pas beaucoup.

12. (a) Pourquoi la fonction $f(x) = |x| - \frac{1}{2}$ a-t-elle ses coefficients $a_{2n} = 0$?

Section 1.3

15. $b_m = \frac{4}{\pi m}$ si m est impair. Tous les autres coefficients sont nuls.

17. $a_0 = 1$ et $a_m = -\frac{4}{\pi^2 m^2}$ si m est impair. Les autres coefficients sont nuls.

18. (a)

$$b_m = \frac{2L^2 \sin\left(\frac{n\pi\theta}{L}\right)}{n^2\pi^2(L-\theta)\theta}.$$

(b) $b_n^g = (-)^{n+1} b_n^f$.

(c) En $\theta \simeq 0.0813L$.

Section 1.4

20. (d) Par induction.

21. (b)

$$c_n = \frac{(-)^n i n}{\pi(n^2 - a^2)} \sin \pi a.$$

22. (b) $c_n = c_{-n}$.
 (d) $c_m = 0$ si $m \neq 0 \pmod n$.

23. (a) f et h le sont.
 (b) Seule f l'est.

24. (a) Intégrer par parties à répétition tant que la dérivée existe.

(b) $c_n = -\frac{6(-)^ni}{n^3}$.

25. (a) $B_1(x) = x - \frac{1}{2}$ et $B_2(x) = x^2 - x + \frac{1}{6}$.

Section 1.5

26. (d) Expérimenter avec *Mathematica* ou *Maple* pour deviner le spectre; puis le démontrer.

29. (a)

$$u_n(x, t) = b_n \sin \frac{n\pi x}{L} e^{-\frac{\pi^2 n^2}{L^2} t}.$$

Chapitre 2

Section 2.1

1. (a) Continue.
 (b) Continue.
 (c) Continue.
3. (a) $a_n \xrightarrow{n \rightarrow \infty} 1$.
 (b) $a_n \xrightarrow{n \rightarrow \infty} 0$.
 (c) a_n tend vers 1 ou -1 selon que $|c| > |d|$ ou $|d| > |c|$. Restent les cas $|c| = |d|$.
4. (a) Converge vers e .
 (b) Ne converge pas.
7. (a) Non.
 (b) Oui.
 (c) Non.

Section 2.2

8. (b) Le noyau de Dirichlet ne vérifie pas les propriétés (i) ni (ii).

Section 2.3

14. (e) Voir l'exercice sur les polynômes de Bernoulli à la section 1.4.

15. (b) $\frac{1}{12}(x^3 - \pi^2 x)$.
 (c) $\Omega = \{(x, y) \in \mathbb{R}^2 \mid x = m\pi \text{ ou } y = n\pi \text{ pour } m, n \in \mathbb{Z}\}$.
 (d) $-\frac{\pi^2}{32}$.

Section 2.4

19. Pour montrer la propriété (iv), écrire la trace en termes d'éléments de matrice.

21. (d) $P_0(x) = 1$, $P_1(x) = x$, $P_2(x) = x^2 - \frac{1}{3}$ et $P_3(x) = x^3 - \frac{3}{5}x$.

22. (d) $2N + 1$.

23. (b) Les familles de fonctions a, b, c et d sont mutuellement orthogonales. Alors, quelque soit $x \in \{a, b, c, d\}$, on a $\epsilon_{x,k,l} = 1$ si k et l sont non-nuls, $\epsilon_{x,k,l} = 2$ si un des deux indices k et l est nul et $\epsilon_{x,k,l} = 4$ si k et l le sont.
 (e) $B_{mn} = C_{nm} = -\frac{2}{n}(-)^n \delta_{m0}$ et les A_{mn} et D_{mn} sont nuls.

24. (a)

$$z_{mn} = \sin mx \sin ny \times (A_{mn} \cos \lambda_{mn} t + B_{mn} \sin \lambda_{mn} t)$$

où les λ_{mn} sont donnés par le spectre (voir (b)).

- (b) $\{-\lambda_{mn}^2 = a\sqrt{m^2 + n^2}, m, n \in \mathbb{N}\}$.
 (d) $\{-\lambda_{mn}^2 = a\sqrt{m^2 + n^2}, m \neq n \text{ et } m, n \in \mathbb{N}\}$.

Section 2.5

26. (a) Non.
 (b) Oui.
 (c) Non.

27. (b) $\frac{\pi^4}{96}$
 (d) $\frac{\pi^2}{16} - \frac{1}{2}$.

28. (c) $\phi_{00} = \frac{1}{2}$ et $\phi_{ij} = -2^{-\frac{3j}{2}-\frac{1}{2}}$ pour $i \geq 1$.

Chapitre 3**Section 3.1**

5. (b) $c = -\frac{1}{2}$ et $d = \frac{3}{2}$.
 (d) $c = 0$ et $d = 1$.

Section 3.2

7. (b) Pour le système parabolique, les fonctions d'une variable $A = A(\xi)$ et $B = B(\eta)$ vérifient

$$\begin{aligned} A'' &= (k^2 \xi^2 + \lambda)A \\ B'' &= (k^2 \eta^2 - \lambda)B \end{aligned}$$

où k est la constante apparaissant dans l'équation de Helmholtz et λ la constante de séparation.

Section 3.3

10. Non. Pourquoi ?

11. (b) Linéaire.
 (d) Non-linéaire.
 (f) Linéaire.

12. (a) Les solutions suivantes sont pour l'équation de Legendre associée : $p(x) = x^2 - 1$, $s(x) = \frac{m^2}{x^2 - 1}$ et $r(x) = 1$.
 (b) $x \in [-1, 1]$.

(c) $(f, g) = \int_{-1}^1 f(x)g(x) dx$ si f et g sont des fonctions à valeurs réelles.

13. (a) $p(x) = e^{-x^2}$, $s(x) = 0$ et $r(x) = e^{-x^2}$.

14. (b) Utiliser un argument graphique.

17. Pour les polynômes de Legendre, on peut consulter la section suivante. Voici les solutions pour l'équation de Chebyshev.

(a) A une constante près, les trois premiers polynômes sont $p_0(x) = 1$, $p_1(x) = x$ et $p_2(x) = x^2 - \frac{1}{2}$.

(b)

$$(f, g) = \int_{-1}^1 \frac{f(x)g(x)}{\sqrt{1-x^2}} dx.$$

Voici celles pour l'équation de Laguerre :

(a) A une constante près, les trois premiers polynômes sont $p_0(x) = 1$, $p_1(x) = -x + (1 + \alpha)$ et $p_2(x) = x^2 - 2(\alpha + 2)x + (\alpha + 2)(\alpha + 1)$.

(b)

$$(f, g) = \int_0^\infty x^\alpha e^{-x} f(x)g(x) dx.$$

Section 3.4

18. (c)

$$S(x, z) = \frac{z \sin x}{1 - 2z \cos x + z^2}.$$

19. (b) $\sum_{m=0}^\infty \sum_{n=m}^\infty B_{nm}$.
 (d) $\sum_{l=0}^\infty \sum_{m=0}^\infty D_{l+m,m}$.

20. (a) Si $f(x) = \sum_{l=0}^\infty a_l P_l(x)$, alors

$$\int_{-1}^1 |f(x)|^2 dx = \sum_{l=0}^\infty \frac{2a_l^2}{2l+1}.$$

24. (a) $H_0(x) = 1$, $H_1(x) = 2x$ et $H_2(x) = 4x^2 - 2$.
(f)

$$(f, g) = \int_{-\infty}^{\infty} e^{-x^2} f(x) g(x) dx.$$

25. (b) $L_0(x) = 1$, $L_1(x) = 1 - x$, $L_2(x) = 1 - 2x + \frac{1}{2}x^2$ et $L_3(x) = 1 - 3x + \frac{3}{2}x^2 - \frac{1}{6}x^3$.

Section 3.5

26. Les deux solutions sont

$$a_0 \sum_{n=0}^{\infty} \frac{(-)^n m^{2n}}{(2n)!}$$

et

$$a_1 \sum_{n=0}^{\infty} \frac{(-)^n m^{2n+1}}{(2n+1)!}.$$

27. (a) Utiliser la représentation intégrale.
(b) Utiliser les relations de récurrence.
(c) Construire les premiers $J_{n+\frac{1}{2}}$ et étudier les termes dominants quand $x \rightarrow \infty$.

29. Pour $m \in \mathbb{N} \cup \{0\}$ et

$$f(x) = \sum_{n=1}^{\infty} a_n J_m(\alpha_{mn} x),$$

l'identité de Parseval est :

$$\int_0^1 |f(x)|^2 x dx = \frac{1}{2} \sum_{n=1}^{\infty} a_n^2 (J_{m+1}(\alpha_{mn}))^2.$$

31. Il est possible d'obtenir le résultat en remplaçant J_0 par sa série de Taylor

et en intégrant terme à terme. (Est-ce rigoureux ?)

35. Cet exercice utilise la fonction génératrice des fonctions de Bessel.

36. $(a^2 + b^2)^{-\frac{1}{2}}$. (Il existe plusieurs façons de faire cet exercice. L'intégration terme à terme en est une.)

37. (a) La solution générale est de la forme

$$u(t, \rho, \theta) = \sum_{m=0}^{\infty} \sum_{n=1}^{\infty} u_{mn}(t, \rho, \theta)$$

où

$$u_{mn}(t, \rho, \theta) = e^{-\alpha_{mn} t} J_m(\alpha_{mn} t) \times (a_{mn} \cos m\theta + b_{mn} \sin m\theta)$$

et α_{mn} est le n -ième zéro positif de J_m .

39. Le nombre β_{mn} est le n -ième zéro de J'_m . Les indices m , n et p sont des entiers vérifiant : $m \geq 0$, $n \geq 1$ et $p \geq 1$.

40. Application de la fonction génératrice.

Chapitre 4

Section 4.1

1. (b) $\frac{4}{\omega^3}(\sin \omega - \omega \cos \omega)$ si $\omega \neq 0$. Ne pas oublier d'obtenir la transformée en $\omega = 0$.

(d) $\frac{\sqrt{\pi}}{\sqrt{2}} \left(\cos \frac{\omega^2}{4} + \sin \frac{\omega^2}{4} \right)$.

(f) $\frac{1}{1+i\omega}$.

(h) La réponse se trouve au chapitre précédent.

2. (b) La valeur absolue des fonctions des exercices (b), (d), (e) et (g) est intégrable. Celle de (a), (c) et (f) ne l'est pas. L'intégrabilité de (h) dépend de la valeur de ν .

4. (b)

$$a_n = \frac{1}{2^{\frac{n}{2}} \sqrt{n!} \pi^{\frac{1}{4}}}.$$

(c)

$$\left(x - \frac{d}{dx} \right) \psi_n = \frac{a_n}{a_{n+1}} \psi_{n+1}.$$

(d)

$$\hat{\psi}_n(\zeta) = \sqrt{2\pi}(-i)^n a_n e^{-\zeta^2/2} H_n(\zeta).$$

6. (c) $\simeq 1.179$.

Section 4.2

7. L'hypothèse (B).

9. (b) La paire appartient à $L^1 \cap C$.

(d) La paire appartient à $L^1 \cap C$.

(f) La paire n'appartient pas à $L^1 \cap C$.

(h) Oui si $\nu > \frac{1}{2}$, non si $\nu \leq -\frac{1}{2}$. (Je ne sais pas pour $-\frac{1}{2} < \nu \leq \frac{1}{2}$.)

Section 4.3

14. (a)

$$f * g(t) =$$

$$= \begin{cases} \frac{2e^{at}}{a} \sinh a, & t \leq -1, \\ \frac{2}{a}(1 - e^{-a} \cosh at), & -1 \leq t \leq 1, \\ \frac{2e^{-at}}{a} \sinh a, & 1 \leq t. \end{cases}$$

15. (a)

$$\hat{f}(\omega) = -\frac{2i\omega}{1 + \omega^2}.$$

$$(c) \int_{-\infty}^{\infty} |f(x)|^2 dx = \frac{1}{2\pi} \int_{-\infty}^{\infty} |\hat{f}(\omega)|^2 d\omega = 1$$

16. (b) Paire.

17. Peut être fait par calcul direct ou en utilisant le théorème de convolution.

19. (a)

$$f * f * f = \begin{cases} 0, & 3 \leq |t|, \\ \frac{1}{2}(|t| - 3)^2, & 1 \leq |t| \leq 3, \\ 3 - t^2, & 0 \leq |t| \leq 1. \end{cases}$$

$$f * f * f * f =$$

$$= \begin{cases} 0, & 4 \leq |t|, \\ \frac{1}{6}(|t| - 4)^3, & 2 \leq |t| \leq 4, \\ -\frac{1}{2}|t|^3 - 2t^2 + \frac{16}{3}, & 0 \leq |t| \leq 2. \end{cases}$$

(c) Pour $n = 1, 2, 3, 4$ les intégrales proposées sont respectivement : π , $\frac{2}{3}\pi$, $\frac{11}{20}\pi$ et $\frac{151}{315}\pi$.

Appendice A

Section A.1

2.

$$\frac{\Gamma(n+2s+1)\Gamma(2n+2)\Gamma(n+s+1)}{(\Gamma(n+1))^2\Gamma(2n+2s+2)\Gamma(s+1)}.$$

3. (b) 1.

9. Pour toutes ces propriétés, démarrer de l'expression

$$B(a, b) = \int_0^1 t^{a-1} (1-t)^{b-1} dt.$$

Section A.2

12. $\frac{2^{2n}}{\sqrt{\pi n^{\frac{3}{2}}}}.$

13. (b) $\frac{e^x}{2}.$

15. $\pi^{-\frac{1}{2}}.$ Se rappeler que

$$\lim_{n \rightarrow \infty} \left(1 \pm \frac{1}{n}\right)^n = e^{\pm 1}.$$

Index

- asymptotique, comportement $-$, 194
- Bernoulli, polynômes de $-$, 22
- Bessel
 - équation de $-$, 94
 - équation de $-$ sphérique, 96
 - fonction de $-$, 124–142
 - comportement
 - asymptotique, 134
 - fonction génératrice, 133
 - normalisation, 139
 - relation
 - de récurrence, 131–133
 - représentation intégrale, 129–131
 - inégalité de $-$, 76
- B (fonction Bêta), 191
- Cesàro, *voir* suite, série
- chaleur
 - équation de la $-$
 - solution par transformation de Fourier, 177–181
 - noyau de la $-$, 182
- Chebyshev
 - équation de $-$, 105, 106
- comportement asymptotique, 194
- conditions aux limites, 99
 - de Dirichlet, 100
 - de Neumann, 100
 - intermédiaires, 100
- continuité, 40
 - par morceaux, 60
 - sur le cercle, 20
- convergence
 - en moyenne, 73
 - ponctuelle, 58
- convolution, 169–174
 - de distributions normales, 175
 - linéarité de la $-$, 170
 - propriété fondamentale de la $-$, 171
- corde
 - pincée, 17
 - plombée, 24–30
 - vibrante, 30–35, 99
- Darwin, 183
- Dirichlet
 - conditions aux limites de $-$, 100
 - noyau de $-$, 55
 - théorème de $-$, 61, 83
- distance moyenne, 77
- domaine symétrique, 9
- équation différentielle
 - aux dérivées partielles, 86
 - homogène, 86
 - linéaire, 86
 - non-homogène, 86
 - ordinaire, 86
 - ordre d'une $-$, 86
 - solution générale d'une $-$, 87
 - solution particulière d'une $-$, 88
- espace vectoriel, 64
 - application linéaire, 64
 - base d'un $-$, 65
 - dimension d'un $-$, 65
 - isomorphisme d'un $-$, 64
 - produit scalaire pour un $-$, 67
- Euler, formule d'un $-$, 18

- Fejér
 noyau de $-$, 50–53
 théorème de $-$, 53
 théorème de $-$
 pour la transformation de Fourier, 157
- fonction
 Bêta (B), voir B
 continue, 40
 continue par morceaux, 60
 continue
 sur le cercle, 20
 convergente uniformément, 46
 en crêneaux, 4
 Gamma (Γ), voir Γ
 impaire, 9
 norme d'une $-$, 74
 périodique, 1
 paire, 9
 parité, 9
 prolongement d'une $-$
 impair, 10
 pair, 10
 régulière, 61
- Fourier, 1, 2
- fréquence
 naturelle, 28
 propre, 28
- Frobenius
 méthode de $-$, 125
- Fubini
 théorème de $-$, 160
- Γ (fonction Gamma), 187–189
 $-$ et la fonction factorielle!, 188
 formule de Stirling, 195, 200
 formule des compléments, 192
 intégrale de Gauss, 190
 pôles de $-$, 189
 relation fonctionnelle, 188
- gaussienne, voir normale
- Gibbs
 phénomène de $-$, 156
- Gibbs, phénomène de $-$, 6
- Heaviside
 fonction de $-$, 175
- Heisenberg
 inégalité de $-$, 167
- Helmholtz
 équation d' $-$, 92
- Hermite
 équation d' $-$, 104, 105
 polynômes d' $-$, 122, 155
- intégrale de Gauss, 190
- intégrale de Riemann, 44–45
- Kelvin, Lord $-$, 183
- Laguerre
 équation de $-$, 105
 polynômes de $-$, 123
- Legendre, 107
 équation de $-$, 105
 formule de
 duplication de $-$, 193
 polynômes de $-$, 106–120
 fonction génératrice, 112
 formule de Rodrigues, 115
 normalisation, 117
 orthogonalité, 115
 relation de
 récurrence, 112–114
- mode propre, 28
- Neumann
 conditions aux limites de $-$, 100
 fonction de $-$, 128
- nombre complexe, 18
 conjugué complexe, 18
 partie imaginaire, 18
 partie réelle, 18
 valeur absolue, 18
- normale
 distribution $-$, 190
 moments de la $-$, 193
- ondelettes de Haar, 56

- opérateur
 - auto-adjoint, 100
 - de Sturm-Liouville, *voir* Sturm-Liouville
- période, 1
- parité, *voir* fonction
- Parseval, identité de –, 78, 105
 - pour la transformation de Fourier, 164, 166
- partition d'un intervalle, 44
- Poisson
 - formule de somme de –, 167
- produit scalaire, 67
- prolongement, *voir* fonction
- séparation des variables, 31, 90–97
- série
 - convergente, 42
 - de Taylor, 2
 - géométrique, 43
 - sommable au sens de Cesàro, 49, 55
- série de Fourier
 - double –, 71
 - complexe, 18
 - convergence en moyenne, 73
 - convergence ponctuelle, 58–62
 - d'exponentielles, *voir* série de Fourier complexe
 - définition, 1–7
 - de période L , 15–17
- somme
 - inférieure, 44
 - supérieure, 44
- spectre, 28, 34
- Stirling, formule de –, 195, 200
- Sturm-Liouville
 - opérateur de –, 98
 - problème de –, 98–103
- suite, 41
 - convergente, 41
 - convergente au sens de Cesàro, 49
- sommable, 42
- système de coordonnées
 - cartésiennes, 97
 - cylindriques, 92
 - elliptiques, 97
 - paraboliques, 97
 - polaires, 97
 - sphériques, 94
- $\mathbb{T}^1$, 21
- tambour
 - carré, 72
 - rond, 135–142
- transformation de Fourier, 145–154
 - convolution, *voir* convolution de la distribution normale, 150, 174
 - identité de Parseval pour –, 164, 166
 - théorème d'inversion pour la –, 163
 - théorème de Fejér pour –, 157