

Polycopié de l'UV MT12
Techniques mathématiques de l'ingénieur

Ahmad El Hajj et Antoine Zurek

19 juin 2026

Table des matières

1	Rappels et mise au point	7
1.1	Intégrale de Riemann	7
1.1.1	Propriétés de l'intégrale de Riemann	9
1.1.2	Calcul intégral	12
1.2	Espace euclidien et espace hermitien	14
1.3	Compléments	21
1.3.1	Preuve du Théorème 2	21
1.3.2	Primitives usuelles	22
1.3.3	Continuité et continuité par morceaux	23
1.3.4	Dérivabilité et dérivabilité par morceaux	25
2	Séries de Fourier	29
2.1	Préambules	29
2.2	Séries de Fourier d'une fonction périodique	30
2.2.1	Polynômes trigonométriques	30
2.2.2	Fonctions périodiques	33
2.3	Convergence des séries de Fourier	43
2.3.1	Convergence ponctuelle des séries de Fourier	43
2.3.2	Suites et séries numériques et suites de fonctions	47
2.3.3	Convergence normale des séries de Fourier	51
2.4	Propriétés supplémentaires et exemples de calcul	53
2.5	Compléments	57
2.5.1	développement d'une fonction sur une base orthogonale	57
2.5.2	Quelques preuves des résultats de la Section 2.3.2	59
2.5.3	Preuve du théorème de Riemann-Lebesgue	60
3	Transformation de Fourier discrète	65
3.1	Introduction de la TFD	65
3.1.1	Calcul de l'erreur d'approximation	68
3.2	Quelques explications sur la FFT	70
3.3	Applications de la TFD/FFT	72
3.3.1	FFT, convolution discrète et débruitage d'un signal	72
3.3.2	FFT et compression d'image	76

4	Intégrale de Lebesgue	81
4.1	Un argument en défaveur de l'intégrale de Riemann	81
4.2	L'intégrale de Lebesgue	82
4.2.1	Le cas des fonctions positives	83
4.2.2	Fonctions Lebesgue intégrables	85
4.2.3	Les théorèmes utiles venant la théorie de l'intégration de Lebesgue .	88
4.3	Les espaces L^p	91
4.4	To be or not to be in L^p ($p = 1$ or 2)	94
4.4.1	Le cas I non borné	94
4.4.2	Le cas I borné	95
4.4.3	Résumé et exemples	96
4.5	Compléments	98
4.5.1	Quelques éléments de la construction de l'intégrale de Lebesgue . .	98
5	Transformée de Fourier	103
5.1	Transformée de Fourier dans L^1	103
5.1.1	Définition formelle de la transformée de Fourier dans L^1	104
5.1.2	Propriétés de la transformée de Fourier	105
5.1.3	Transformée de Fourier inverse	109
5.2	Transformée de Fourier dans L^2	111
5.3	Transformée de Fourier et convolution	112
5.4	Applications et illustrations	115
5.4.1	Résolution de l'équation de la chaleur sur un domaine non borné . .	115
5.4.2	Transformée de Fourier à fenêtre glissante (transformée de Gabor) .	117
5.5	Compléments	121
5.5.1	Transformées de Fourier usuelles	121
6	Introduction à l'échantillonnage	123
6.1	Impulsion de Dirac et peigne de Dirac	123
6.1.1	Impulsion de Dirac	123
6.1.2	Peigne de Dirac et principales propriétés	125
6.2	Echantillonnage d'un signal	126
6.2.1	Principal général	126
6.2.2	Echantillonnage d'un signal à bande limitée	127
7	Transformée de Laplace	131
7.1	Transformée de Laplace	131
7.1.1	Propriétés de la transformée de Laplace	133
7.1.2	Exemples de calculs	136
7.2	Résolution d'équations différentielles	138
7.3	Compléments	142
7.3.1	Tableau de certaines transformée de Laplace	142

Remerciements

Nous tenons à remercier Ghislaine Gayraud pour ses lectures attentives !

Chapitre 1

Rappels et mise au point

Résumé. L’objectif principal de ce premier chapitre est de rappeler (ou de présenter) des notions élémentaires utiles pour toute la suite du cours. En particulier, nous revenons sur l’intégrale de Riemann ainsi que sur la notion d’espace euclidien et hermitien.

Références. Concernant la Section 1.1, la plupart des résultats et preuves viennent des livres :

- L. Moonens. *Intégration, de Riemann à Kurzweil et Henstock*, Ellipses (2017), chapitre 2,
- F. Boschet. *Séries de fonctions ; Intégrale de Riemann*, Masson (1995), chapitre 2 et chapitre 3.

Par ailleurs, si vous ressentez le besoin de revenir sur la pratique du calcul intégral (en plus des TD de ce cours) je conseille la fiche d’exercices (corrigés) suivante

- Calculs d’intégrales disponible sur le site [exo7](#) dans la section première année.

Les définitions et résultats de la Section 1.2 viennent principalement des documents :

- M. El Amrani. *Analyse de Fourier dans les espaces fonctionnels*, Ellipses (2008), chapitre 1,
- Poly MT23, chapitre 5.

1.1 Intégrale de Riemann

Rappelons dans un premier temps les idées principales de la construction de Riemann pour définir l’intégrale d’une fonction $f : [a, b] \rightarrow \mathbb{R}$ avec $-\infty < a < b < +\infty$. Pour ce faire, on considère dans un premier temps un découpage uniforme de l’intervalle fermé et borné $[a, b]$ de la forme $a = x_0 < x_1 < \dots < x_m = b$ et des points $t^j \in [x_{j-1}, x_j]$, pour $1 \leq j \leq m$. On considère ensuite la somme

$$S_m = \sum_{j=1}^m f(t^j)(x_j - x_{j-1}),$$

qui représente la somme des aires des rectangles de base $[x_{j-1}, x_j]$ et de hauteur $|f(t^j)|$. On dit alors, en suivant Riemann, que f est intégrable sur $[a, b]$ si la somme S_m tend vers une valeur finie lorsque les découpages de $[a, b]$ sont de plus en plus « petits » (i.e. quand $m \rightarrow \infty$ alors les longueurs $x_j - x_{j-1} \rightarrow 0$).

Afin de formaliser cette approche et d'être en mesure de définir rigoureusement la notion de fonction Riemann intégrable (ou R-intégrable), on introduit dans la suite un peu de vocabulaire et on formalise l'idée de « découpage » d'un intervalle de $\mathbb{R}$.

Définition 1. On a les définitions suivantes :

- Un intervalle I de $\mathbb{R}$ est dit dégénéré si I est l'intervalle vide ou réduit à un point.
- Deux intervalles I et J non dégénérés sont presque disjoints si $I \cap J$ est dégénéré.
- Soit $I = [a, b]$ ($a < b$) un intervalle de $\mathbb{R}$ borné. On appelle partition de I une famille finie $\mathcal{P} = \{I^j : 1 \leq j \leq m\}$ d'intervalles fermés, bornés, non dégénérés et presque disjoints deux à deux (I^k et I^p sont presque disjoints pour tout $1 \leq k, p \leq m$ avec $k \neq p$) vérifiant $I = \cup_{j=1}^m I^j$. Le pas de la partition $\mathcal{P}$, noté $\|\mathcal{P}\|$ est défini par

$$\|\mathcal{P}\| = \max\{\ell(I^j) : 1 \leq j \leq m\},$$

où $\ell(I^j)$ désigne la longueur de l'intervalle I^j .

- Enfin, on dit qu'une famille $\Pi = \{(x^j, I^j) : 1 \leq j \leq m\}$ est une P -partition de I si $\mathcal{P} = \{I^j : 1 \leq j \leq m\}$ est une partition de I et pour tout $1 \leq j \leq m$ on a $x^j \in I^j$. On appelle pas de Π le pas de la partition sous-jacente, i.e., $\|\Pi\| = \|\mathcal{P}\|$.

On considère à présent une fonction $f : I \rightarrow \mathbb{R}$ (avec $I = [a, b]$ non dégénéré et borné) et $\Pi = \{(x^j, I^j) : 1 \leq j \leq m\}$ une P -partition de I . On appelle somme de Riemann associée à f , I et Π le nombre donné par

$$S(I, f, \Pi) = \sum_{j=1}^m f(x^j) \ell(I^j).$$

On est alors en mesure de définir proprement la notion d'intégrabilité au sens de Riemann pour une fonction à valeurs dans $\mathbb{R}$:

Définition 2. Une fonction $f : I \rightarrow \mathbb{R}$, avec $I = [a, b]$ non dégénéré et **borné**, est R-intégrable sur I si il existe $\alpha \in \mathbb{R}$ tel que pour tout $\varepsilon > 0$, il existe $\delta > 0$ vérifiant

$$|S(I, f, \Pi) - \alpha| \leq \varepsilon,$$

Pour Π une P -partition de I vérifiant $\|\Pi\| \leq \delta$.

Remarque 1. Soit $f : [a, b] \rightarrow \mathbb{C}$ où $[a, b]$ est un intervalle non dégénéré et borné de $\mathbb{R}$. On dit que f est R-intégrable sur $[a, b]$ si et seulement si la partie réelle et imaginaire de f sont R-intégrable sur $[a, b]$.

Si $f : I \rightarrow \mathbb{R}$ est R-intégrable sur I , on appelle alors R-intégrale de f l'unique¹ $\alpha \in \mathbb{R}$ vérifiant la définition précédente. On note alors ce réel par

$$\int_I f(x) dx \quad \text{ou} \quad \int_{\min I}^{\max I} f(x) dx.$$

On posera aussi par **convention** :

$$\int_{\max I}^{\min I} f(x) dx = - \int_I f(x) dx.$$

Avant d'étudier les premières propriétés de l'intégrale de Riemann, voici une condition nécessaire de R-intégrabilité.

Théorème 2. Soient $I = [a, b]$ ($a < b$) un intervalle fermé et borné de $\mathbb{R}$ et $f : I \rightarrow \mathbb{R}$ une fonction. Si f est R-intégrable sur I , alors f est bornée sur I .

Démonstration. Voir Section 1.3.1. □

Remarque 3. Il existe des fonctions bornées non R-intégrables ! Par exemple $f = \mathbf{1}_{\mathbb{Q} \cap [0,1]}$.

1.1.1 Propriétés de l'intégrale de Riemann

Proposition 4. Soit $I = [a, b]$ ($a < b$) un intervalle borné. On a les propriétés élémentaires suivantes :

1. La fonction constante $f : I \rightarrow \mathbb{R}$ avec $f(x) = c$ (pour $c \in \mathbb{R}$) est R-intégrable sur I et $\int_I f(x) dx = c(b - a)$.
2. Soient f et g deux fonctions R-intégrables sur I et $c \in \mathbb{R}$. Alors cf et $f + g$ sont R-intégrables sur I et on a

$$\int_I cf(x) dx = c \int_I f(x) dx \quad \text{et} \quad \int_I (f(x) + g(x)) dx = \int_I f(x) dx + \int_I g(x) dx.$$

1. Supposons par l'absurde l'existence de deux réels α et β avec $\alpha \neq \beta$ vérifiant la Définition 2. Fixons alors un $\varepsilon > 0$ arbitraire, d'après la définition il existe $\delta' > 0$ tel que

$$|S(I, f, \Pi) - \alpha| \leq \varepsilon/2,$$

pour toute P-partition Π vérifiant $\|\Pi\| \leq \delta'$. Il existe également $\delta'' > 0$ vérifiant

$$|S(I, f, \Pi) - \beta| \leq \varepsilon/2,$$

pour toute P-partition Π vérifiant $\|\Pi\| \leq \delta''$. Soit alors $\delta = \min(\delta', \delta'')$ et Π une P-partition de I vérifiant $\|\Pi\| \leq \delta$. On obtient alors

$$|\alpha - \beta| \leq |S(I, f, \Pi) - \alpha| + |S(I, f, \Pi) - \beta| \leq \varepsilon.$$

Ceci montre que $\alpha = \beta$ car ε est arbitraire.

3. Soient f et g deux fonctions R -intégrables sur I vérifiant $|f(x)| \leq g(x)$ pour tout $x \in I$. Alors on a

$$\left| \int_I f(x) dx \right| \leq \int_I g(x) dx.$$

4. Soit f une fonction R -intégrable sur I telle que $|f|$ est R -intégrable sur I . Alors on a

$$\left| \int_I f(x) dx \right| \leq \int_I |f(x)| dx.$$

5. Soit f une fonction R -intégrable sur I vérifiant $f(x) \geq 0$ pour tout $x \in I$. Alors on a $\int_I f(x) dx \geq 0$.
6. Soient $f, g : I \rightarrow \mathbb{R}$ deux fonction R -intégrables sur I . Si on a $f(x) \leq g(x)$ pour tout $x \in I$, alors on a :

$$\int_I f(x) dx \leq \int_I g(x) dx.$$

7. Soient $f : I \rightarrow \mathbb{R}$ et $\mathcal{P} = \{I^j : 1 \leq j \leq m\}$ une partition de I . Si f est R -intégrable sur I^j pour tout $j = 1, \dots, m$, alors f est R -intégrable sur I et on a

$$\int_I f(x) dx = \sum_{j=1}^m \int_{I^j} f(x) dx.$$

Remarque 5. Dans le point (4) de la Proposition 4, l'hypothèse $|f|$ est R -intégrable est en fait inutile. En effet, on peut démontrer² que si f est une fonction R -intégrable sur un intervalle fermé et borné I de $\mathbb{R}$ alors $|f|$ est également R -intégrable sur I .

Démonstration. On procède point par point.

1. Pour toute P -partition Π de I , on obtient facilement que $S(I, f, \Pi) = c(b - a)$. En effet, comme par hypothèse on a $f(x) = c$ pour tout $x \in [a, b]$, on obtient

$$S(I, f, \Pi) = \sum_{j=1}^m f(x^j) \ell(I^j) = c \sum_{j=1}^m \ell(I^j) = c(b - a).$$

Ainsi, on a trivialement que

$$|S(I, f, \Pi) - c(b - a)| = 0.$$

Donc f est R -intégrable sur I d'intégrale $c(b - a)$.

2. Pour ce faire il faut d'abord démontrer qu'une fonction est R -intégrable si et seulement si elle est bornée et continue presque partout (notion que nous étudierons plus tard au chapitre 4).

2. Montrons que cf est R-intégrable, d'intégrale $c \int_I f(x) dx$. On fixe $\varepsilon > 0$, alors comme f est R-intégrable sur I il existe $\delta > 0$ tel que

$$\left| S(I, f, \Pi) - \int_I f(x) dx \right| \leq \frac{\varepsilon}{|c| + 1},$$

pour Π une P-partition de I vérifiant $\|\Pi\| \leq \delta$. De plus comme $S(I, cf, \Pi) = cS(I, f, \Pi)$, on a

$$\left| S(I, cf, \Pi) - c \int_I f(x) dx \right| = |c| \left| S(I, f, \Pi) - \int_I f(x) dx \right| \leq \frac{|c|}{|c| + 1} \varepsilon \leq \varepsilon.$$

Donc cf est R-intégrable sur I , d'intégrale $c \int_I f(x) dx$. Prouvons à présent que $f + g$ est R-intégrable sur I si f et g sont R-intégrables sur I . Comme f est R-intégrable sur I alors pour tout $\varepsilon > 0$ il existe un $\delta' > 0$ tel que

$$\left| S(I, f, \Pi') - \int_I f(x) dx \right| \leq \frac{\varepsilon}{2},$$

où Π' est une P-partition de I vérifiant $\|\Pi'\| \leq \delta'$. De même comme g est R-intégrable sur I alors pour tout $\varepsilon > 0$ il existe un $\delta'' > 0$ tel que

$$\left| S(I, g, \Pi'') - \int_I g(x) dx \right| \leq \frac{\varepsilon}{2},$$

avec Π'' une P-partition de I vérifiant $\|\Pi''\| \leq \delta''$. Soit à présent Π une P-partition de I avec $\|\Pi\| \leq \delta = \min\{\delta', \delta''\}$. On remarque par linéarité de la somme que

$$\begin{aligned} S(I, f + g, \Pi) &= \sum_{j=1}^m (f(x^j) + g(x^j)) \ell(I^j) = \sum_{j=1}^m f(x^j) \ell(I^j) + \sum_{j=1}^m g(x^j) \ell(I^j) \\ &= S(I, f, \Pi) + S(I, g, \Pi). \end{aligned}$$

Ainsi on en déduit que

$$\begin{aligned} \left| S(I, f + g, \Pi) - \int_I f(x) dx - \int_I g(x) dx \right| \\ \leq \left| S(I, f, \Pi) - \int_I f(x) dx \right| + \left| S(I, g, \Pi) - \int_I g(x) dx \right| \leq \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon. \end{aligned}$$

On en conclut que $f + g$ est R-intégrable sur I d'intégrale $\int_I f(x) dx + \int_I g(x) dx$.

3. On fixe $\varepsilon > 0$, soit $\delta' > 0$ un réel associé à f et $\varepsilon/2$ par la définition de R-intégrabilité de f sur I . De même soit $\delta'' > 0$ un réel associé à g et $\varepsilon/2$ par la définition de R-intégrabilité de g sur I . Enfin notons $\delta = \min\{\delta', \delta''\}$ et Π une P-partition de I avec $\|\Pi\| = \delta$. En remarquant que

$$|S(I, f, \Pi)| = \left| \sum_{j=1}^m f(x^j) \ell(I^j) \right| \leq \sum_{j=1}^m |f(x^j)| \ell(I^j) \leq \sum_{j=1}^m g(x^j) \ell(I^j) = S(I, g, \Pi),$$

on en déduit alors que

$$\begin{aligned}
\left| \int_I f(x) dx \right| &\leq \left| \int_I f(x) dx - S(I, f, \Pi) \right| + |S(I, f, \Pi)| \\
&\leq \frac{\varepsilon}{2} + S(I, g, \Pi) \\
&\leq \frac{\varepsilon}{2} + \left| S(I, g, \Pi) - \int_I g(x) dx \right| + \int_I g(x) dx \\
&\leq \varepsilon + \int_I g(x) dx.
\end{aligned}$$

Cette inégalité étant vérifiée pour $\varepsilon > 0$ alors le point (3) est vérifié.

4. On applique le point (3) à f et $g = |f|$.
5. On applique le point (3) à 0 et $g = f$.
6. On applique le point (5) à $g - f$.
7. On admet ce point.

Ce qui conclut la preuve du résultat. $\square$

Proposition 6. Soient $I = [a, b]$ ($a < b$) un intervalle borné et $f : I \rightarrow \mathbb{R}$ une fonction. Si f est continue (resp. monotone, i.e., croissante ou décroissante) sur I alors f est R-intégrable sur I .

Démonstration. On admet ce résultat. $\square$

Corollaire 7. Soient $I = [a, b]$ ($a < b$) un intervalle borné de $\mathbb{R}$ et $f : I \rightarrow \mathbb{R}$ une fonction continue par morceaux. Alors f est R-intégrable sur I .

Remarque importante 1. Voir (ou revoir) la Définition 8 pour la notion de continuité par morceaux. Je renvoie également à la Définition 9 pour la notion de dérivabilité par morceaux. Ces deux notions seront très importantes dans la suite du cours.

Démonstration. C'est une conséquence immédiate de la définition 8 et du point (7) de la Proposition 4. $\square$

1.1.2 Calcul intégral

On donne dans cette section des outils pour calculer l'intégrale de Riemann d'une fonction au moins continue (éventuellement par morceaux). Soient $I = [a, b]$ ($a < b$) un intervalle borné de $\mathbb{R}$ et $f : I \rightarrow \mathbb{R}$ une fonction continue sur I donc R-intégrable sur I . De manière générale pour calculer $\int_I f(x) dx$, on cherche une primitive de f , i.e., une fonction F vérifiant $F'(x) = f(x)$ pour tout $x \in I$. Dans ce cas on a alors

$$\int_I f(x) dx = [F(x)]_a^b = F(b) - F(a).$$

Pour trouver une primitive de f on peut dans certains cas se ramener au Tableau 1.1 (Section 1.3.2) et dans d'autres il faut utiliser des techniques différentes. Commençons par rappeler dans un premier temps la formule d'intégration par parties :

Proposition 8 (formule intégration par parties). *Soient $I = [a, b]$ ($a < b$) un intervalle borné de $\mathbb{R}$ et $u, v : I \rightarrow \mathbb{R}$ deux fonctions C^1 (i.e. dérivables avec u' et v' continues). Alors on a*

$$\int_a^b u'(x) v(x) dx = [u(x) v(x)]_a^b - \int_a^b u(x) v'(x) dx.$$

Démonstration. On a pour tout $x \in I$

$$(uv)'(x) = u'(x)v(x) + u(x)v'(x).$$

En intégrant cette relation on obtient

$$[u(x)v(x)]_a^b = \int_a^b (uv)'(x) dx = \int_a^b u'(x)v(x) dx + \int_a^b u(x)v'(x) dx.$$

Ce qui conclut la preuve de la proposition □

Exemple 1. *On cherche à calculer l'intégrale de la fonction $f : [0, 1] \rightarrow \mathbb{R}$ donnée par $f(x) = x \exp(-x)$. On remarque dans un premier temps que f est continue sur $[0, 1]$ et donc R -intégrable sur cet intervalle. Par ailleurs il n'est pas évident de trouver une primitive de $f \ll$ de tête $\gg$. En posant $u(x) = \exp(-x)$ et $v(x) = x$, on obtient par intégration par parties (IPP) :*

$$\begin{aligned} \int_0^1 x \exp(-x) dx &= [-x \exp(-x)]_0^1 + \int_0^1 \exp(-x) dx \\ &= -\exp(-1) + [-\exp(-x)]_0^1 = 1 - 2\exp(-1). \end{aligned}$$

Exemple 2. *On cherche maintenant à déterminer l'intégrale de la fonction $f : [1, 2] \rightarrow \mathbb{R}$ donnée par $f(x) = x \ln(x)$. Comme dans l'exemple précédent f est R -intégrable et on pose $u(x) = x$ et $v(x) = \ln(x)$, la formule d'IPP donne alors*

$$\int_1^2 x \ln(x) = \left[\frac{x^2}{2} \ln(x) \right]_1^2 - \int_1^2 \frac{x}{2} dx = 2 \ln(2) - \left[\frac{x^2}{4} \right]_1^2 = 2 \ln(2) - 1 + \frac{1}{4} = 2 \ln(2) - \frac{3}{4}.$$

Rappelons à présent la formule de changement de variable :

Proposition 9 (formule de changement de variable). *Soient I et J deux intervalles fermés, bornés et non dégénérés et $\varphi : J \rightarrow I$ une fonction monotone C^1 sur J . Soit f une fonction continue par morceaux sur I alors on a*

$$\int_I f(x) dx = \int_J (f \circ \varphi)(y) |\varphi'(y)| dy.$$

Démonstration. Admise. □

Exemple 3. On veut calculer l'intégrale suivante

$$\int_0^1 \frac{\exp(x)}{\exp(2x) + 1} dx.$$

On considère alors les fonctions $f : x \mapsto 1/(x^2 + 1)$ et $\varphi : x \mapsto \exp(x)$. Alors φ est C^1 et croissante sur $[0, 1]$ et on a $\varphi([0, 1]) = [1, \exp(1)]$. De plus f est continue sur $[1, \exp(1)]$ et de plus, pour tout $x \in [0, 1]$, on a

$$\frac{\exp(x)}{\exp(2x) + 1} = (f \circ \varphi)(x) \varphi'(x).$$

Ainsi d'après la formule de changement de variable on a

$$\int_0^1 \frac{\exp(x)}{\exp(2x) + 1} dx = \int_1^{\exp(1)} \frac{dx}{x^2 + 1} = [\arctan(x)]_1^{\exp(1)} = \arctan(\exp(1)) - \frac{\pi}{4}.$$

1.2 Espace euclidien et espace hermitien

Soit $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$. Dans toute cette section on notera E un $\mathbb{K}$ espace vectoriel et je vais supposer les notions d'espace vectoriel, de base et de dimension connues.

Définition 3. On appelle produit scalaire dans un espace vectoriel **réel** E , une application de $E \times E$ dans $\mathbb{R}$ notée $(x, y) \mapsto \langle x, y \rangle$ vérifiant les propriétés suivantes :

1. *bilinéarité* : $\forall x_1, x_2, y_1, y_2 \in E$ et $\forall \alpha_1, \alpha_2 \in \mathbb{R}$

$$\begin{aligned} \langle \alpha_1 x_1 + \alpha_2 x_2, y_1 \rangle &= \alpha_1 \langle x_1, y_1 \rangle + \alpha_2 \langle x_2, y_1 \rangle, \\ \langle x_1, \alpha_1 y_1 + \alpha_2 y_2 \rangle &= \alpha_1 \langle x_1, y_1 \rangle + \alpha_2 \langle x_1, y_2 \rangle, \end{aligned}$$

2. *Symétrie* : $\forall x, y \in E$ on a $\langle x, y \rangle = \langle y, x \rangle$,

3. *Caractère définie positif* :

$$\forall x \in E, \quad \langle x, x \rangle \geq 0, \quad \text{et} \quad \langle x, x \rangle = 0 \Rightarrow x = 0.$$

Exemple 4. Voici deux exemples de produit scalaire :

— Dans $\mathbb{R}^n$ ($n \geq 1$), l'application $\langle \cdot, \cdot \rangle$ de $\mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ donnée par

$$\langle x, y \rangle = \sum_{i=1}^n x_i y_i = x^\top y, \quad \forall x, y \in \mathbb{R}^n,$$

est un produit scalaire (appelé produit scalaire canonique de $\mathbb{R}^n$).

- Dans $C^0([a, b])$, l'ensemble des fonctions continues sur l'intervalle fermé, borné et non dégénéré $[a, b]$, l'application $\langle \cdot, \cdot \rangle$ donnée par

$$\langle f, g \rangle = \int_a^b f(x) g(x) dx, \quad \forall f, g \in C^0([a, b]),$$

est un produit scalaire.

Définition 4. Un espace euclidien est un $\mathbb{R}$ espace vectoriel de dimension finie, muni d'un produit scalaire. Dans le cas où E est de dimension infinie E est dit préhilbertien.

Exemple 5. Voici quelques exemples :

- L'espace $\mathbb{R}^n$ est un espace euclidien pour le produit scalaire canonique, pour rappel ce produit scalaire est donné par la formule suivante :

$$\langle x, y \rangle = \sum_{i=1}^n x_i y_i, \quad \forall x, y \in \mathbb{R}^n,$$

où les x_i et y_i désignent les coordonnées des vecteurs x et y .

- L'espace $\mathbb{R}_n[X]$ désignant l'ensemble des polynômes à coefficients dans $\mathbb{R}$ et de degré inférieur ou égale à n est un espace euclidien pour le produit scalaire suivant

$$\langle P, Q \rangle = \int_0^1 P(x) Q(x) dx, \quad \forall P, Q \in \mathbb{R}_n[X].$$

- Enfin, l'espace $C^0([a, b])$ doté du produit scalaire

$$\langle f, g \rangle = \int_a^b f(x) g(x) dx, \quad \forall f, g \in C^0([a, b]),$$

est un espace préhilbertien ($C^0([a, b])$ n'est pas de dimension finie).

Proposition 10 (Inégalité de Cauchy-Schwarz). Soit E un $\mathbb{R}$ espace vectoriel doté d'un produit scalaire $\langle \cdot, \cdot \rangle$. Alors pour tout $x, y \in E$ on a

$$|\langle x, y \rangle| \leq \sqrt{\langle x, x \rangle} \sqrt{\langle y, y \rangle},$$

avec égalité si x et y sont colinéaires.

Démonstration. Soient x et $y \in E$. Pour tout $t \in \mathbb{R}$ on considère la quantité $\langle tx + y, tx + y \rangle$. Alors d'après le point (3) de la Définition 3 on a

$$\langle tx + y, tx + y \rangle \geq 0,$$

et en développant

$$\langle tx + y, tx + y \rangle = t^2 \langle x, x \rangle + t \langle x, y \rangle + t \langle y, x \rangle + \langle y, y \rangle$$

$$= t^2 \langle x, x \rangle + 2t \langle x, y \rangle + \langle y, y \rangle.$$

Il y a alors plusieurs cas à considérer. Si $\langle tx + y, tx + y \rangle = 0$ alors soit $x = y = 0$ (et dans ce cas les vecteurs sont colinéaires et il y a clairement égalité dans l'inégalité de Cauchy-Schwarz) soit il existe un réel $t \neq 0$ tel que $tx + y = 0$ (encore une fois x et y sont colinéaires). Autrement dit, $y = -tx$ et alors

$$\begin{aligned} |\langle x, y \rangle| &= |-t| \langle x, x \rangle = |t| \langle x, x \rangle, \\ \sqrt{\langle y, y \rangle} &= \sqrt{t^2 \langle x, x \rangle} = |t| \sqrt{\langle x, x \rangle}, \end{aligned}$$

et donc l'égalité est vérifiée. Si par contre $\langle tx + y, tx + y \rangle \neq 0$ alors la fonction polynômiale $t \in \mathbb{R} \mapsto t^2 \langle x, x \rangle + 2t \langle x, y \rangle + \langle y, y \rangle$ est strictement positive, ce qui implique que son discriminant est strictement négatif, i.e.,

$$4 \langle x, y \rangle^2 \leq 4 \langle x, x \rangle \langle y, y \rangle \Rightarrow |\langle x, y \rangle| \leq \sqrt{\langle x, x \rangle} \sqrt{\langle y, y \rangle}.$$

Ce qui termine la preuve de l'inégalité de Cauchy-Schwarz. $\square$

Définition 5. On appelle norme sur un $\mathbb{R}$ espace vectoriel E , une application de E dans $\mathbb{R}_+$ notée $x \mapsto \|x\|$ et possédant les propriétés suivantes :

1. $\|x\| = 0 \Rightarrow x = 0$,
2. $\|\alpha x\| = |\alpha| \|x\|, \forall \alpha \in \mathbb{R}$,
3. $\|x + y\| \leq \|x\| + \|y\|$ (inégalité triangulaire).

Si E admet une norme on dit que E est un espace vectoriel normé.

Exemple 6. Dans $E = \mathbb{R}^n$, on a la norme euclidienne $\|x\| = (\sum_{i=1}^n |x_i|^2)^{1/2}$ ou encore

$$\|x\| = \left(\sum_{i=1}^n x_i x_i \right)^{1/2} = \sqrt{\langle x, x \rangle},$$

où $\langle \cdot, \cdot \rangle$ désigne le produit scalaire canonique de $\mathbb{R}^n$. Montrons que $\|\cdot\|$ est bien une norme. Les points (1) et (2) sont clairs. Montrons que l'inégalité triangulaire est vérifiée. Pour tout x et $y \in E$, on remarque dans un premier temps les équivalences suivantes

$$\begin{aligned} \|x + y\| \leq \|x\| + \|y\| &\Leftrightarrow \|x + y\|^2 \leq (\|x\| + \|y\|)^2 \\ &\Leftrightarrow \|x + y\|^2 \leq \|x\|^2 + 2\|x\| \|y\| + \|y\|^2. \end{aligned}$$

Ainsi, démontrer l'inégalité triangulaire revient à démontrer que pour tout x et $y \in E$ on a :

$$\|x + y\|^2 \leq \|x\|^2 + 2\|x\| \|y\| + \|y\|^2.$$

Or, on remarque que

$$\|x + y\|^2 = \sum_{i=1}^n (x_i + y_i)^2 = \sum_{i=1}^n (x_i^2 + 2x_i y_i + y_i^2) = \sum_{i=1}^n x_i^2 + 2 \sum_{i=1}^n x_i y_i + \sum_{i=1}^n y_i^2$$

$$\begin{aligned}
&= ||x||^2 + 2\langle x, y \rangle + ||y||^2 \\
&\leq ||x||^2 + 2|\langle x, y \rangle| + ||y||^2.
\end{aligned}$$

À présent, on applique l'inégalité de Cauchy-Schwarz au deuxième terme du membre de droite et on obtient

$$||x + y||^2 \leq ||x||^2 + 2\sqrt{\langle x, x \rangle} \sqrt{\langle y, y \rangle} + ||y||^2 \leq ||x||^2 + 2||x|| ||y|| + ||y||^2.$$

On en déduit que l'inégalité triangulaire est satisfaite.

Corollaire 11. Soit $(E, \langle \cdot, \cdot \rangle)$ un espace euclidien (ou préhilbertien) alors E est normé, la norme de E est donnée par $||x|| = \langle x, x \rangle^{1/2}$.

Remarque 12. La notion de norme et le Corollaire ci-dessus permettent de réécrire l'inégalité de Cauchy-Schwarz sous la forme :

$$|\langle x, y \rangle| \leq ||x|| ||y||, \quad \forall x, y \in E.$$

Preuve du Corollaire 11. On vérifie les trois points de la Définition 5. Si $||x|| = 0$ alors $||x||^2 = 0$ et dans ce cas $\langle x, x \rangle = 0$ ce qui implique $x = 0$ car le produit scalaire est défini positif. Soient $x \in E$ et $\alpha \in \mathbb{R}$ alors

$$||\alpha x||^2 = \langle \alpha x, \alpha x \rangle = \alpha^2 \langle x, x \rangle = \alpha^2 ||x||^2.$$

Ce qui prouve le second point. Enfin pour le dernier point il est équivalent de démontrer que $||x + y||^2 \leq (||x|| + ||y||)^2$ pour tout $x, y \in E$ (il suffit en fait de reprendre la preuve de l'inégalité triangulaire de l'Exemple 6). On a alors

$$\begin{aligned}
||x + y||^2 - (||x|| + ||y||)^2 &= \langle x + y, x + y \rangle - ||x||^2 - 2||x|| ||y|| - ||y||^2 \\
&= ||x||^2 + \langle x, y \rangle + \langle y, x \rangle + ||y||^2 - ||x||^2 - 2||x|| ||y|| - ||y||^2 \\
&= 2\langle x, y \rangle - 2||x|| ||y||.
\end{aligned}$$

À présent si $\langle x, y \rangle \leq 0$ alors on a $||x + y||^2 - (||x|| + ||y||)^2 \leq 0$ ce qui conclut la preuve. Si par contre $\langle x, y \rangle > 0$ alors d'après l'inégalité de Cauchy-Schwarz on a

$$\langle x, y \rangle \leq ||x|| ||y||,$$

et donc dans ce cas on obtient encore une fois $||x + y||^2 - (||x|| + ||y||)^2 \leq 0$. $\square$

Exemple 7. Dans $\mathbb{R}^n$ la norme euclidienne est donnée par la racine carrée du produit scalaire canonique.

Il faut tout de même prendre garde. Tout produit scalaire permet de construire une norme mais pour une norme donnée il n'existe pas nécessairement un produit scalaire associé. On retiendra

produit scalaire $\Rightarrow$ norme,
norme $\nRightarrow$ produit scalaire.

Par exemple dans $\mathbb{R}^n$, l'application $\|x\|_1 = \sum_{i=1}^n |x_i|$ est une norme mais ne découle pas d'un produit scalaire (voir la fiche de TD pour un résultat à ce propos).

Définition 6. Soit $(E, \langle \cdot, \cdot \rangle)$ un espace euclidien.

1. Deux vecteurs x et y sont orthogonaux si $\langle x, y \rangle = 0$.
2. Une famille de vecteurs $(x_1, \dots, x_n)$ est orthogonale si $\langle x_i, x_j \rangle = 0$ pour tout $i \neq j$.
3. une famille de vecteurs $(x_1, \dots, x_n)$ est orthonormée ou orthonormale si elle est orthogonale et $\|x_i\| = 1$ pour tout $i \in \{1, \dots, n\}$.

Exemple 8. Dans $\mathbb{R}^2$ muni du produit scalaire canonique les vecteurs $x = (1, 0)^\top$ et $y = (0, 1)^\top$ forment une famille orthonormée.

Proposition 13. Une famille orthogonale de vecteurs non nuls d'un espace euclidien E est une famille libre.

Démonstration. On note $n = \dim(E)$ et $(e_1, \dots, e_k)$ une famille orthogonale de vecteurs de E (avec $1 \leq k \leq n$). Soient $\alpha_1, \dots, \alpha_k \in \mathbb{R}$ vérifiant

$$\alpha_1 e_1 + \dots + \alpha_k e_k = 0.$$

En utilisant le produit scalaire de E on a

$$0 = \langle 0, e_1 \rangle = \langle \alpha_1 e_1 + \dots + \alpha_k e_k, e_1 \rangle = \alpha_1 \underbrace{\langle e_1, e_1 \rangle}_{\neq 0} + \alpha_2 \underbrace{\langle e_2, e_1 \rangle}_{=0} + \dots + \alpha_k \underbrace{\langle e_k, e_1 \rangle}_{=0}.$$

On en déduit que $\alpha_1 = 0$. En itérant ce processus on en déduit alors que $\alpha_1 = \dots = \alpha_k = 0$ et donc que la famille $(e_1, \dots, e_k)$ est libre. $\square$

Remarque 14. Soit $(E, \langle \cdot, \cdot \rangle)$ un espace euclidien doté d'une base $(e_1, \dots, e_n)$ orthogonale. Dans ce cas pour un vecteur quelconque $x \in E$ il existe $x_1, \dots, x_n \in \mathbb{R}$ (ses coordonnées) vérifiant

$$x = \sum_{i=1}^n x_i e_i.$$

On peut ensuite exprimer les coordonnées de x via le produit scalaire. En effet on a

$$\langle x, e_1 \rangle = \sum_{i=1}^n x_i \langle e_i, e_1 \rangle = x_1 \|e_1\|^2 + \sum_{i=1}^n x_i \underbrace{\langle e_i, e_1 \rangle}_{=0} = x_1 \|e_1\|^2.$$

De manière générale on a

$$x_i = \frac{\langle x, e_i \rangle}{\|e_i\|^2}, \quad \forall 1 \leq i \leq n.$$

On en déduit que pour déterminer les coordonnées de x on projette x sur les vecteurs de la base de E .

Terminons ce chapitre par l'introduction des espaces hermitiens. En effet, remarquons que la définition du produit scalaire, Définition 3, impose que E soit un $\mathbb{R}$ espace vectoriel. Comment adapter cette définition si E est un $\mathbb{C}$ espace vectoriel ? On a la définition suivante :

Définition 7. Soit E un $\mathbb{K}$ espace vectoriel (avec $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$).

- Une application $\langle \cdot, \cdot \rangle : E \times E \rightarrow \mathbb{K}$ est dite *sesquilinéaire* si elle est linéaire par rapport à sa première variable et anti-linéaire par rapport à la deuxième. Autrement dit, pour tout x, y et $z \in E$ et $\alpha \in \mathbb{K}$ on a

$$\begin{aligned}\langle x + y, z \rangle &= \langle x, z \rangle + \langle y, z \rangle, & \langle \alpha x, y \rangle &= \alpha \langle x, y \rangle, \\ \langle x, y + z \rangle &= \langle x, y \rangle + \langle x, z \rangle, & \langle x, \alpha y \rangle &= \overline{\alpha} \langle x, y \rangle.\end{aligned}$$

- Une forme sesquilinéaire $\langle \cdot, \cdot \rangle$ sur $E \times E$ est dite *hermitienne* si

$$\langle x, y \rangle = \overline{\langle y, x \rangle}, \quad \forall x, y \in E.$$

- Une forme sesquilinéaire hermitienne $\langle \cdot, \cdot \rangle$ sur $E \times E$ est dite *positive* si

$$\langle x, x \rangle \geq 0, \quad \forall x \in E.$$

- Une forme sesquilinéaire hermitienne $\langle \cdot, \cdot \rangle$ sur $E \times E$ est dite *définie positive* si elle est positive et si

$$\langle x, x \rangle = 0 \Rightarrow x = 0_E.$$

- On dit que $\langle \cdot, \cdot \rangle$ est un *produit scalaire hermitien* sur E si cette application est une forme sesquilinéaire hermitienne définie positive.
- Enfin on appelle E *espace hermitien* si il est de dimension finie et doté d'un produit scalaire hermitien.

Bien entendu pour E un espace hermitien on étend facilement la notion de norme en définissant

$$\|x\| = \langle x, x \rangle^{1/2}, \quad \forall x \in E.$$

Il convient de remarquer que pour un produit scalaire hermitien comme

$$\langle x, x \rangle = \overline{\langle x, x \rangle}, \quad \forall x \in E,$$

alors $\langle x, x \rangle \in \mathbb{R}$. On peut également étendre dans le cas d'un espace hermitien l'inégalité de Cauchy-Schwarz (la valeur absolue est alors remplacée par le module) ainsi que la notion de famille orthogonale et orthonormée.

Exemple 9. Soit $\mathcal{T}_N$ l'espace des polynômes trigonométrique de degré au plus N ($N \geq 1$), i.e., l'espace composé des polynômes de la forme

$$p(t) = \sum_{n=-N}^N c_n e_n(t) = \sum_{n=-N}^N c_n \exp(2i\pi n t) = \sum_{n=-N}^N c_n (\cos(2\pi n t) + i \sin(2\pi n t)), \quad (1.1)$$

où $c_n \in \mathbb{C}$. L'espace $\mathcal{T}_N$ est un $\mathbb{C}$ espace vectoriel de dimension finie. Pour le voir on note $\mathcal{B}_N = (e_{-N}, e_{-(N-1)}, \dots, e_0, \dots, e_{N-1}, e_N)$. Alors d'après (1.1)

$$\mathcal{T}_N = \text{Vect}\langle \mathcal{B}_N \rangle.$$

La famille $\mathcal{B}_N$ est en fait une base de l'espace vectoriel $\mathcal{T}_N$. En effet on définit sur l'espace $\mathcal{T}_N$ le produit scalaire hermitien (à vérifier) suivant :

$$\langle p, q \rangle = \int_0^1 p(t) \bar{q}(t) \, dt, \quad \forall p, q \in \mathcal{T}_N.$$

Ainsi $\mathcal{T}_N$ est un espace hermitien. De plus, comme dans le cas des espaces euclidiens, on peut munir $\mathcal{T}_N$ d'une norme via la formule

$$\|p\| = \sqrt{\langle p, p \rangle} = \left(\int_0^1 |p(t)|^2 \, dt \right)^{1/2}, \quad \forall p \in \mathcal{T}_N.$$

On remarque ensuite que pour tout n et $m \in \mathbb{Z}$ alors

$$\langle e_n, e_m \rangle = \int_0^1 e_n(t) \overline{e_m}(t) \, dt = \begin{cases} 0 & \text{si } n \neq m \\ 1 & \text{si } n = m. \end{cases}$$

En adaptant la preuve de la Proposition 13 au cas d'un espace hermitien, on déduit alors facilement des relations précédentes que la famille $\mathcal{B}_N$ est une famille libre contenant $2N+1$ éléments de $\mathcal{T}_N$. C'est donc une base de cet espace et $\dim(\mathcal{T}_N) = 2N+1$. Remarquons que les polynômes de $\mathcal{T}_N$ sont 1-périodiques ($p(t+1) = p(t)$ pour tout t), on reverra l'espace $\mathcal{T}_N$ au chapitre 2...

1.3 Compléments

1.3.1 Preuve du Théorème 2

Si f est R-intégrable sur I alors il existe $\delta > 0$ associé à $\varepsilon = 1$ et $\alpha = \int_I f(x) dx$ d'après la Définition 2. On fixe un $m \in \mathbb{N}^*$ vérifiant $\ell(I) \leq m\delta$ et on considère pour tout $1 \leq j \leq m$ les intervalles

$$I^j = \left[\min(I) + \frac{j-1}{m}\ell(I), \min(I) + \frac{j}{m}\ell(I) \right],$$

et $x^j = \min(I^j)$. Alors $\{I^j : 1 \leq j \leq m\}$ est une partition de I . De plus la famille

$$\Pi = \{(x^j, I^j) : 1 \leq j \leq m\},$$

est une P-partition de I de pas $\|\Pi\| \leq \delta$. En effet, on remarque que

$$\ell(I^j) = \min(I) + \frac{j}{m}\ell(I) - \left(\min(I) + \frac{j-1}{m}\ell(I) \right) = \frac{\ell(I)}{m} \leq \delta.$$

Comme $\|\Pi\| \leq \delta$ alors $|S(I, f, \Pi) - \int_I f(x) dx| \leq 1$. On pose à présent $M = \max_{1 \leq j \leq m} |f(x^j)|$. On remarque que M est nécessairement fini car

$$\left| S(I, f, \Pi) - \int_I f(x) dx \right| \leq 1 \Rightarrow \left| \sum_{j=1}^m f(x^j) \underbrace{\ell(I^j)}_{< \infty} - \underbrace{\int_I f(x) dx}_{< \infty} \right| \leq 1.$$

Montrons que f est bornée par une constante dépendante de M , m et $\ell(I)$. Pour ce faire on fixe $x \in I$. Comme la famille $\{I^j : 1 \leq j \leq m\}$ est une partition de I alors il existe un indice l tel que $x \in I^l$. On pose alors

$$\Pi_x = \{(x^1, I^1), \dots, (x^{l-1}, I^{l-1}), (x, I^l), (x^{l+1}, I^{l+1}), \dots, (x^m, I^m)\}.$$

La famille Π_x forme une P-partition de I avec $\|\Pi_x\| \leq \delta$, ce qui implique que

$$\left| S(I, f, \Pi_x) - \int_I f(x) dx \right| \leq 1.$$

Remarquons à présent que

$$(f(x) - f(x^l))\ell(I^l) = (f(x) - f(x^l))\ell(I^l) + \sum_{\substack{j=1 \\ j \neq l}}^m f(x^j)\ell(I^j) - \sum_{\substack{j=1 \\ j \neq l}}^m f(x^j)\ell(I^j). \quad (1.2)$$

On a alors

$$f(x)\ell(I^l) = [f(x) - f(x^l)]\ell(I^l) + f(x^l)\ell(I^l) \stackrel{(1.2)}{=} S(I, f, \Pi_x) - S(I, f, \Pi) + f(x^l)\ell(I^l).$$

Par conséquent on obtient

$$\begin{aligned} |f(x)\ell(I^l)| &\leq \left| S(I, f, \Pi_x) - \int_I f(x)dx \right| + \left| \int_I f(x)dx - S(I, f, \Pi) \right| + |f(x^l)| \frac{\ell(I)}{m} \\ &\leq 2 + M \frac{\ell(I)}{m}. \end{aligned}$$

Enfin comme $\ell(I^l) = \ell(I)/m$ on en déduit que $|f(x)| \leq M + 2m/\ell(I)$. Or comme x est quelconque dans I on en conclut que f est bornée.

1.3.2 Primitives usuelles

Fonctions	Primitive
$x \mapsto x^\alpha \quad \alpha \in \mathbb{R} \setminus \{-1\}, x \in \mathbb{R}_+^*$	$x \mapsto \frac{x^{\alpha+1}}{\alpha+1}$
$x \mapsto \frac{1}{x} \quad x \in \mathbb{R}_+^* \text{ ou } x \in \mathbb{R}_-^*$	$x \mapsto \ln(x)$
$x \mapsto \exp(x)$	$x \mapsto \exp(x)$
$x \mapsto \cos(x)$	$x \mapsto \sin(x)$
$x \mapsto \sin(x)$	$x \mapsto -\cos(x)$
$x \mapsto 1 + \tan^2(x) = \frac{1}{\cos^2(x)} \quad x \in]-\pi/2, \pi/2[$	$x \mapsto \tan(x)$
$x \mapsto \operatorname{ch}(x)$	$x \mapsto \operatorname{sh}(x)$
$x \mapsto \operatorname{sh}(x)$	$x \mapsto \operatorname{ch}(x)$
$x \mapsto 1 - \operatorname{th}^2(x) = \frac{1}{\operatorname{ch}^2(x)}$	$x \mapsto \operatorname{th}(x)$
$x \mapsto \frac{1}{\sqrt{1-x^2}} \quad x \in]-1, 1[$	$x \mapsto \arcsin(x)$
$x \mapsto -\frac{1}{\sqrt{1-x^2}} \quad x \in]-1, 1[$	$x \mapsto \arccos(x)$
$x \mapsto \frac{1}{1+x^2}$	$x \mapsto \arctan(x)$

TABLE 1.1 – Primitives usuelles.

1.3.3 Continuité et continuité par morceaux

Dans la suite je considère une fonction f définie sur $\mathbb{R}$, i.e., on a

$$f : \mathbb{R} \rightarrow \mathbb{R}.$$

Pour tout point $x_0 \in \mathbb{R}$ on appelle limite à droite en x_0 de f la quantité

$$\lim_{\substack{x \rightarrow x_0 \\ x > x_0}} f(x).$$

De même on appelle limite à gauche en x_0 de f la quantité

$$\lim_{\substack{x \rightarrow x_0 \\ x < x_0}} f(x).$$

Si ces limites existent on note

$$f(x_0^+) = \lim_{\substack{x \rightarrow x_0 \\ x > x_0}} f(x) \quad \text{et} \quad f(x_0^-) = \lim_{\substack{x \rightarrow x_0 \\ x < x_0}} f(x).$$

Exemple 10. — Soit f la fonction définie sur $\mathbb{R}$ avec $f(x) = 1$ pour tout $x \in \mathbb{R}$. Alors la limite à droite et à gauche de f en tout point $x_0 \in \mathbb{R}$ est égale à 1.

— Soit f la fonction définie sur $\mathbb{R}$ avec

$$f(x) = \begin{cases} 1 & \text{si } x > 0, \\ -1 & \text{si } x \leq 0. \end{cases}$$

La limite à droite de f en 0 vaut 1 et la limite à gauche de f en 0 vaut -1

— Soit f la fonction définie sur $\mathbb{R}$ avec

$$f(x) = \begin{cases} x & \text{si } x > 1, \\ x - 1 & \text{si } x \leq 1. \end{cases}$$

On a

$$\lim_{\substack{x \rightarrow 1 \\ x > 1}} f(x) = \lim_{\substack{x \rightarrow 1 \\ x > 1}} x = 1,$$

et

$$\lim_{\substack{x \rightarrow 1 \\ x < 1}} f(x) = \lim_{\substack{x \rightarrow 1 \\ x < 1}} (x - 1) = 0.$$

Par le biais de la définition de limite à droite et à gauche d'une fonction $f : \mathbb{R} \rightarrow \mathbb{R}$, on en déduit que f est continue en un point $x_0 \in \mathbb{R}$ si

$$\lim_{\substack{x \rightarrow x_0 \\ x > x_0}} f(x) = \lim_{\substack{x \rightarrow x_0 \\ x < x_0}} f(x).$$

Si ces limites sont égales on a

$$f(x_0) = \lim_{\substack{x \rightarrow x_0 \\ x > x_0}} f(x) = \lim_{\substack{x \rightarrow x_0 \\ x < x_0}} f(x).$$

De plus, si une fonction est continue en tout point de $\mathbb{R}$ on dit que cette fonction est dans l'espace $C^0(\mathbb{R})$.

La première fonction de l'Exemple 10 est continue en tout point de $\mathbb{R}$. Par contre la deuxième fonction n'est pas continue en 0, elle est par contre continue sur $] -\infty, 0[\cup]0, +\infty[$. Idem la troisième fonction de l'Exemple 10 n'est pas continue en 1 car

$$\lim_{\substack{x \rightarrow 1 \\ x > 1}} f(x) = 1 \neq 0 = \lim_{\substack{x \rightarrow 1 \\ x < 1}} f(x).$$

Par contre elle est continue sur $] -\infty, 1[\cup]1, +\infty[$. On arrive ainsi la notion (plus faible que celle de continuité) de continuité par morceaux :

Définition 8. Soit $f : \mathbb{R} \rightarrow \mathbb{R}$ une fonction. On dit que f est continue par morceaux sur $\mathbb{R}$, si à chaque fois que f est restreinte à un intervalle non vide $[a, b]$ de $\mathbb{R}$ il existe une suite finie de points $(x_i)_{1 \leq i \leq m}$ de $[a, b]$ vérifiant

$$a = x_1 < x_2 < \dots < x_m = b,$$

tel que la restriction de f à chaque intervalle $]x_{i-1}, x_i[$ ($2 \leq i \leq m$) se prolonge en une fonction continue sur $[x_{i-1}, x_i]$. Une fonction continue par morceaux sur $\mathbb{R}$ est également appelée C^0 par morceaux (sur $\mathbb{R}$).

Si on considère la première fonction de l'Exemple 10 alors f est continue par morceaux sur $\mathbb{R}$ car elle est continue sur $\mathbb{R}$. Pour la deuxième fonction, si on restreint f à tout intervalle évitant le point 0 alors la fonction est continue sur cet intervalle. Par contre, si on considère par exemple f restreinte à l'intervalle $[-1, 1]$ alors on définit la suite de trois points $x_1 = -1$, $x_2 = 0$ et $x_3 = 1$. Dans ce cas la fonction f restreinte à $]x_1, x_2[$ et $]x_2, x_3[$ est continue et est prolongeable par continuité sur les intervalles fermés $[x_1, x_2]$ et $[x_2, x_3]$. En effet, on a les limites suivantes

$$\begin{aligned} \lim_{\substack{x \rightarrow x_1 \\ x > x_1}} f(x) &= \lim_{\substack{x \rightarrow -1 \\ x > -1}} -1 = -1, \\ \lim_{\substack{x \rightarrow x_2 \\ x < x_2}} f(x) &= \lim_{\substack{x \rightarrow 0 \\ x < 0}} -1 = -1, \\ \lim_{\substack{x \rightarrow x_2 \\ x > x_2}} f(x) &= \lim_{\substack{x \rightarrow 0 \\ x > 0}} 1 = 1, \\ \lim_{\substack{x \rightarrow x_3 \\ x > x_3}} f(x) &= \lim_{\substack{x \rightarrow 1 \\ x > 1}} 1 = 1. \end{aligned}$$

Ainsi cette fonction est continue par morceaux sur $\mathbb{R}$. On raisonne de la même manière pour la troisième fonction (la discontinuité est située dans ce cas en 1). Donnons à présent l'exemple d'une fonction qui n'est pas continue par morceaux sur $\mathbb{R}$

Exemple 11. Soit la fonction définie sur $\mathbb{R}$ par

$$f(x) = \begin{cases} \frac{1}{x} & \text{si } x > 0, \\ 0 & \text{si } x \leq 0. \end{cases}$$

Ici f n'est pas continue par morceaux sur $\mathbb{R}$. En effet, la restriction de f à chaque intervalle évitant 0 est continue par contre si (par exemple) on restreint f à l'intervalle $[-2, 1]$ et que l'on définit la suite finie de trois points $x_1 = -2$, $x_2 = 0$ et $x_3 = 1$ alors f est continue sur $]x_1, x_2[$ et $]x_2, x_3[$ par contre f est continue sur $[x_1, x_2]$ mais pas $[x_2, x_3]$. En effet, on a

$$\lim_{\substack{x \rightarrow x_2 \\ x > x_2}} f(x) = \lim_{\substack{x \rightarrow 0 \\ x > 0}} \frac{1}{x} = +\infty.$$

Ainsi f n'admet pas de prolongement continu sur $[0, 1]$. Donc f n'est pas continue par morceaux.

On peut ainsi retenir qu'une fonction f est continue par morceaux sur $\mathbb{R}$ si à chaque point x_0 où la fonction est discontinue les quantités $f(x_0^-)$ et $f(x_0^+)$ existent. En particulier une fonction continue sur $\mathbb{R}$ est également continue par morceaux sur $\mathbb{R}$.

1.3.4 Dérivabilité et dérivabilité par morceaux

Soit $f : \mathbb{R} \rightarrow \mathbb{R}$ une fonction. On suppose f continue en $x_0 \in \mathbb{R}$. On dit que f est dérivable en $x_0 \in \mathbb{R}$ si la limite

$$\lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0},$$

existe. Si cette limite existe on la note $f'(x_0)$. Si f est dérivable en tout point de $\mathbb{R}$ on dit que f est dérivable sur $\mathbb{R}$. Considérons à présent une fonction $f : \mathbb{R} \rightarrow \mathbb{R}$ **continue ou continue par morceaux** sur $\mathbb{R}$ (ce qui explique la présence de $f(x_0^+)$ et $f(x_0^-)$ dans la suite). Comme dans le cas de la continuité, on définit les tangentes à droite et à gauche de f en un point $x_0 \in \mathbb{R}$ via

$$\lim_{\substack{x \rightarrow x_0 \\ x > x_0}} \frac{f(x) - f(x_0^+)}{x - x_0},$$

et

$$\lim_{\substack{x \rightarrow x_0 \\ x < x_0}} \frac{f(x) - f(x_0^-)}{x - x_0}.$$

Bien entendu si f est continue en x_0 on peut remplacer $f(x_0^+)$ et $f(x_0^-)$ par $f(x_0)$. Dans tous les cas, si ces limites existent on note

$$f'(x_0^+) = \lim_{\substack{x \rightarrow x_0 \\ x > x_0}} \frac{f(x) - f(x_0^+)}{x - x_0},$$

et

$$f'(x_0^-) = \lim_{\substack{x \rightarrow x_0 \\ x < x_0}} \frac{f(x) - f(x_0^-)}{x - x_0}.$$

Si f est continue et vérifie

$$f'(x_0) = f'(x_0^+) = f'(x_0^-),$$

alors f est dérivable en x_0 . Si par contre f est continue ou continue par morceaux sur $\mathbb{R}$, et que les quantités $f'(x_0^+)$ et $f'(x_0^-)$ existent mais $f'(x_0^+) \neq f'(x_0^-)$ on arrive alors à la notion de dérivabilité par morceaux.

Définition 9. Soit $f : \mathbb{R} \rightarrow \mathbb{R}$ une fonction continue ou continue par morceaux sur $\mathbb{R}$. On dit que f est dérivable par morceaux sur $\mathbb{R}$, si à chaque fois que f est restreinte à un intervalle non vide $[a, b]$ de $\mathbb{R}$ il existe un nombre fini de points $(x_i)_{1 \leq i \leq n}$ de $[a, b]$ vérifiant

$$a = x_1 < x_2 < \dots < x_m = b,$$

tel que la restriction de f à chaque intervalle $]x_i, x_{i+1}[$ ($1 \leq i \leq m-1$) se prolonge en une fonction dérivable sur $[x_i, x_{i+1}]$.

Si on reprend la première fonction de l'Exemple 10 alors f est continue sur $\mathbb{R}$ et f est dérivable par morceaux sur $\mathbb{R}$ car dérivable sur $\mathbb{R}$ avec $f'(x) = 0$ pour tout $x \in \mathbb{R}$. Si on considère à présent la troisième fonction, toujours de l'Exemple 10, cette fonction est continue par morceaux et dérivable par morceaux sur $\mathbb{R}$. En effet, si on restreint f à tout intervalle évitant 1 alors f est dérivable. Par ailleurs, si on considère sa restriction à l'intervalle $[0, 2]$ (par exemple) et la suite $x_1 = 0$, $x_2 = 1$ et $x_3 = 2$ alors la restriction de f à $]x_1, x_2[$ et $]x_2, x_3[$ est dérivable (sur les intervalles ouverts la fonction f est juste affine). Par ailleurs, on a

$$\begin{aligned} \lim_{\substack{x \rightarrow x_1 \\ x > x_1}} \frac{f(x) - f(x_1^+)}{x - x_1} &= \lim_{\substack{x \rightarrow 0 \\ x > 0}} \frac{(x-1) - (0-1)}{x-0} = 1, \\ \lim_{\substack{x \rightarrow x_2 \\ x < x_2}} \frac{f(x) - f(x_2^-)}{x - x_2} &= \lim_{\substack{x \rightarrow 1 \\ x < 1}} \frac{(x-1) - (1-1)}{x-1} = 1, \\ \lim_{\substack{x \rightarrow x_2 \\ x > x_2}} \frac{f(x) - f(x_2^+)}{x - x_2} &= \lim_{\substack{x \rightarrow 1 \\ x > 1}} \frac{x-1}{x-1} = 1, \\ \lim_{\substack{x \rightarrow x_3 \\ x < x_3}} \frac{f(x) - f(x_3^-)}{x - x_3} &= \lim_{\substack{x \rightarrow 2 \\ x < 2}} \frac{x-2}{x-2} = 1. \end{aligned}$$

Ainsi f se prolonge en une fonction dérivable sur $[x_1, x_2]$ et $[x_2, x_3]$ et donc f est dérivable par morceaux sur $\mathbb{R}$. Dans cet exemple il faut faire attention, ici $f'(x_2^-) = f'(x_2^+)$ mais la

fonction n'est pas dérivable en x_2 car celle-ci n'est pas continue en x_2 ! On peut utiliser un raisonnement assez similaire pour démontrer que la deuxième fonction de l'Exemple 10 est également dérivable par morceaux sur $\mathbb{R}$.

Donnons à présent un exemple de fonction continue sur $\mathbb{R}$ mais pas dérivable par morceaux sur $\mathbb{R}$.

Exemple 12. Soit f la fonction définie sur $\mathbb{R}$ par

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

Ici f est continue sur $\mathbb{R}$ mais pas dérivable par morceaux sur $\mathbb{R}$. En effet, il est clair que f est continue sur $] -\infty, 0[\cup]0, +\infty[$ et en 0 on a

$$\lim_{\substack{x \rightarrow 0 \\ x < 0}} f(x) = 0 = \lim_{\substack{x \rightarrow 0 \\ x > 0}} f(x) = \left(\lim_{\substack{x \rightarrow 0 \\ x > 0}} \sqrt{x} \right).$$

Par ailleurs, la restriction de f à tout intervalle évitant le point 0 est dérivable. Considérons à présent la restriction de f à (par exemple) l'intervalle $[-2, 1]$. On considère également la suite de points $x_1 = -2$, $x_2 = 0$ et $x_3 = 1$. La fonction f restreinte aux intervalles ouverts $]x_1, x_2[$ et $]x_2, x_3[$ est dérivable. On montre facilement que f se prolonge en une fonction dérivable sur l'intervalle fermé $[x_1, x_2]$ (dans ce cas f est juste la fonction nulle sur cet intervalle). Par contre, f ne se prolonge pas en une fonction dérivable sur l'intervalle fermé $[x_2, x_3]$. En effet, on a

$$\lim_{\substack{x \rightarrow 0 \\ x > 0}} \frac{f(x) - f(0)}{x - 0} = \lim_{\substack{x \rightarrow 0 \\ x > 0}} \frac{\sqrt{x}}{x} = \lim_{\substack{x \rightarrow 0 \\ x > 0}} \frac{1}{\sqrt{x}} = +\infty.$$

Donc $f'(0^+)$ n'existe pas et la fonction f n'est pas dérivable par morceaux.

Comme dans le cas d'une fonction continue ou continue par morceaux sur $\mathbb{R}$, si f est dérivable ou dérivable par morceaux on peut ensuite s'intéresser à la continuité ou la continuité par morceaux de f' . Dans le cas où f est continue, dérivable et que f' est continue sur $\mathbb{R}$ on dit que f est de classe C^1 sur $\mathbb{R}$ (par extension de la notation C^0 pour les fonctions continues). Dans le cas où f n'est pas C^1 sur $\mathbb{R}$ on peut assouplir cette notion via celle de C^1 par morceaux :

Définition 10. Soit $f : \mathbb{R} \rightarrow \mathbb{R}$ une fonction continue ou continue par morceaux sur $\mathbb{R}$. On dit que f est C^1 par morceaux sur $\mathbb{R}$, si à chaque fois que f est restreinte à un intervalle non vide $[a, b]$ de $\mathbb{R}$ il existe un nombre fini de points $(x_i)_{1 \leq i \leq n}$ de $[a, b]$ vérifiant

$$a = x_1 < x_2 < \dots < x_m = b,$$

tel que la restriction de f à chaque intervalle $]x_i, x_{i+1}[$ ($1 \leq i \leq m-1$) se prolonge en une fonction C^1 sur $[x_i, x_{i+1}]$.

La construction d'une fonction dérivable par morceaux sur $\mathbb{R}$ mais pas C^1 par morceaux sur $\mathbb{R}$ n'est complètement évidente. Une première idée peut être de considérer la primitive d'une fonction non continue par morceaux sur $\mathbb{R}$. En suivant cette idée, on considère la fonction non continue par morceaux sur $\mathbb{R}$ donnée par l'Exemple 11. Cependant, la primitive de cette fonction est de la forme

$$\begin{cases} \ln(x) & \text{si } x > 0, \\ 0 & \text{si } x \leq 0. \end{cases}$$

Or il est clair que cette fonction n'est pas non plus continue par morceaux sur $\mathbb{R}$. Il faut donc considérer un exemple plus raffiné. On a l'exemple classique suivant :

Exemple 13. Soit f la fonction définie sur $\mathbb{R}$ par

$$f(x) = \begin{cases} x^2 \sin\left(\frac{1}{x}\right) & \text{si } x \in \mathbb{R}^*, \\ 0 & \text{si } x = 0. \end{cases}$$

On commence par remarquer que f est continue sur $\mathbb{R}^*$ et pour tout $x \in \mathbb{R}^*$ on a

$$|f(x)| = \left| x^2 \sin\left(\frac{1}{x}\right) \right| \leq x^2.$$

On en déduit que $f(0^+) = f(0^-) = 0$ et donc f est bien continue sur $\mathbb{R}$. De même f est dérivable sur $\mathbb{R}^*$ comme produit et composée de fonctions dérivables sur $\mathbb{R}^*$ et pour tout $x \in \mathbb{R}^*$ on a

$$\frac{f(x) - f(0)}{x - 0} = x \sin\left(\frac{1}{x}\right).$$

Ainsi, on obtient que $f'(0^+) = f'(0^-) = 0$. La fonction f est donc dérivable sur $\mathbb{R}$ (et a fortiori dérivable par morceaux sur $\mathbb{R}$). Par contre cette fonction n'est pas C^1 par morceaux sur $\mathbb{R}$ ou autrement dit f' n'est pas continue par morceaux sur $\mathbb{R}$. En effet, pour tout $x \in \mathbb{R}^*$ on a

$$f'(x) = 2x \sin\left(\frac{1}{x}\right) - \cos\left(\frac{1}{x}\right).$$

Ici la fonction $x \mapsto 2x \sin\left(\frac{1}{x}\right)$ admet pour limite 0 quand x tend vers 0, par contre la fonction $x \mapsto \cos\left(\frac{1}{x}\right)$ n'admet pas de limite en 0. Ainsi, f' n'est pas continue en 0, i.e., f est continue et dérivable sur $\mathbb{R}$ mais pas C^1 par morceaux sur $\mathbb{R}$.

Dans la suite et sauf mention explicite du contraire, nous considérerons **toujours** des fonctions continues ou continues par morceaux et C^1 par morceaux (donc a fortiori dérivable par morceaux sur $\mathbb{R}$).

Remarque 15. Pour étudier la continuité, dérivabilité, continuité ou dérivabilité par morceaux sur $\mathbb{R}$ d'une fonction a -périodique il suffit de faire des études similaires à celles faites dans les exemples précédents sur un intervalle de longueur a , comme par exemple $[-a/2, a/2]$ ou $[0, a]$.

Bien entendu, on peut généraliser les notions précédentes en étudiant la dérivabilité ou la dérivabilité par morceaux de f' puis de f'' etc. Ceci conduit naturellement à la notion de fonction de classe C^p ou C^p par morceaux sur $\mathbb{R}$ pour tout $p \geq 1$.

Chapitre 2

Séries de Fourier

Résumé. Dans ce chapitre on étudie comment décomposer les fonctions périodiques comme somme de polynômes trigonométriques. Dans la première partie on explique « géométriquement » l'idée derrière les séries de Fourier. Puis on étudie le problème de la représentation ponctuelle d'une fonction périodique par sa série de Fourier.

Références. Pour ce chapitre la plupart des résultats et preuves viennent des livres :

- C. Gasquet et P. Witomski. *Analyse de Fourier et applications*, Dunod (1990), chapitre 2 (disponible à la BUTC),
- F. Boschet. *Séries de fonctions ; Intégrale de Riemann*, Masson (1995), chapitre 2 et chapitre 3,
- M. El Amrani. *Analyse de Fourier dans les espaces fonctionnels*, Ellipses (2008).

2.1 Préambules

Dans ce chapitre nous allons développer la théorie des séries de Fourier pour des fonctions (ou signaux) a -périodiques ($a > 0$) et continues par morceaux. Cette hypothèse est relativement contraignante d'un point de vue théorique mais beaucoup moins d'un point de vue pratique. On verra à la suite du chapitre 3 que la théorie développée ci-après peut être étendue aux fonctions appartenant à l'espace vectoriel

$$L_p^2(0, a) = \left\{ f : \mathbb{R} \rightarrow \mathbb{C} \text{ mesurable et } a\text{-périodique avec } \int_0^a |f(x)|^2 dx < \infty \right\}.$$

Le terme « mesurable » est bien entendu pour l'instant cryptique (et le restera même peut-être un peu après ce cours...). Admettons pour l'instant que toutes les fonctions raisonnables¹ définies sur $\mathbb{R}$ sont mesurables. Alors l'espace $L_p^2(0, a)$ désigne l'ensemble des signaux périodiques de puissance moyenne totale finie, i.e.,

$$\frac{1}{a} \int_0^a |f(x)|^2 dx < \infty, \quad \forall f \in L_p^2(0, a).$$

1. Par exemple les fonctions continues, continues par morceaux etc.

Sur l'espace $L_p^2(0, a)$ il est naturel de considérer la norme suivante

$$\|f\|_{L_p^2(0, a)} = \left(\int_0^a |f(x)|^2 dx \right)^{1/2}, \quad \forall f \in L_p^2(0, a).$$

Nous admettrons pour le moment que $L_p^2(0, a)$ est un espace normé et on remarque que cette norme résulte du produit scalaire hermitien :

$$\langle f, g \rangle_{L_p^2(0, a)} = \int_0^a f(x) \overline{g}(x) dx, \quad \forall f, g \in L_p^2(0, a).$$

L'espace $L_p^2(0, a)$ est donc un espace préhilbertien (il est en fait même un peu plus). Pour l'instant nous allons nous contenter de développer la théorie des séries de Fourier dans un sous-espace vectoriel de $L_p^2(0, a)$. Remarquons dans un premier temps que l'espace $C_p^0([0, a]; \mathbb{C})$ des fonctions a -périodiques continues à valeurs dans $\mathbb{C}$ est un sous-espace vectoriel de $L_p^2(0, a)$. En effet, si $f \in C_p^0([0, a])$ alors f est R-intégrable sur $[0, a]$ et donc bornée, i.e., il existe $M \in \mathbb{R}$ vérifiant $|f(x)| \leq M$. Ainsi

$$\int_0^a |f(x)|^2 dx \leq M^2 a < \infty,$$

et donc $C_p^0([0, a]; \mathbb{C}) \subset L_p^2([0, a])$. De plus cet espace est préhilbertien pour le produit scalaire $\langle \cdot, \cdot \rangle_{L_p^2(0, a)}$. En argumentant de la même manière on prouve que l'espace vectoriel des fonctions a -périodiques et continues par morceaux à valeurs dans $\mathbb{C}$, noté $C_{p, m}^0([0, a]; \mathbb{C})$, est un sous-espace vectoriel préhilbertien (pour le produit scalaire $\langle \cdot, \cdot \rangle_{L_p^2(0, a)}$) de $L_p^2([0, a])$. **Dans tout la suite, sauf mention explicite du contraire, afin de simplifier les notations on désignera par $\mathcal{H}$ l'espace $C_{p, m}^0([0, a]; \mathbb{C})$.**

2.2 Séries de Fourier d'une fonction périodique

L'objectif premier des séries de Fourier est de décomposer par le biais de fonctions "simples" une fonction (ou signal) périodique. Pour ce faire nous allons tout d'abord considérer les fonctions périodiques les plus naturelles, à savoir les fonctions cosinus et sinus. Nous aborderons ensuite la décomposition de fonctions périodiques. Nous aborderons également le développement d'une fonction sur une base orthogonale.

2.2.1 Polynômes trigonométriques

Soit e_n la fonction a périodique (avec $a > 0$) définie sur $\mathbb{R}$ par

$$e_n(t) = \exp\left(2i\pi n \frac{t}{a}\right), \quad n \in \mathbb{Z}, \forall t \in \mathbb{R}.$$

Alors la fonction p définie sur $\mathbb{R}$ par

$$p(t) = \sum_{n=-N}^N c_n e_n(t) = \sum_{n=-N}^N c_n \exp\left(2i\pi n \frac{t}{a}\right), \quad (2.1)$$

où $N \geq 1$ et $c_n \in \mathbb{C}$, est une fonction également a périodique. On appelle cette fonction un polynôme trigonométrique. On a déjà parlé de ces polynômes au chapitre 1.

Remarquons dans un premier temps qu'on peut réécrire ces fonctions sous la forme

$$p(t) = c_0 + \sum_{n=1}^N \left(c_n \exp\left(2i\pi n \frac{t}{a}\right) + c_{-n} \exp\left(-2i\pi n \frac{t}{a}\right) \right), \quad \forall t \in \mathbb{R},$$

ou encore sous la forme

$$p(t) = \frac{1}{2} a_0 + \sum_{n=1}^N \left(a_n \cos\left(2\pi n \frac{t}{a}\right) + b_n \sin\left(2\pi n \frac{t}{a}\right) \right), \quad \forall t \in \mathbb{R},$$

avec pour $n \geq 0$

$$\begin{aligned} a_n &= c_n + c_{-n}, \\ b_n &= i(c_n - c_{-n}). \end{aligned}$$

On retrouve également les formules inverses

$$\begin{aligned} c_n &= (a_n - ib_n)/2, \\ c_{-n} &= (a_n + ib_n)/2. \end{aligned}$$

La question qui se pose à présent est la suivante : soit p un polynôme trigonométrique quelconque de la forme (2.1), comment déterminer l'expression de ses coefficients c_n pour $n = -N, \dots, N$?

En suivant les notations du chapitre 1, notons $\mathcal{T}_N^a$ l'espace vectoriel des polynômes trigonométriques p (de la forme (2.1)) de degré inférieur ou égal à N et de période a . D'après le premier chapitre nous savons que ce $\mathbb{C}$ espace vectoriel est de dimension finie avec $\dim(\mathcal{T}_N^a) = 2N + 1$. Nous avons également vu que la famille

$$\mathcal{B}_N = (e_{-N}, e_{-(N-1)}, \dots, e_0, \dots, e_{N-1}, e_N),$$

constitue une base de $\mathcal{T}_N^a$. De plus, comme $\mathcal{T}_N^a \subset C_p^0([0, a]; \mathbb{C})$ alors cet espace est hermitien pour le produit scalaire de $L_p^2(0, a)$:

$$\langle p, q \rangle_{L_p^2(0, a)} = \int_0^a p(t) \bar{q}(t) \, dt, \quad \forall p, q \in \mathcal{T}_N^a.$$

Par conséquent $\mathcal{T}_N^a$ est normé, avec,

$$\|p\|_{L_p^2(0,a)} = \sqrt{\langle p, p \rangle} = \left(\int_0^a |p(t)|^2 dt \right)^{1/2}, \quad \forall p \in \mathcal{T}_N^a.$$

Il reste alors à remarquer que $\mathcal{B}_N$ est une famille orthogonale, i.e.,

$$\langle e_n, e_m \rangle_{L_p^2(0,a)} = \int_0^a e_n(t) \overline{e_m(t)} dt = \begin{cases} 0 & \text{si } n \neq m \\ a & \text{si } n = m. \end{cases}$$

Ainsi d'après la Remarque 12 du chapitre 1 pour déterminer les coefficients du polynôme p il suffit de « projeter » p via la produit scalaire sur les éléments de la base $\mathcal{B}_N$. On a

$$\langle p, e_n \rangle_{L_p^2(0,a)} = c_n \langle e_n, e_n \rangle_{L_p^2(0,a)} = c_n \|e_n\|_{L_p^2(0,a)}^2 = a c_n, \quad \forall n \in \{-N, \dots, N\}. \quad (2.2)$$

Ainsi on obtient alors la formule, dite de Fourier,

$$c_n = \frac{\langle p, e_n \rangle_{L_p^2(0,a)}}{\|e_n\|_{L_p^2(0,a)}^2} = \frac{1}{a} \int_0^a p(t) \exp \left(-2i\pi n \frac{t}{a} \right) dt, \quad \forall n \in \{-N, \dots, N\}.$$

Remarquons également que par le biais de (2.2) on trouve également les formules :

$$\begin{aligned} a_n &= \frac{2}{a} \int_0^a p(t) \cos \left(2\pi n \frac{t}{a} \right) dt, \\ b_n &= \frac{2}{a} \int_0^a p(t) \sin \left(2\pi n \frac{t}{a} \right) dt. \end{aligned}$$

Ainsi connaissant les valeurs $p(t)$ pour tout t dans un intervalle de longueur a (avec $p \in \mathcal{T}_N^a$) on peut obtenir la valeurs de ces coefficients c_n et ainsi reconstruire l'expression analytique de la fonction p .

Enfin pour conclure cette section remarquons que pour un polynôme $p \in \mathcal{T}_N^a$, nous avons

$$\|p\|_{L_p^2(0,a)}^2 = \int_0^a p(t) \overline{p(t)} dt = \int_0^a \left(\sum_{n=-N}^N c_n e_n(t) \right) \left(\sum_{m=-N}^N \overline{c_m} \overline{e_m(t)} \right) dt.$$

Par orthogonalité de la base $\mathcal{B}_N$ on obtient :

$$\|p\|_{L_p^2(0,a)}^2 = \sum_{n=-N}^N \sum_{m=-N}^N c_n \overline{c_m} \int_0^a e_n(t) \overline{e_m(t)} dt = \sum_{n=-N}^N \sum_{m=-N}^N c_n \overline{c_m} \langle e_n, e_m \rangle_{L_p^2(0,a)},$$

ce qui implique

$$\|p\|_{L_p^2(0,a)}^2 = a \sum_{n=-N}^N c_n \overline{c_n}.$$

On obtient alors l'égalité, dite de Parseval, pour un polynôme trigonométrique

$$\sum_{n=-N}^N |c_n|^2 = \frac{1}{a} \int_0^a |p(t)|^2 dt = \frac{1}{a} \|p\|_{L_p^2(0,a)}^2. \quad (2.3)$$

2.2.2 Fonctions périodiques

Considérons à présent une fonction $f \in \mathcal{H}$ pas nécessairement de la forme (2.1). Est-il alors possible d'écrire cette fonction sous la forme

$$f(t) = \sum_{n \in I} c_n \exp\left(2i\pi n \frac{t}{a}\right) \quad ? \quad (2.4)$$

Dans cette somme I est un sous ensemble de $\mathbb{Z}$. L'idée ici est de décomposer f en utilisant des fonctions périodiques simples e_n . La réponse à la question précédente est généralement négative si I est un sous ensemble de cardinal fini de $\mathbb{Z}$. En effet, comme les fonctions e_n sont indéfiniment dérivables (C^∞), si I est de la forme $\{-N, \dots, N\}$ alors f est nécessairement indéfiniment dérivable. Or si $f \in \mathcal{H}$ alors f n'a aucune raison d'être C^∞ , comme par exemple sur la Figure 2.1.

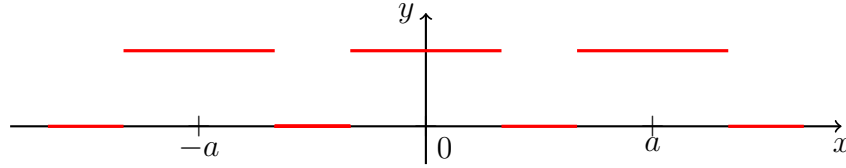


FIGURE 2.1 – Représentation graphique d'une fonction dans $\mathcal{H}$ mais pas dans $C_p^0([0, a])$.

Ainsi on comprend que pour que la réponse à la question initiale soit positive il faut généralement supposer que $I = \mathbb{Z}$. Mais dans ce cas quel sens donner à cette somme infinie ? Pour quelle fonction f peut on écrire ponctuellement l'égalité (2.4) ? Nous reviendrons sur ces questions concernant la convergence dans la Section 2.3. Pour l'instant nous allons nous contenter d'approcher une fonction $f \in \mathcal{H}$ par des polynômes trigonométriques de $\mathcal{T}_N^a$.

Pour ce faire, comme dans section précédente, on rappelle que $\mathcal{H}$ est un espace préhilbertien pour le produit scalaire suivant :

$$\langle f, g \rangle_{L_p^2(0,a)} = \int_0^a f(t) \overline{g(t)} dt, \quad \forall f, g \in \mathcal{H}.$$

Ce qui donne la norme

$$\|f\|_{L_p^2(0,a)} = \left(\int_0^a |f(t)|^2 dt \right)^{1/2}, \quad \forall f \in \mathcal{H}.$$

Remarquons que l'espace vectoriel $\mathcal{H}$ est de dimension infinie. En effet pour tout $N \geq 1$ la famille $(e_n)_{n \in \{-N, \dots, N\}}$ est libre dans $\mathcal{H}$ car orthogonale. Ainsi si on suppose que $\dim(\mathcal{H}) = 2N_0 + 1 \geq 1$ comme la famille orthogonale $(e_n)_{n \in \{-(N_0+1), \dots, N_0+1\}}$ est libre dans $\mathcal{H}$ on aboutit à une contradiction.

Pour approcher une fonction $f \in \mathcal{H}$ par un polynôme trigonométrique, l'idée principale est d'approcher f par l'élément de $\mathcal{T}_N^a$ le plus proche possible au sens de la norme $\|\cdot\|_{L_p^2(0,a)}$.

On se fixe un entier $N \geq 1$ et on cherche alors des coefficients $x_{-N}, \dots, x_N$ telle que la quantité suivante

$$\left\| f - \sum_{n=-N}^N x_n e_n \right\|_{L_p^2(0,a)},$$

soit la plus petite possible. On peut interpréter ce problème de manière géométrique comme suit : soit $f \in \mathcal{H}$ trouver f_N appartenant à l'espace vectoriel de dimension finie $\mathcal{T}_N^a$ qui soit à distance minimale de f . Si un tel polynôme f_N existe on dit que c'est la meilleure approximation de f dans $\mathcal{T}_N^a$.

Pour résoudre ce problème on considère $p \in \mathcal{T}_N^a$ quelconque de la forme

$$p(t) = \sum_{n=-N}^N x_n e_n(t), \quad \forall t \in \mathbb{R}.$$

Cherchons à évaluer la distance au sens $\|\cdot\|_{L_p^2(0,a)}$ de f à p . En utilisant les propriétés d'un produit scalaire hermitien on a :

$$\begin{aligned} \|f - p\|_{L_p^2(0,a)}^2 &= \langle f - p, f - p \rangle_{L_p^2(0,a)} \\ &= \|f\|_{L_p^2(0,a)}^2 - \langle f, p \rangle_{L_p^2(0,a)} - \langle p, f \rangle_{L_p^2(0,a)} + \|p\|_{L_p^2(0,a)}^2 \\ &= \|f\|_{L_p^2(0,a)}^2 - \left(\langle f, p \rangle_{L_p^2(0,a)} + \overline{\langle f, p \rangle_{L_p^2(0,a)}} \right) + \|p\|_{L_p^2(0,a)}^2. \end{aligned}$$

En utilisant à présent les propriétés élémentaires sur les nombres complexes, on obtient

$$\|f - p\|_{L_p^2(0,a)}^2 = \|f\|_{L_p^2(0,a)}^2 - 2 \operatorname{Re}(\langle f, p \rangle_{L_p^2(0,a)}) + \|p\|_{L_p^2(0,a)}^2.$$

Ensuite d'après la relation de Parseval pour les polynômes trigonométriques (2.3) on a

$$\|p\|_{L_p^2(0,a)}^2 = a \sum_{n=-N}^N |x_n|^2.$$

On a également la relation

$$\langle f, p \rangle_{L_p^2(0,a)} = \sum_{n=-N}^N \bar{x}_n \langle f, e_n \rangle_{L_p^2(0,a)}.$$

On introduit maintenant les coefficients de Fourier de f

$$c_n(f) = \frac{\langle f, e_n \rangle_{L_p^2(0,a)}}{\|e_n\|_{L_p^2(0,a)}^2} = \frac{1}{a} \int_0^a f(t) \overline{e_n}(t) dt = \frac{1}{a} \int_0^a f(t) \exp\left(-2i\pi n \frac{t}{a}\right) dt,$$

pour simplifier les notations on écrira régulièrement c_n à la place de $c_n(f)$. De plus comme dans la section précédente on peut aussi introduire les coefficients $a_n(f)$ et $b_n(f)$:

$$\begin{aligned} a_n(f) &= \frac{2}{a} \int_0^a f(t) \cos\left(2\pi n \frac{t}{a}\right) dt, \\ b_n(f) &= \frac{2}{a} \int_0^a f(t) \sin\left(2\pi n \frac{t}{a}\right) dt, \end{aligned}$$

avec les relations suivantes entre les coefficients

$$\begin{aligned} c_n(f) &= \frac{a_n(f) - ib_n(f)}{2}, \\ c_{-n}(f) &= \frac{a_n(f) + ib_n(f)}{2}. \end{aligned}$$

Revenons à présent à l'estimation de $\|f - p\|_{L_p^2(0,a)}$. On a

$$\begin{aligned} &\|f - p\|_{L_p^2(0,a)}^2 \\ &= \|f\|_{L_p^2(0,a)}^2 + a \sum_{n=-N}^N |x_n|^2 - 2 \operatorname{Re} \left(\sum_{n=-N}^N \bar{x}_n \langle f, e_n \rangle_{L_p^2(0,a)} \right), \\ &= \|f\|_{L_p^2(0,a)}^2 + a \sum_{n=-N}^N |x_n|^2 - 2 \operatorname{Re} \left(a \sum_{n=-N}^N \bar{x}_n c_n(f) \right) \\ &= \|f\|_{L_p^2(0,a)}^2 + a \sum_{n=-N}^N |x_n|^2 - 2a \sum_{n=-N}^N \operatorname{Re}(\bar{x}_n c_n(f)) \\ &= \|f\|_{L_p^2(0,a)}^2 + a \sum_{n=-N}^N |x_n|^2 - 2a \sum_{n=-N}^N \operatorname{Re}(\bar{x}_n c_n(f)) + a \sum_{n=-N}^N |c_n(f)|^2 - a \sum_{n=-N}^N |c_n(f)|^2 \\ &= \|f\|_{L_p^2(0,a)}^2 + a \sum_{n=-N}^N (|x_n|^2 - 2 \operatorname{Re}(\bar{x}_n c_n(f)) + |c_n(f)|^2) - a \sum_{n=-N}^N |c_n(f)|^2 \\ &= \|f\|_{L_p^2(0,a)}^2 + a \sum_{n=-N}^N |x_n - c_n(f)|^2 - a \sum_{n=-N}^N |c_n(f)|^2. \end{aligned}$$

D'où

$$\|f - p\|_{L_p^2(0,a)}^2 = \|f\|_{L_p^2(0,a)}^2 + a \sum_{n=-N}^N (|x_n - c_n(f)|^2 - |c_n(f)|^2). \quad (2.5)$$

La quantité (2.5) est minimale si $x_n = c_n(f)$ pour $n = -N, \dots, N$. En effet, notons $\tilde{p} \in \mathcal{T}_N^a$ un polynôme de la forme

$$\tilde{p}(t) = \sum_{n=-N}^N c_n(f) e_n(t), \quad \forall t \in \mathbb{R}.$$

Alors on a d'après (2.5)

$$\begin{aligned}
\|f - p\|_{L_p^2(0,a)}^2 - \|f - \tilde{p}\|_{L_p^2(0,a)}^2 &= \|f\|_{L_p^2(0,a)}^2 + a \sum_{n=-N}^N (|x_n - c_n(f)|^2 - |c_n(f)|^2) \\
&\quad - \|f\|_{L_p^2(0,a)}^2 + a \sum_{n=-N}^N |c_n(f)|^2 \\
&= a \sum_{n=-N}^N |x_n - c_n(f)|^2 \geq 0.
\end{aligned} \tag{2.6}$$

Donc pour tout $p \in \mathcal{T}_N^a$ on a

$$\|f - p\|_{L_p^2(0,a)}^2 \geq \|f - \tilde{p}\|_{L_p^2(0,a)}^2.$$

Ainsi il existe au moins un polynôme trigonométrique f_N réalisant la meilleure approximation de f dans $\mathcal{T}_N^a$ donné par

$$f_N(t) = \sum_{n=-N}^N c_n(f) e_n(t) = \frac{1}{a} \sum_{n=-N}^N \langle f, e_n \rangle_{L_p^2(0,a)} e_n(t), \quad \forall t \in \mathbb{R}.$$

Le polynôme f_N est obtenu en projetant, via le produit scalaire $\langle \cdot, \cdot \rangle_{L_p^2(0,a)}$, f sur le sous-espace vectoriel de dimension finie $\mathcal{T}_N^a$. De plus remarquons que le polynôme f_N réalisant la meilleure approximation de f dans $\mathcal{T}_N^a$ est unique. En effet supposons l'existence d'un polynôme $\tilde{p}(t) = \sum_{n=-N}^N x_n e_n(t)$ réalisant

$$\|f - \tilde{p}\|_{L_p^2(0,a)} = \min\{\|f - p\|_{L_p^2(0,a)}, p \in \mathcal{T}_N^a\} (= \|f_N - f\|_{L_p^2(0,a)}),$$

avec $\tilde{p}(t) \neq f_N(t)$ pour tout $t \in \mathbb{R}$, alors

$$\|f - f_N\|_{L_p^2(0,a)}^2 = \|f - \tilde{p}\|_{L_p^2(0,a)}^2.$$

Par conséquent, d'après (2.6), on a

$$\sum_{n=-N}^N |x_n - c_n(f)|^2 = 0 \Rightarrow x_n = c_n(f).$$

Donc f_N est unique. Dans la suite on appellera f_N la somme partielle d'ordre N de la série de Fourier de f .

On vient de démontrer une (grande) partie du résultat important suivant :

Théorème 16. *Soit $f \in \mathcal{H}$ il existe un unique polynôme trigonométrique f_N dans $\mathcal{T}_N^a$ vérifiant*

$$\|f - f_N\|_{L_p^2(0,a)} = \min\{\|f - p\|_{L_p^2(0,a)}, p \in \mathcal{T}_N^a\}.$$

Ce polynôme f_N , appelé somme partielle d'ordre N de la série de Fourier de f , est donné par

$$f_N(t) = \sum_{n=-N}^N c_n(f) \exp\left(2i\pi n \frac{t}{a}\right), \quad \forall t \in \mathbb{R},$$

où

$$c_n(f) = \frac{\langle f, e_n \rangle_{L_p^2(0,a)}}{\|e_n\|_{L_p^2(0,a)}^2} = \frac{1}{a} \int_0^a f(t) \exp\left(-2i\pi n \frac{t}{a}\right) dt, \quad n = -N, \dots, N.$$

Le polynôme f_N peut également s'écrire sous la forme

$$f_N(t) = \frac{a_0(f)}{2} + \sum_{n=1}^N \left(a_n(f) \cos\left(2\pi n \frac{t}{a}\right) + b_n(f) \sin\left(2\pi n \frac{t}{a}\right) \right),$$

où

$$\begin{aligned} a_0(f) &= 2c_0(f) = \frac{2}{a} \int_0^a f(t) dt, \\ a_n(f) &= c_n(f) + c_{-n}(f) = \frac{2}{a} \int_0^a f(t) \cos\left(2\pi n \frac{t}{a}\right) dt, \\ b_n(f) &= i(c_n(f) - c_{-n}(f)) = \frac{2}{a} \int_0^a f(t) \sin\left(2\pi n \frac{t}{a}\right) dt. \end{aligned}$$

Les fonctions f et f_N vérifient la relation

$$\|f - f_N\|_{L_p^2(0,a)}^2 = \|f\|_{L_p^2(0,a)}^2 - a \sum_{n=-N}^N |c_n(f)|^2. \quad (2.7)$$

De plus, pour tout $f \in \mathcal{H}$ on a convergence en moyenne quadratique de $(f_N)_{N \in \mathbb{N}}$ (sa série de Fourier) vers f , i.e.,

$$\|f - f_N\|_{L_p^2(0,a)} = \left(\int_0^a |f(x) - f_N(x)|^2 dx \right)^{1/2} \rightarrow 0, \quad \text{quand } N \rightarrow +\infty. \quad (2.8)$$

Suite à l'énoncé de ce théorème plusieurs commentaires s'imposent :

- On admet dans l'énoncé du théorème le résultat portant sur la convergence en moyenne quadratique (2.8). Suite à notre construction des séries de Fourier, ce résultat n'est pas étonnant et correspond à l'intuition « plus N est grand mieux on approche f ».
- Il faut cependant faire attention, la convergence en moyenne quadratique n'implique pas que

$$f(t) = \lim_{N \rightarrow +\infty} f_N(t) = \sum_{n \in \mathbb{Z}} c_n(f) e_n(t), \quad \forall t \in \mathbb{R},$$

voir l'exemple 15. On verra dans la Section 2.3, sous quelles conditions cette égalité est vérifiée.

- Enfin pour les plus sceptiques donnons quelques mots sur l'idée de la preuve de (2.8). L'espace vectoriel $\mathcal{H}$ est de dimension infinie et il est tentant de voir la famille $(e_n)_{n \in \mathbb{Z}}$ comme une base de $\mathcal{H}$. Cette affirmation est fausse mais pas si fausse que ça... En effet si la famille $(e_n)_{n \in \mathbb{Z}}$ ne constitue pas une base au sens algébrique (le sens « usuel » en algèbre linéaire) de $\mathcal{H}$, elle forme une famille dite totale (ou base topologique de $\mathcal{H}$). La différence étant la suivante : on ne peut pas toujours écrire, pour un certain $N \geq 1$, $f \in \mathcal{H}$ comme

$$f(t) = \sum_{n=-N}^N x_n e_n(t), \quad \forall t \in \mathbb{R}, \quad x_{-N}, \dots, x_N \in \mathbb{C},$$

mais on peut approcher d'aussi près que l'on veut (au sens de la norme de $\mathcal{H}$) toute fonction f de $\mathcal{H}$ par des combinaisons linéaires finies de la famille $(e_n)_{n \in \mathbb{Z}}$. Autrement dit, pour tout $\varepsilon > 0$, il existe un entier N_1 et des scalaires $x_{-N_1}, \dots, x_{N_1}$ vérifiant

$$\left\| f - \sum_{n=-N_1}^{N_1} x_n e_n \right\|_{L_p^2(0,a)} \leq \varepsilon.$$

Vocabulaire 17 (Notion de spectre d'un signal périodique). *Pour un signal f , de période a et développé en série de Fourier :*

$$f(x) = \sum_{n \in \mathbb{Z}} c_n(f) \exp\left(2i\pi n \frac{x}{a}\right),$$

on appelle spectre de f l'ensemble des couples $(n/a, c_n(f))_{n \in \mathbb{Z}}$. Ainsi tout signal périodique de $\mathcal{H}$ admet une représentation en fréquence (voir Figure 2.2). Ce spectre est constitué de raies régulièrement espacées à la fréquence $1/a$. Pour $|n| = 1$, les raies c_1 et c_{-1} correspondent à la fréquence dite fondamentale. Les autres raies s'appellent les harmoniques du signal. Enfin, on remarque, d'après les formules

$$c_n(f) = \frac{a_n(f) - ib_n(f)}{2}, \quad c_{-n}(f) = \frac{a_n(f) + ib_n(f)}{2},$$

pour tout $n \geq 1$, que $|c_n(f)| = |c_{-n}(f)|$.

Donnons deux conséquences importantes du Théorème 16.

Corollaire 18 (Inégalité de Bessel). *Pour tout $f \in \mathcal{H}$, alors*

$$\sum_{n=-N}^N |c_n(f)|^2 \leq \frac{1}{a} \int_0^a |f(t)|^2 dt, \quad \forall N \in \mathbb{N}.$$

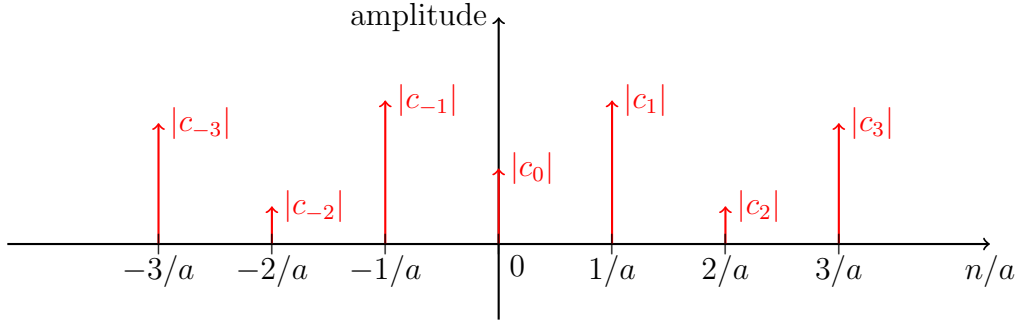


FIGURE 2.2 – Spectre d'amplitude d'un signal.

Démonstration. D'après l'égalité (2.7) du Théorème 16, on a

$$\|f\|_{L_p^2(0,a)}^2 - a \sum_{n=-N}^N |c_n(f)|^2 = \|f - f_N\|_{L_p^2(0,a)}^2 \geq 0,$$

donc

$$\|f\|_{L_p^2(0,a)}^2 \geq a \sum_{n=-N}^N |c_n(f)|^2.$$

Ce qui conclut la preuve du résultat. □

Corollaire 19 (Égalité de Parseval). *Pour tout $f \in \mathcal{H}$, on a l'égalité de Parseval*

$$\sum_{n \in \mathbb{Z}} |c_n(f)|^2 = \frac{1}{a} \int_0^a |f(t)|^2 dt. \quad (2.9)$$

On peut réécrire cette égalité sous la forme

$$\frac{1}{a} \int_0^a |f(t)|^2 dt = \frac{|a_0(f)|^2}{4} + \frac{1}{2} \sum_{n=1}^{+\infty} (|a_n(f)|^2 + |b_n(f)|^2). \quad (2.10)$$

Démonstration. L'égalité (2.9) est une conséquence de la relation (2.7) du Théorème 16. En effet, pour tout $f \in \mathcal{H}$, on a

$$\|f\|_{L_p^2(0,a)}^2 - a \sum_{n=-N}^N |c_n(f)|^2 = \|f - f_N\|_{L_p^2(0,a)}^2 \rightarrow 0, \quad \text{quand } N \rightarrow +\infty.$$

Pour l'égalité (2.10), on remarque que

$$c_n(f) = \frac{a_n(f) - ib_n(f)}{2}, \quad n \geq 1,$$

$$c_{-n}(f) = \frac{a_n(f) + ib_n(f)}{2}, \quad n \geq 1.$$

De plus comme $a_0(f) = 2c_0(f)$ on a, pour tout $N \geq 1$

$$\begin{aligned} \sum_{n=-N}^N |c_n(f)|^2 &= |c_0(f)|^2 + \sum_{n=1}^N (|c_n(f)|^2 + |c_{-n}(f)|^2) \\ &= \frac{|a_0(f)|^2}{4} + \sum_{i=1}^N \left(\frac{|a_n(f)|^2 + |b_n(f)|^2}{4} + \frac{|a_n(f)|^2 + |b_n(f)|^2}{4} \right) \\ &= \frac{|a_0(f)|^2}{4} + \frac{1}{2} \sum_{i=1}^N (|a_n(f)|^2 + |b_n(f)|^2). \end{aligned}$$

On peut ainsi réécrire la relation (2.7) du Théorème 16 sous la forme

$$\|f\|_{L_p^2(0,a)}^2 - a \left(\frac{|a_0(f)|^2}{4} + \frac{1}{2} \sum_{i=1}^N (|a_n(f)|^2 + |b_n(f)|^2) \right) = \|f - f_N\|_{L_p^2(0,a)}^2.$$

Comme $\|f - f_N\|_{L_p^2(0,a)}^2 \rightarrow 0$ quand $N \rightarrow +\infty$, ceci conclut la preuve du résultat. $\square$

Exemple 14. *Donnons un premier exemple. Soit $f : \mathbb{R} \rightarrow \mathbb{R}$, la fonction 2π -périodique constante et égale à un, i.e., $f(x) = 1$ pour tout $x \in \mathbb{R}$. Si dans la théorie on utilise les coefficients de Fourier $c_n(f)$, dans la pratique on calcule plutôt les coefficients $a_n(f)$ et $b_n(f)$. Ici on a*

$$a_0(f) = \frac{2}{2\pi} \int_0^{2\pi} f(x) dx = \frac{1}{\pi} \int_0^{2\pi} dx = 2.$$

On a ensuite, pour tout $n \geq 1$,

$$\begin{aligned} a_n(f) &= \frac{2}{2\pi} \int_0^{2\pi} f(x) \cos\left(2\pi n \frac{x}{2\pi}\right) dx = \frac{1}{\pi} \int_0^{2\pi} \cos(nx) dx \\ &= \frac{1}{\pi} \left[\frac{\sin(nx)}{n} \right]_0^{2\pi} \\ &= 0 \end{aligned}$$

et

$$\begin{aligned} b_n(f) &= \frac{2}{2\pi} \int_0^{2\pi} f(x) \sin\left(2\pi n \frac{x}{2\pi}\right) dx = \frac{1}{\pi} \int_0^{2\pi} \sin(nx) dx \\ &= \frac{1}{\pi} \left[\frac{-\cos(nx)}{n} \right]_0^{2\pi} \\ &= \frac{-\cos(2\pi n) + 1}{n\pi} \end{aligned}$$

$$= 0,$$

car pour tout $n \geq 1$, l'entier $2n$ est pair et donc $\cos(2\pi n) = 1$. En résumé, pour tout $N \geq 1$, la somme partielle d'ordre N de la série de Fourier de f est donnée par

$$f_N(x) = \frac{a_0(f)}{2} = 1.$$

En particulier, dans ce cas, on a $\lim_{N \rightarrow +\infty} f_N(x) = 1 = f(x)$ pour tout $x \in \mathbb{R}$.

Exemple 15. Donnons un second exemple illustrant les calculs des coefficients de Fourier d'une fonction 2π -périodique et continue par morceaux. Soit f la fonction 2π -périodique :

$$f(x) = \begin{cases} 1 & \text{si } 0 \leq x < \pi, \\ -1 & \text{si } \pi \leq x < 2\pi. \end{cases}$$

On commence par calculer le coefficient $a_0(f)$, on a

$$a_0(f) = \frac{2}{2\pi} \int_0^{2\pi} f(x) dx = \frac{1}{\pi} \int_0^{\pi} dx - \frac{1}{\pi} \int_{\pi}^{2\pi} dx = 1 - 1 = 0.$$

On a ensuite, pour $n \geq 1$,

$$\begin{aligned} a_n(f) &= \frac{2}{2\pi} \int_0^{2\pi} f(x) \cos\left(2\pi n \frac{x}{2\pi}\right) dx \\ &= \frac{1}{\pi} \int_0^{\pi} \cos(nx) dx - \frac{1}{\pi} \int_{\pi}^{2\pi} \cos(nx) dx \\ &= \frac{1}{n\pi} [\sin(nx)]_0^{\pi} - \frac{1}{n\pi} [\sin(nx)]_{\pi}^{2\pi} = 0, \end{aligned}$$

et

$$\begin{aligned} b_n(f) &= \frac{2}{2\pi} \int_0^{2\pi} f(x) \sin\left(2\pi n \frac{x}{2\pi}\right) dx \\ &= \frac{1}{\pi} \int_0^{\pi} \sin(nx) dx - \frac{1}{\pi} \int_{\pi}^{2\pi} \sin(nx) dx \\ &= \frac{1}{n\pi} [-\cos(nx)]_0^{\pi} - \frac{1}{n\pi} [-\cos(nx)]_{\pi}^{2\pi} \\ &= \frac{1}{n\pi} (-\cos(n\pi) + 1) - \frac{1}{n\pi} (-\cos(2n\pi) + \cos(n\pi)) \\ &= \frac{1}{n\pi} (-(-1)^n + 1) - \frac{1}{n\pi} (-1 + (-1)^n) \\ &= \frac{2}{n\pi} ((-1)^{n+1} + 1). \end{aligned}$$

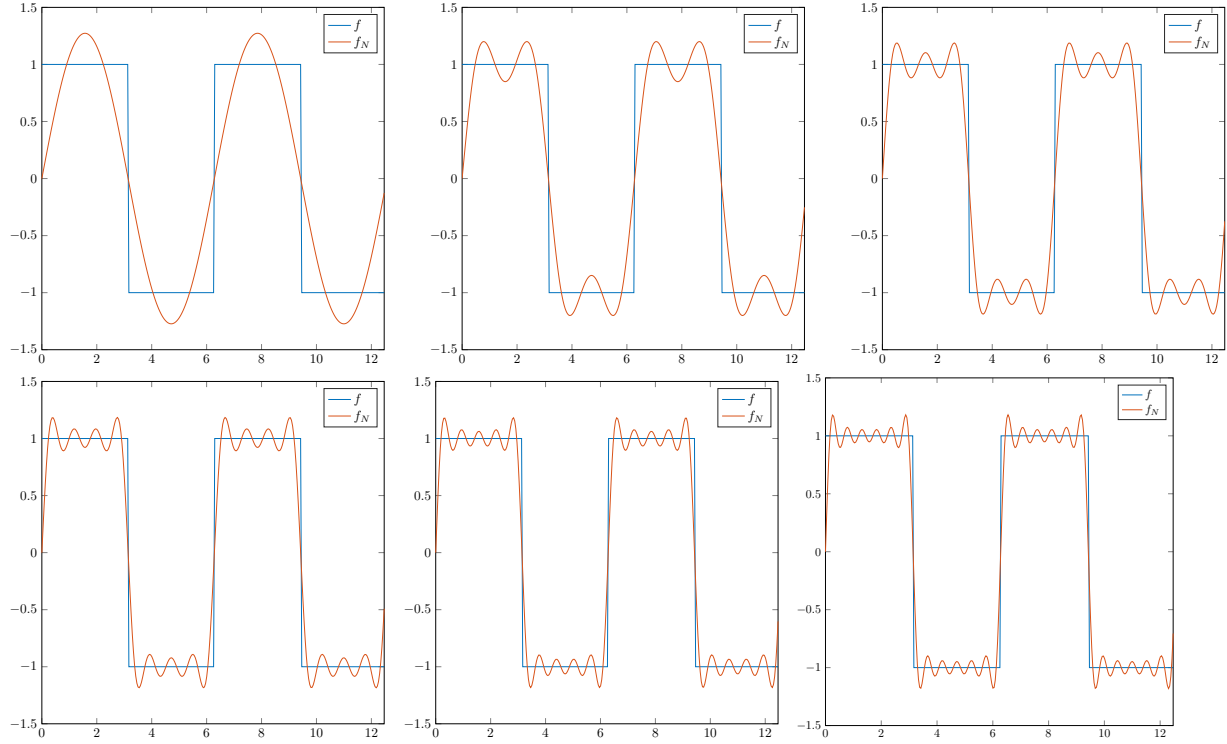


FIGURE 2.3 – Illustration de l'exemple 15 concernant la convergence de f_N vers f pour $N = 1, 3, 5, 7, 9$ et 11 .

Pour conclure, pour tout $N \geq 1$, la somme partielle d'ordre N de la série de Fourier de f est donnée par

$$f_N(x) = \frac{2}{\pi} \sum_{n=1}^N \frac{((-1)^{n+1} + 1)}{n} \sin(nx).$$

La Figure 2.3, illustre la convergence de f_N vers f . On remarque également que $f_N(\pi) = 0$, pour tout $N \geq 1$. Cependant $f(\pi) = -1$. Ainsi il est impossible d'avoir

$$f(\pi) = \lim_{N \rightarrow +\infty} f_N(\pi) (= 0).$$

Suite à l'Exemple 15 revenons sur les explications du Théorème 16. Ce théorème implique que la suite de nombres réels $(\|f - f_N\|_{L_p^2(0,2\pi)})_{N \geq 1}$ converge vers 0 quand $N \rightarrow +\infty$, i.e., on a

$$\|f - f_N\|_{L_p^2(0,2\pi)} \rightarrow 0, \quad \text{quand } N \rightarrow +\infty.$$

Cependant cette convergence en moyenne quadratique n'implique pas nécessairement la convergence vers 0 (quand $N \rightarrow +\infty$) de la suite $(|f(x) - f_N(x)|)_{N \geq 1}$ pour tout $x \in \mathbb{R}$. Ainsi on a pas nécessairement que pour tout $x \in \mathbb{R}$

$$|f(x) - f_N(x)| \rightarrow 0, \quad \text{quand } N \rightarrow +\infty,$$

où ici $|\cdot|$ désigne le module. Il faut juste comprendre ici que la convergence au sens de la norme quadratique est plus « faible » que la convergence ponctuelle. Nous allons revenir en détail dans la suite sur cette question de convergence ponctuelle.

2.3 Convergence des séries de Fourier

Dans cette section nous allons étudier deux types de convergence des séries de Fourier.

2.3.1 Convergence ponctuelle des séries de Fourier

Dans cette partie nous voulons étudier sous quelles conditions sur $f \in \mathcal{H}$ il est possible d'écrire

$$f(t) = \sum_{n \in \mathbb{Z}} c_n(f) e^{2i\pi n \frac{t}{a}}, \quad \text{pour } t \in \mathbb{R}. \quad (2.11)$$

Répetons, encore une fois, que pour une fonction $f \in \mathcal{H}$ (continue par morceaux sur $\mathbb{R}$ et a -périodique) le Théorème 16 implique uniquement la convergence de la suite $(f_N)_{N \geq 1}$ des sommes partielles au sens de la moyenne quadratique vers f . Cette notion de convergence n'implique pas nécessairement la convergence ponctuelle de la suite $(f_N(t))_{N \geq 1}$ vers $f(t)$ pour $t \in \mathbb{R}$. Ainsi, il faut a priori supposer un peu plus de régularité sur la fonction f pour avoir une chance de pouvoir écrire (2.11).

Vocabulaire 20. *Pour une fonction donnée $f \in \mathcal{H}$ et un point $t \in \mathbb{R}$, si $f(t)$ peut s'exprimer à l'aide de la somme « infinie »*

$$\sum_{n \in \mathbb{Z}} c_n(f) e^{2i\pi n \frac{t}{a}},$$

on dira que f est développable en sa série de Fourier en t . Ainsi, on dira que f est développable en sa série de Fourier sur $\mathbb{R}$ si $f(t)$ peut s'exprimer via cette somme en tout point t de $\mathbb{R}$.

Pour parler rigoureusement de la notion de somme « infinie » (i.e. la somme apparaissant dans (2.11)), il convient d'aborder la notion de suite (déjà vue dans votre cursus) et la notion de série (peut-être moins connue). Nous renvoyons à la Section 2.3.2 pour une présentation détaillée de ces concepts, mais sans perdre de temps nous pouvons énoncer et démontrer un théorème, dû à Dirichlet, permettant de répondre à la question de la convergence ponctuelle des séries de Fourier.

Théorème 21 (Théorème de Dirichlet). *Soit $f \in \mathcal{H}$. Si f est dérivable par morceaux, alors pour tout $x_0 \in \mathbb{R}$ on a*

$$f_N(x_0) \rightarrow \frac{1}{2} (f(x_0^+) + f(x_0^-)) \quad \text{quand } N \rightarrow +\infty.$$

De plus si f est continue en x_0 alors $f_N(x_0)$ tend vers $f(x_0)$ quand $N \rightarrow +\infty$.

Remarque 22. *L'important théorème de Dirichlet amène plusieurs commentaires.*

- *Il est important de comprendre que si f est continue en $x_0 \in \mathbb{R}$ (et dérivable par morceaux sur $\mathbb{R}$) alors f est développable en sa série de Fourier en x_0 . Donc, si f est continue sur $\mathbb{R}$ (et dérivable par morceaux) f est développable en sa série de Fourier sur $\mathbb{R}$.*
- *On peut s'interroger sur la nécessité des hypothèses du théorème de Dirichlet. En effet, ce résultat n'établit en aucun cas des conditions nécessaires et suffisantes pour qu'une fonction soit développable en série de Fourier. Il est par exemple possible de construire des fonctions périodiques, continues et non dérivables par morceaux tout de même développables en série de Fourier (voir médian P26). Par ailleurs, Du Bois-Reymond (1876) puis Leopold Fejér (1911) ont donné deux exemples de fonctions continues dont les séries de Fourier divergent en 0^2 .*
- *Le théorème de Dirichlet permet d'expliquer le résultat obtenu dans l'Exemple 15. En effet nous avons remarqué que $\lim_{N \rightarrow +\infty} f_N(\pi) = 0$ et $f(\pi) = -1$. Il n'y avait donc pas égalité. Par contre on a bien*

$$\lim_{N \rightarrow +\infty} f_N(\pi) = 0 = \frac{f(\pi^+) + f(\pi^-)}{2} = \frac{-1 + 1}{2} = 0.$$

Méthode. Il est important de comprendre et de retenir que dans la suite du cours à chaque fois que la question de la convergence simple (ou ponctuelle) de la série de Fourier de f est posée, il faut étudier si les hypothèses de régularité sur la fonction f du théorème de Dirichlet sont vérifiées ou non. Par exemple si $f \in \mathcal{H}$ avec f dérivable par morceaux alors

$$\begin{aligned} \sum_{n \in \mathbb{Z}} c_n(f) e^{2i\pi x_0 \frac{n}{a}} &= f(x_0) \quad \text{si } f \text{ est continue en } x_0, \\ \sum_{n \in \mathbb{Z}} c_n(f) e^{2i\pi x_0 \frac{n}{a}} &= \frac{f(x_0^+) + f(x_0^-)}{2} \quad \text{si } f \text{ n'est pas continue en } x_0. \end{aligned}$$

En particulier, si f est continue et dérivable par morceaux sur $\mathbb{R}$ on peut affirmer que f est développable en sa série de Fourier sur $\mathbb{R}$.

Démonstration du théorème de Dirichlet. Pour ce faire considérons la somme partielle de Fourier de f en $x_0 \in \mathbb{R}$, on a

$$\begin{aligned} f_N(x_0) &= \sum_{n=-N}^N c_n(f) e_n(x_0) = \frac{1}{a} \sum_{n=-N}^N \left(\int_{-a/2}^{a/2} f(x) e^{-2i\pi n x/a} dx \right) e^{2i\pi n x_0/a} \\ &= \frac{1}{a} \int_{-a/2}^{a/2} \left(\sum_{n=-N}^N e^{2i\pi n(x_0-x)/a} \right) f(x) dx, \end{aligned} \quad (2.12)$$

2. Pour les étudiants motivés, je renvoie à la page personnelle de Robert Rolland pour une présentation de l'exemple de L. Fejér.

où pour la définition des coefficients $c_n(f)$ on a utilisé la Proposition 34. On remarque également que pour tout $t \in \mathbb{R}$ on a

$$\begin{aligned}
 \sum_{n=-N}^N e^{2i\pi nt} &= e^{-2i\pi Nt} \sum_{n=0}^{2N} e^{2i\pi nt} = e^{-2i\pi Nt} \sum_{n=0}^{2N} (e^{2i\pi t})^n \\
 &= e^{-2i\pi Nt} \frac{1 - e^{2i\pi(2N+1)t}}{1 - e^{2i\pi t}} \\
 &= e^{-2i\pi Nt} \frac{e^{i\pi(2N+1)t} e^{-i\pi(2N+1)t} - e^{i\pi(2N+1)t}}{e^{i\pi t} e^{-i\pi t} - e^{i\pi t}} \\
 &= \frac{-2i \sin((2N+1)\pi t)}{-2i \sin(\pi t)} \\
 &= \frac{\sin((2N+1)\pi t)}{\sin(\pi t)}.
 \end{aligned}$$

En appliquant cette formule à (2.12) en le point $t = (x_0 - x)/a$, on obtient

$$f_N(x_0) = \frac{1}{a} \int_{-a/2}^{a/2} \frac{\sin((2N+1)\pi(x_0 - x)/a)}{\sin(\pi(x_0 - x)/a)} f(x) dx.$$

On considère alors le changement de variable suivant $x = x_0 + y$, on a

$$f_N(x_0) = \frac{1}{a} \int_{-x_0-a/2}^{-x_0+a/2} \frac{\sin((2N+1)\pi y/a)}{\sin(\pi y/a)} f(y + x_0) dy.$$

Par ailleurs, en utilisant les formules trigonométriques on a

$$\frac{\sin\left(\frac{(2N+1)\pi y}{a} + (2N+1)\pi\right)}{\sin\left(\frac{\pi y}{a} + \pi\right)} = \frac{-\sin\left(\frac{(2N+1)\pi y}{a}\right)}{-\sin\left(\frac{\pi y}{a}\right)} = \frac{\sin\left(\frac{(2N+1)\pi y}{a}\right)}{\sin\left(\frac{\pi y}{a}\right)},$$

et comme la fonction f est a -périodique alors d'après la Proposition 34

$$\begin{aligned}
 f_N(x_0) &= \frac{1}{a} \int_{-a/2}^{a/2} \frac{\sin((2N+1)\pi y/a)}{\sin(\pi y/a)} f(y + x_0) dy \\
 &= \frac{1}{a} \left(\int_0^{a/2} \frac{\sin((2N+1)\pi y/a)}{\sin(\pi y/a)} f(y + x_0) dy + \int_{-a/2}^0 \frac{\sin((2N+1)\pi y/a)}{\sin(\pi y/a)} f(y + x_0) dy \right).
 \end{aligned}$$

On utilise le changement de variable $z = -y$ dans la deuxième intégrale et on a

$$f_N(x_0) = \frac{1}{a} \left(\int_0^{a/2} \frac{\sin((2N+1)\pi y/a)}{\sin(\pi y/a)} f(y + x_0) dy + \int_0^{a/2} \frac{\sin((2N+1)\pi z/a)}{\sin(\pi z/a)} f(x_0 - z) dz \right),$$

d'où

$$f_N(x_0) = \frac{1}{a} \int_0^{a/2} (f(x_0 + x) + f(x_0 - x)) \frac{\sin((2N+1)\pi x/a)}{\sin(\pi x/a)} dx.$$

On remarque que si $f(x) = 1$ pour tout $x \in \mathbb{R}$ alors $c_0(f) = 1$ et $c_n(f) = 0$ pour tout $n \in \mathbb{N}^*$ et en particulier $f_N(x) = 1$ (voir également l'Exemple 14). On en déduit alors la formule

$$\frac{1}{2} = \frac{1}{a} \int_0^{a/2} \frac{\sin((2N+1)\pi x/a)}{\sin(\pi x/a)} dx.$$

On pose alors $y_0 = (f(x_0^+) + f(x_0^-))/2$ et on obtient

$$f_N(x_0) - y_0 = \frac{1}{a} \int_0^{a/2} \left(f(x_0 + x) - f(x_0^+) + f(x_0 - x) - f(x_0^-) \right) \frac{\sin((2N+1)\pi x/a)}{\sin(\pi x/a)} dx.$$

Soit φ la fonction définie sur $[0, a/2]$ par

$$\varphi(x) = \frac{f(x_0 + x) - f(x_0^+) + f(x_0 - x) - f(x_0^-)}{x} \frac{x}{\sin(\pi x/a)}.$$

Comme f est supposée dérivable par morceaux alors

$$\lim_{x \rightarrow 0^+} \frac{f(x_0 + x) - f(x_0^+) + f(x_0 - x) - f(x_0^-)}{x} = f'(x_0^+) + f'(x_0^-),$$

et cette même fonction est continue par morceaux pour tout $x \in]0, a/2]$. De plus, la fonction $x \mapsto x/\sin(\pi x/a)$ est continue sur $]0, a/2]$ et un développement limité de cette fonction au voisinage de 0 montre que $\lim_{x \rightarrow 0^+} x/\sin(\pi x/a) = a/\pi$ donc la fonction est continue sur tout $[0, a/2]$. Par conséquent φ est continue par morceaux. On a donc

$$f_N(x_0) - y_0 = \frac{1}{a} \int_0^{a/2} \varphi(x) \sin((2N+1)\pi x/a) dx,$$

et le Théorème de Riemann-Lebesgue (voir Théorème 23) permet de conclure que

$$\lim_{N \rightarrow +\infty} (f_N(x_0) - y_0) = 0.$$

Ce qui conclut la preuve du résultat. □

Théorème 23 (Théorème de Riemann-Lebesgue). *Soit f une fonction continue par morceaux sur un intervalle $[a, b]$ de $\mathbb{R}$. Alors on a*

$$\int_a^b f(x) e^{2i\pi n x} dx \rightarrow 0 \quad \text{quand } |n| \rightarrow +\infty.$$

En particulier, on a aussi

$$\begin{aligned} \int_a^b f(x) \cos(2\pi n x) dx &\rightarrow 0 \quad \text{quand } |n| \rightarrow +\infty, \\ \int_a^b f(x) \sin(2\pi n x) dx &\rightarrow 0 \quad \text{quand } |n| \rightarrow +\infty. \end{aligned}$$

Démonstration. Voir Section 2.5.3. □

L'objectif des deux sections suivantes est d'étudier une notion de convergence plus forte des séries de Fourier. Pour cela, et comme annoncé plus haut, nous avons besoin de revenir sur les notions de suite et de série.

2.3.2 Suites et séries numériques et suites de fonctions

Définition 11. Une suite $(u_n)_{n \in \mathbb{N}}$ de nombres réels ou complexes est dite convergente si il existe un $\ell \in \mathbb{C}$ tel que

$$\forall \varepsilon > 0, \exists n_0 \in \mathbb{N} : (\forall n \geq n_0 \Rightarrow |u_n - \ell| < \varepsilon)$$

On dit que ℓ est la limite de la suite $(u_n)_{n \in \mathbb{N}}$. On dit encore que la suite $(u_n)_{n \in \mathbb{N}}$ converge vers ℓ et on note $\lim_{n \rightarrow +\infty} u_n = \ell$.

Exemple 16. Soit $u_n = 1/n$ pour tout $n \geq 1$. Alors la suite $(u_n)_{n \geq 1}$ est convergente avec $\lim_{n \rightarrow +\infty} u_n = 0$. Remarquons que toute suite n'est pas nécessairement convergente. En effet la suite $(n)_{n \in \mathbb{N}}$ n'est pas convergente puisque $\lim_{n \rightarrow +\infty} n = +\infty$. De même la suite $(\cos(n\pi))_{n \in \mathbb{N}}$ n'est pas convergente car si n est pair alors $\cos(n\pi) = 1$ et si n est impair alors $\cos(n\pi) = -1$. Dans les deux cas précédents on dit que la suite n'est pas convergente ou que la suite est divergente.

Proposition 24. Pour les suites convergentes on a les propriétés suivantes :

- Si la suite $(u_n)_{n \in \mathbb{N}}$ est convergente, sa limite est unique.
- Toute suite convergente est bornée, mais la réciproque est fausse.

Démonstration. Voir Section 2.5.2 □

Exemple 17. Comme vu précédemment la suite $(\cos(n\pi))_{n \in \mathbb{N}}$ n'est pas convergente. Cependant, pour tout $n \in \mathbb{N}$, on a $|\cos(n\pi)| \leq 1$.

Théorème 25. Toute suite croissante de réels $(u_n)_{n \in \mathbb{N}}$ est convergente si et seulement si elle est majorée. De même toute suite décroissante de réels $(u_n)_{n \in \mathbb{N}}$ est convergente si et seulement si elle est minorée.

Démonstration. Voir Section 2.5.2 □

Exemple 18. La suite $(1/n^2)_{n \geq 1}$ est décroissante et minorée par 0, elle est donc convergente.

Définition 12. Soit $(u_n)_{n \in \mathbb{N}}$ une suite de nombres réels ou complexes et $(S_n)_{n \in \mathbb{N}}$ la suite définie par

$$S_n = \sum_{k=0}^n u_k, \quad \forall n \in \mathbb{N}.$$

La suite de couples $((u_n, S_n))_{n \in \mathbb{N}}$ est appelée série de terme général u_n et de somme partielle S_n et est notée $\sum u_n$. On dit que la série $\sum u_n$ converge si et seulement si la suite des sommes partielles $(S_n)_{n \in \mathbb{N}}$ converge au sens de la Définition 11. Dans ce cas la limite S de $(S_n)_{n \in \mathbb{N}}$ est appelée la somme de la série, on la note

$$S = \sum_{k=0}^{+\infty} u_k.$$

Enfin si la suite $(S_n)_{n \in \mathbb{N}}$ diverge, on dira que la série de terme général u_n est divergente.

Exemple 19. Voici deux exemples de séries convergentes à connaître :

- Pour $r \in \mathbb{R}$ considérons la série, dite géométrique, $\sum r^n$. On note $(S_n)_{n \in \mathbb{N}}$ la suite des sommes partielles associées à la série, on a

$$S_n = \sum_{k=0}^n r^k = \frac{1 - r^{n+1}}{1 - r}.$$

Ainsi si $|r| < 1$ alors la suite $(S_n)_{n \in \mathbb{N}}$ est convergente avec pour limite

$$S = \sum_{k=0}^{+\infty} r^k = \lim_{n \rightarrow +\infty} S_n = \frac{1}{1 - r}.$$

Par contre si $|r| \geq 1$ alors la série $\sum r^n$ est divergente.

- L'autre exemple est donné par les séries de Riemann. On démontrera plus tard dans le cours le résultat suivant :

la série $\sum 1/n^\alpha$ est convergente si et seulement si $\alpha > 1$.

On renvoie à l'Exemple 26 pour le calcul de $\sum 1/n^2$.

Proposition 26. Une condition nécessaire pour que la série de terme général u_n converge est $\lim_{n \rightarrow +\infty} u_n = 0$.

Démonstration. On remarque que pour tout $N \geq 1$

$$S_{N+1} - S_N = \sum_{n=0}^{N+1} u_n - \sum_{n=0}^N u_n = u_{N+1}.$$

Or la série est supposée convergente donc $S_{N+1} - S_N \rightarrow 0$ quand $N \rightarrow +\infty$. □

Remarque 27. La condition de la précédente proposition est nécessaire mais pas suffisante. Par exemple la série de Riemann $\sum 1/n$ est divergente même si $\lim_{n \rightarrow +\infty} 1/n = 0$.

Il est important de comprendre que pour montrer qu'une série converge, il faut montrer que la suite de ses sommes partielles converge. En particulier, on peut appliquer les résultats de convergence des suites rappelés précédemment. À titre d'exemple on a le résultat suivant :

Proposition 28. Soit $f \in \mathcal{H}$, alors la série $\sum |c_n(f)|^2$, où $c_n(f)$ désigne le coefficient de Fourier de f d'ordre n , est convergente. En particulier, on a

$$c_n(f) \rightarrow 0 \quad \text{quand } |n| \rightarrow +\infty.$$

Démonstration. D'après l'inégalité de Bessel, on a pour tout $N \in \mathbb{N}$

$$S_N = \sum_{n=-N}^N |c_n(f)|^2 \leq \frac{1}{a} \int_0^a |f(t)|^2 dt.$$

Donc la suite des sommes partielles $(S_N)_{N \in \mathbb{N}}$ est une suite croissante de réels majorée. Ainsi d'après le Théorème 25 la suite $(S_N)_{N \in \mathbb{N}}$ est une suite convergente. Donc la série $\sum |c_n(f)|^2$ est convergente. On déduit de la Proposition 26 que

$$c_n(f) \rightarrow 0 \quad \text{quand } |n| \rightarrow +\infty.$$

Ce qui achève la preuve du résultat. □

Définition 13. On dit que la série $\sum u_n$ converge absolument (ou est absolument convergente) si la série de terme général $|u_n|$ converge. Autrement dit $\sum u_n$ est absolument convergente si $\sum |u_n|$ est convergente

Exemple 20. Soit $r \in [0, 1[$, on considère la série $\sum (-1)^n r^n$. Cette série est absolument convergente car pour tout $n \in \mathbb{N}$, on a $|(-1)^n r^n| = r^n$. Donc comme la série $\sum |(-1)^n r^n|$ est convergente alors la série $\sum (-1)^n r^n$ est absolument convergente.

Proposition 29. On a les propriétés suivantes :

- Si la série de terme général u_n converge absolument, alors la série $\sum u_n$ est convergente et on a l'inégalité suivante

$$\left| \sum_{n=0}^{+\infty} u_n \right| \leq \sum_{n=0}^{+\infty} |u_n|.$$

- Soient u_n et v_n les termes généraux de deux séries. Si la série $\sum v_n$ est convergente et que

$$|u_n| \leq v_n, \quad \forall n \in \mathbb{N},$$

alors la série $\sum u_n$ est absolument convergente.

Démonstration. Admise. □

Exemple 21. On déduit de la question précédente et de l'Exemple 20 que la série $\sum (-1)^n r^n$ est convergente. De même toujours pour $r \in [0, 1[$, on considère la série $\sum r^n \cos(rn)$. Cette série est convergente car pour tout $n \in \mathbb{N}$ on a

$$|r^n \cos(rn)| \leq r^n.$$

Or la série géométrique $\sum r^n$ est convergente donc la série $\sum r^n \cos(rn)$ est absolument convergente et donc convergente.

On va maintenant considérer le cas des suites de fonctions. On va ainsi remplacer les suites de nombres u_n par des suites de fonctions f_n .

Définition 14. Soient I un intervalle non vide de $\mathbb{R}$, et $(f_n)_{n \in \mathbb{N}}$ une suite de fonctions avec $f_n : I \rightarrow \mathbb{R}$ ou $\mathbb{C}$. On dit que la suite $(f_n)_{n \in \mathbb{N}}$ converge simplement (ou ponctuellement) vers f (une fonction) sur I , si la suite numérique $(f_n(x))_{n \in \mathbb{N}}$ converge vers $f(x)$ quand $n \rightarrow +\infty$ et ce pour tout $x \in I$. Autrement dit, on a

$$\forall x \in I, \left[\forall \varepsilon > 0, \exists n_0 \in \mathbb{N} : (\forall n \geq n_0 \Rightarrow |f_n(x) - f(x)| < \varepsilon) \right].$$

Exemple 22. Soit $f_n : [0, 1] \rightarrow \mathbb{R}$ avec $f_n(x) = x^n$. Alors pour tout $x \in [0, 1[$, on a $f_n(x) \rightarrow 0$ quand $n \rightarrow +\infty$. Pour $x = 1$, on a $f_n(1) = 1 \rightarrow 1$ quand $n \rightarrow +\infty$. Ainsi la suite de fonctions $(f_n)_{n \in \mathbb{N}}$ converge simplement vers f sur $[0, 1]$ avec

$$f(x) = \begin{cases} 0 & \text{si } x \in [0, 1[, \\ 1 & \text{si } x = 1. \end{cases}$$

Définition 15. Soient I un intervalle non vide de $\mathbb{R}$, et $(f_n)_{n \in \mathbb{N}}$ une suite de fonctions avec $f_n : I \rightarrow \mathbb{R}$ ou $\mathbb{C}$. On dit que la suite $(f_n)_{n \in \mathbb{N}}$ converge uniformément vers f (une fonction) sur I , si

$$\lim_{n \rightarrow +\infty} \sup_{x \in I} |f_n(x) - f(x)| = 0$$

Autrement dit, on a

$$\forall \varepsilon > 0, \exists n_0 \in \mathbb{N} : (\forall n \geq n_0 \text{ et } \forall x \in I \Rightarrow |f_n(x) - f(x)| < \varepsilon).$$

Exemple 23. Soit $I = [0, 1]$, on considère la suite de fonctions $(f_n)_{n \in \mathbb{N}}$ avec $f_n(x) = x/(x+n)$. On remarque que pour chaque $x \in I$, on a

$$f_n(x) = \frac{x}{x+n} \rightarrow 0 \quad \text{quand } n \rightarrow +\infty.$$

En notant f la fonction identiquement nulle sur I , on obtient

$$\sup_{x \in I} |f_n(x) - f(x)| = \sup_{x \in I} \left| \frac{x}{x+n} \right| = \frac{1}{1+n} \rightarrow 0 \quad \text{quand } n \rightarrow +\infty.$$

Donc $(f_n)_{n \in \mathbb{N}}$ converge uniformément vers f .

Il convient de faire attention entre la notion de convergence simple et de convergence uniforme, la différence semble subtile mais très importante. Dans la définition de la convergence simple, pour chaque $x \in I$ il existe un rang $n_0 \in \mathbb{N}$ tel que pour tout $n \geq n_0$ on ait $|f_n(x) - f(x)| < \varepsilon$. Ce rang n_0 peut dépendre de x . À l'inverse pour la convergence uniforme il existe un rang $n_0 \in \mathbb{N}$ tel que pour tout $n \geq n_0$ et pour tout $x \in I$ on ait $|f_n(x) - f(x)| < \varepsilon$. Ainsi, dans le cas de la convergence uniforme le rang n_0 **ne dépend pas de x** . Il faut comprendre que dans le cas de la convergence simple, les suites $(f_n(x))_{n \in \mathbb{N}}$ convergent avec des « vitesses » potentiellement différentes vers $f(x)$, cette « vitesse » ne dépendant que du point x . Par contre dans le cas de la convergence uniforme les suites $(f_n(x))_{n \in \mathbb{N}}$ convergent toutes à la même « vitesse » vers $f(x)$.

Proposition 30. Soient I un intervalle non vide de $\mathbb{R}$, et $(f_n)_{n \in \mathbb{N}}$ une suite de fonctions avec $f_n : I \rightarrow \mathbb{R}$ ou $\mathbb{C}$. Si $(f_n)_{n \in \mathbb{N}}$ converge uniformément vers f sur I alors $(f_n)_{n \in \mathbb{N}}$ converge simplement vers f sur I . La réciproque est par contre fausse.

Démonstration. Il suffit de remarquer que pour tout $x \in I$, on a

$$|f_n(x) - f(x)| \leq \sup_{x \in I} |f_n(x) - f(x)|.$$

Ce qui conclut la preuve du résultat. $\square$

Proposition 31. Voici deux propriétés importantes de la convergence uniforme.

- Soit $(f_n)_{n \in \mathbb{N}}$ une suite de fonctions définies et continues sur un intervalle I de $\mathbb{R}$. Si $(f_n)_{n \in \mathbb{N}}$ converge uniformément vers une fonction f sur I , alors, f est une fonction continue sur I .
- Soit $(f_n)_{n \in \mathbb{N}}$ une suite de fonctions définies et C^1 sur un intervalle I de $\mathbb{R}$. Si pour un point $x_0 \in I$ la suite de nombres $(f_n(x_0))_{n \in \mathbb{N}}$ est convergente et que la suite $(f'_n)_{n \in \mathbb{N}}$ converge uniformément sur I . Alors la limite f de $(f_n)_{n \in \mathbb{N}}$ est dérivable sur I et on a

$$f'(x) = \lim_{n \rightarrow +\infty} f'_n(x) \quad \forall x \in I.$$

Démonstration. Admise. $\square$

2.3.3 Convergence normale des séries de Fourier

On va maintenant s'intéresser aux séries de fonctions. On remplace les suites de nombres u_n par des suites de fonctions f_n dans le terme général des séries. Comme dans le cas des séries numériques, la convergence des séries de fonctions est obtenue par la convergence des sommes partielles. Ces sommes partielles étant vues comme des suites de fonctions. En particulier, les notions précédentes de convergence simple, convergence uniformes ainsi que les propriétés énoncées s'appliquent directement aux séries de fonctions via les sommes partielles. Pour les séries de fonctions nous introduisons une dernière notion propre aux séries de fonctions.

Définition 16. On dit qu'une série de fonctions $\sum f_n$, où f_n est une fonction définie sur un intervalle I de $\mathbb{R}$, converge normalement sur I , si il existe une série numérique $\sum \alpha_n$ (avec $\alpha_n \geq 0$ pour tout n) convergente et vérifiant

$$|f_n(x)| \leq \alpha_n, \quad \forall x \in I, \forall n \in \mathbb{N}.$$

Exemple 24. Pour tout $n \geq 1$, soit $f_n : \mathbb{R} \rightarrow \mathbb{R}$ une fonction définie par $f_n(x) = \sin(x)/n^2$. On remarque que pour tout $n \geq 1$ on a

$$|f_n(x)| \leq \frac{1}{n^2}, \quad \forall x \in \mathbb{R}.$$

la série (numérique) $\sum 1/n^2$ étant une série de Riemann convergente, on en déduit que la série de fonctions $\sum f_n$ est normalement convergente sur $\mathbb{R}$.

Théorème 32. *Si la série de fonctions $\sum f_n$ définies sur I un intervalle de $\mathbb{R}$ converge normalement sur I alors elle converge uniformément et donc simplement sur I .*

Démonstration. Admise. □

Théorème 33 (Corollaire du théorème de Dirichlet). *Soit f une fonction a -périodique, continue et C^1 par morceaux. Alors on a les propriétés suivantes :*

1. *La fonction f est développable en sa série de Fourier sur $\mathbb{R}$. Autrement dit, pour tout $x \in \mathbb{R}$, on a*

$$f(x) = \sum_{n \in \mathbb{Z}} c_n(f) e^{2i\pi n \frac{x}{a}}.$$

2. *La série de Fourier de f converge normalement (et donc absolument et uniformément) sur $\mathbb{R}$ vers f .*

Démonstration. Comme f est continue et C^1 par morceaux sur $\mathbb{R}$ on déduit directement du théorème de Dirichlet, voir Théorème 21, que f est développable en sa série de Fourier sur $\mathbb{R}$. Pour le deuxième point du résultat, supposons dans un premier temps que f est de classe C^1 sur $\mathbb{R}$ et montrons que sa série de Fourier

$$\sum c_n(f) \exp\left(2i\pi n \frac{t}{a}\right),$$

converge normalement sur $\mathbb{R}$. Un fois la convergence normale de la série de Fourier de f obtenue on déduit du premier point que cette série converge normalement vers f sur $\mathbb{R}$. Remarquons dans un premier temps que pour tout $t \in \mathbb{R}$ la majoration suivante est vérifiée :

$$\left| c_n(f) \exp\left(-2i\pi n \frac{t}{a}\right) \right| \leq |c_n(f)|, \quad \forall n \in \mathbb{Z}.$$

De plus, pour tout $n \in \mathbb{Z}^*$, en utilisant une IPP (possible car f est C^1) on a

$$\begin{aligned} c_n(f) &= \frac{1}{a} \int_0^a f(t) \exp\left(-2i\pi n \frac{t}{a}\right) dt \\ &= \frac{1}{a} \left(\left[-a \frac{f(t) \exp(-2i\pi n t/a)}{2i\pi n} \right]_0^a + \frac{a}{2i\pi n} \int_0^a f'(t) \exp\left(-2i\pi n \frac{t}{a}\right) dt \right) \\ &= \frac{1}{2i\pi n} ((-f(a) + f(0)) + a c_n(f')). \end{aligned}$$

Or f est C^1 et a -périodique donc $f(a) = f(0)$. D'où

$$c_n(f) = \frac{a}{2i\pi n} c_n(f') \Rightarrow |c_n(f)| = \frac{a}{2\pi |n|} |c_n(f')|, \quad \forall n \in \mathbb{Z}^*$$

Maintenant, il reste à utiliser l'inégalité, dite de Young, $bc \leq (b^2 + c^2)/2$ pour tout $b, c \geq 0$. Pour démontrer cette inégalité, il faut montrer que pour $b, c \geq 0$, on a $(b^2 + c^2)/2 - bc \geq 0$. Or on a facilement

$$\frac{b^2 + c^2}{2} - bc = \frac{1}{2}(b^2 - 2bc + c^2) = \frac{1}{2}(b - c)^2 \geq 0.$$

En appliquant à présent cette inégalité avec $b = 1/|n|$ et $c = |c_n(f')|$, on obtient

$$|c_n(f)| = \frac{a}{2\pi|n|}|c_n(f')| \leq \frac{a}{4\pi} \left(\frac{1}{n^2} + |c_n(f')|^2 \right), \quad \forall n \in \mathbb{Z}^*.$$

On en déduit que la série $\sum |c_n(f)|$ est convergente car la série de Riemann $\sum 1/n^2$ est convergente et d'après la Proposition 28 la série $\sum |c_n(f')|^2$ est également convergente. On en conclut que la série de Fourier de f converge normalement sur $\mathbb{R}$. Si à présent on suppose que f est uniquement continue et C^1 par morceaux sur $\mathbb{R}$, on montre que la formule

$$c_n(f) = \frac{a}{2i\pi n} c_n(f') \quad \forall n \in \mathbb{Z}^*,$$

est toujours valable (voir TD chapitre 2) et le reste de la preuve est identique. $\square$

2.4 Propriétés supplémentaires et exemples de calcul

Dans cette section on donne deux résultats utiles pour le calcul des coefficients de Fourier d'une fonction ainsi que deux exemples.

Proposition 34. *Soit $f \in \mathcal{H}$. Alors l'intégrale de f est constante sur tout intervalle de longueur a , i.e.,*

$$\int_{\alpha}^{\alpha+a} f(x) dx = \int_0^a f(x) dx, \quad \forall \alpha \in \mathbb{R}.$$

Démonstration. Soit $\alpha \in \mathbb{R}$, alors on a

$$\int_{\alpha}^{\alpha+a} f(x) dx = \int_{\alpha}^0 f(x) dx + \int_0^a f(x) dx + \int_a^{\alpha+a} f(x) dx.$$

En utilisant le changement de variable $y = x - a$ dans la dernière intégrale on obtient

$$\int_{\alpha}^{\alpha+a} f(x) dx = \int_{\alpha}^0 f(x) dx + \int_0^a f(x) dx + \int_0^{\alpha} f(y+a) dy.$$

Or f est a -périodique donc $f(x+a) = f(x)$ et

$$\int_{\alpha}^{\alpha+a} f(x) dx = \int_{\alpha}^0 f(x) dx + \int_0^a f(x) dx + \int_0^{\alpha} f(x) dx$$

$$= \int_{\alpha}^0 f(x) dx + \int_0^a f(x) dx - \int_{\alpha}^0 f(x) dx.$$

D'où

$$\int_{\alpha}^{\alpha+a} f(x) dx = \int_0^a f(x) dx.$$

Ce qui conclut la preuve du résultat. □

Proposition 35. Soit $f \in \mathcal{H}$. Alors on a

- Si f est paire alors $b_n(f) = 0$ pour tout $n \geq 1$.
- Si f est impaire alors $a_n(f) = 0$ pour tout $n \in \mathbb{N}$.

Démonstration. Supposons dans un premier temps f paire. On a d'après le résultat précédent pour tout $n \geq 1$

$$\begin{aligned} b_n(f) &= \frac{2}{a} \int_{-a/2}^{a/2} f(x) \sin\left(2\pi n \frac{x}{a}\right) dx \\ &= \frac{2}{a} \int_{-a/2}^0 f(x) \sin\left(2\pi n \frac{x}{a}\right) dx + \frac{2}{a} \int_0^{a/2} f(x) \sin\left(2\pi n \frac{x}{a}\right) dx. \end{aligned}$$

En utilisant dans la première intégrale le changement de variable $y = -x$, la parité de f et l'imparité du sinus, on obtient

$$\begin{aligned} b_n(f) &= \frac{2}{a} \int_{a/2}^0 f(-y) \sin\left(-2\pi n \frac{y}{a}\right) (-dy) + \frac{2}{a} \int_0^{a/2} f(x) \sin\left(2\pi n \frac{x}{a}\right) dx \\ &= \frac{2}{a} \int_0^{a/2} f(y) \sin\left(-2\pi n \frac{y}{a}\right) dy + \frac{2}{a} \int_0^{a/2} f(x) \sin\left(2\pi n \frac{x}{a}\right) dx \\ &= -\frac{2}{a} \int_0^{a/2} f(x) \sin\left(2\pi n \frac{x}{a}\right) dx + \frac{2}{a} \int_0^{a/2} f(x) \sin\left(2\pi n \frac{x}{a}\right) dx = 0. \end{aligned}$$

Si maintenant f est impaire alors pour tout $n \in \mathbb{N}$

$$\begin{aligned} a_n(f) &= \frac{2}{a} \int_{-a/2}^{a/2} f(x) \cos\left(2\pi n \frac{x}{a}\right) dx \\ &= \frac{2}{a} \int_{-a/2}^0 f(x) \cos\left(2\pi n \frac{x}{a}\right) dx + \frac{2}{a} \int_0^{a/2} f(x) \cos\left(2\pi n \frac{x}{a}\right) dx. \end{aligned}$$

Comme dans le cas précédent en utilisant dans la première intégrale le changement de variable $y = -x$, l'imparité de f et la parité du cosinus, on obtient

$$a_n(f) = \frac{2}{a} \int_{a/2}^0 f(-y) \cos\left(-2\pi n \frac{y}{a}\right) (-dy) + \frac{2}{a} \int_0^{a/2} f(x) \cos\left(2\pi n \frac{x}{a}\right) dx$$

$$\begin{aligned}
&= \frac{2}{a} \int_0^{a/2} f(-y) \cos\left(2\pi n \frac{y}{a}\right) dy + \frac{2}{a} \int_0^{a/2} f(x) \cos\left(2\pi n \frac{x}{a}\right) dx \\
&= -\frac{2}{a} \int_0^{a/2} f(x) \cos\left(2\pi n \frac{x}{a}\right) dx + \frac{2}{a} \int_0^{a/2} f(x) \cos\left(2\pi n \frac{x}{a}\right) dx = 0.
\end{aligned}$$

Ce qui conclut la preuve du résultat. $\square$

Exemple 25. Soit f la fonction 2-périodique avec $f(x) = |x|$ si $|x| < 1$. Comme f est paire alors on a

$$b_n(f) = 0, \quad \forall n \geq 1.$$

On a également

$$a_0(f) = \frac{2}{2} \int_{-1}^1 f(x) dx = - \int_{-1}^0 x dx + \int_0^1 x dx = - \left[\frac{x^2}{2} \right]_{-1}^0 + \left[\frac{x^2}{2} \right]_0^1 = \frac{1}{2} + \frac{1}{2} = 1,$$

et comme la fonction $x \mapsto f(x) \cos(\pi n x)$ est paire

$$\begin{aligned}
a_n(f) &= - \int_{-1}^0 x \cos(\pi n x) dx + \int_0^1 x \cos(\pi n x) dx \\
&= 2 \int_0^1 x \cos(\pi n x) dx \\
&= 2 \left[\frac{x \sin(\pi n x)}{\pi n} \right]_0^1 - 2 \int_0^1 \frac{\sin(\pi n x)}{\pi n} dx \\
&= -\frac{2}{\pi n} \left[\frac{-\cos(\pi n x)}{\pi n} \right]_0^1 \\
&= \frac{2}{(\pi n)^2} (\cos(\pi n) - 1) \\
&= \frac{2}{(\pi n)^2} ((-1)^n - 1).
\end{aligned}$$

Ainsi si n est pair $a_n(f) = 0$ et si n est impaire $a_n(f) = -4/(\pi n)^2$. On obtient alors l'expression de la somme partielle d'ordre N de la série de Fourier de f :

$$f_N(x) = \frac{1}{2} + \sum_{n=1}^N \frac{2}{(\pi n)^2} ((-1)^n - 1) \cos(\pi n x).$$

Par ailleurs comme f est continue sur $\mathbb{R}$ et C^1 par morceaux alors d'après le Théorème 33 la série de Fourier de f converge normalement sur $\mathbb{R}$ vers f et pour tout $x \in \mathbb{R}$ on a

$$f(x) = \frac{1}{2} + \sum_{n=1}^{+\infty} \frac{2}{(\pi n)^2} ((-1)^n - 1) \cos(\pi n x).$$

Exemple 26. Soit f la fonction a -périodique avec $f(x) = x/a$ pour tout $0 \leq x < a$. Ici la fonction n'est ni paire, ni impaire. On a

$$a_0(f) = \frac{2}{a} \int_0^a \frac{x}{a} dx = \frac{2}{a^2} \left[\frac{x^2}{2} \right]_0^a = 1,$$

et par IPP pour tout $n \geq 1$

$$\begin{aligned} a_n(f) &= \frac{2}{a} \int_0^a \frac{x}{a} \cos\left(2\pi n \frac{x}{a}\right) dx = \frac{2}{a^2} \underbrace{\left[a \frac{x \sin\left(2\pi n \frac{x}{a}\right)}{2\pi n} \right]_0^a}_{=0} - \frac{1}{a\pi n} \int_0^a \sin\left(2\pi n \frac{x}{a}\right) dx \\ &= \frac{1}{2(\pi n)^2} \left[\cos\left(2\pi n \frac{x}{a}\right) \right]_0^a = 0. \end{aligned}$$

De même pour tout $n \geq 1$ on a

$$\begin{aligned} b_n(f) &= \frac{2}{a} \int_0^a \frac{x}{a} \sin\left(2\pi n \frac{x}{a}\right) dx = \frac{1}{a\pi n} \left[-x \cos\left(2\pi n \frac{x}{a}\right) \right]_0^a + \frac{1}{a\pi n} \int_0^a \cos\left(2\pi n \frac{x}{a}\right) dx \\ &= -\frac{1}{\pi n} + \frac{1}{2(\pi n)^2} \left[\sin\left(2\pi n \frac{x}{a}\right) \right]_0^a \\ &= -\frac{1}{\pi n}. \end{aligned}$$

On obtient alors l'expression de la somme partielle d'ordre N de la série de Fourier de f :

$$f_N(x) = \frac{1}{2} - \sum_{n=1}^N \frac{1}{\pi n} \sin\left(2\pi n \frac{x}{a}\right).$$

Ici f n'est pas continue, on a alors convergence simple de la suite $(f_N(x))_{N \in \mathbb{N}}$ vers $f(x)$ pour tout les points $x \in \mathbb{R} \setminus \{na, n \in \mathbb{Z}\}$ et pour tout $x \in \{na, n \in \mathbb{Z}\}$ on a

$$\lim_{N \rightarrow +\infty} f_N(na) = \frac{f(na^+) + f(na^-)}{2} = \frac{1}{2}.$$

De plus d'après l'égalité de Parseval (Corollaire 19) on a

$$\frac{1}{a} \int_0^a |f(x)|^2 dx = \frac{1}{4} + \frac{1}{2} \sum_{n=1}^{+\infty} \frac{1}{(\pi n)^2}.$$

Or

$$\frac{1}{a} \int_0^a |f(x)|^2 dx = \frac{1}{a^3} \int_0^a x^2 dx = \frac{1}{3}.$$

Donc

$$\sum_{n=1}^{+\infty} \frac{1}{n^2} = \frac{\pi^2}{6}.$$

Proposition 36. Soit f une fonction a -périodique de classe C^p alors pour $1 \leq n \leq p$ on a

$$c_k(f^{(n)}) = \left(\frac{2i\pi k}{a}\right)^n c_k(f), \quad \forall k \in \mathbb{Z}. \quad (2.13)$$

En particulier si f est C^p alors $\lim_{|k| \rightarrow +\infty} k^p c_k(f) = 0$.

Démonstration. Voir le TD du chapitre 2 pour l'égalité (2.13). Pour le deuxième point on remarque que

$$k^p c_k(f) = \left(\frac{a}{2i\pi}\right)^p c_k(f^{(p)}) \quad \forall k \in \mathbb{Z}.$$

La fonction $f^{(p)}$ étant continue, on applique le théorème de Riemann-Lebesgue (voir Théorème 23) et on obtient $\lim_{|k| \rightarrow +\infty} c_k(f^{(p)}) = 0$. Ce qui conclut la preuve du résultat. $\square$

Le résultat suivant est une sorte de réciproque.

Proposition 37. Soit f une fonction a -périodique, si pour $p \geq 2$ la suite $(|k^p c_k(f)|)_{k \in \mathbb{Z}}$ est bornée alors f est de classe C^{p-2} .

Démonstration. Admise. $\square$

2.5 Compléments

Nous présentons dans cette partie des résultats supplémentaires pour les étudiantes et étudiants curieux.

2.5.1 développement d'une fonction sur une base orthogonale

La technique utilisée dans Section 2.2.2 afin d'approcher une fonction périodique peut se généraliser afin d'approcher de manière similaire d'autres types de fonctions. Pour cela, il ne faut plus considérer la famille des exponentielles complexes $(e_n)_{n \in \mathbb{Z}}$ mais une autre famille de polynômes orthogonaux (des fonctions simples).

De manière générale considérons $I =]a, b[$ ($a < b$) un intervalle ouvert borné de $\mathbb{R}$ et w une fonction continue sur $]a, b[$ à valeurs dans $]0, +\infty[$. Cette fonction w s'appelle un poids. On suppose que pour tout $n \in \mathbb{N}$ la fonction w vérifie

$$\int_a^b |x|^n w(t) dt < \infty.$$

On note alors $\mathcal{H}_w$ l'espace vectoriel (de dimension infinie) définit par

$$\mathcal{H}_w = \left\{ f :]a, b[\rightarrow \mathbb{C}, f \text{ est continue par morceaux et } \int_a^b |f(t)|^2 w(t) dt < \infty \right\}.$$

On munit $\mathcal{H}_w$ du produit scalaire suivant

$$\langle f, g \rangle_w = \int_a^b f(t)g(t)w(t)dt, \quad \forall f, g \in \mathcal{H}_w.$$

On en déduit l'expression de la norme sur $\mathcal{H}_w$ via

$$\|f\|_w = \sqrt{\langle f, f \rangle_w} = \left(\int_a^b |f(t)|^2 w(t) dt \right)^{1/2}, \quad \forall f \in \mathcal{H}_w.$$

Comme dans le cas de l'espace $\mathcal{H}$, supposons disposer d'une famille $(\varphi_n)_{n \in \mathbb{N}}$ de polynômes orthogonaux, i.e., vérifiant

$$\langle \varphi_n, \varphi_m \rangle_w = \int_0^a \varphi_n(t) \varphi_m(t) w(t) dt = \begin{cases} 0 & \text{si } n \neq m, \\ \|\varphi_n\|_w^2 & \text{si } n = m. \end{cases}$$

Pour $f \in \mathcal{H}_w$ on cherche à déterminer l'expression du polynôme f_N réalisant la meilleure approximation de f dans le sous-espace vectoriel $V_N = \text{Vect}\langle \varphi_0, \dots, \varphi_N \rangle$. En réalisant des calculs similaires au cas des fonctions périodiques, on en déduit que ce polynôme f_N est donné par

$$f_N(t) = \sum_{n=0}^N c_n \varphi_n(t), \quad \forall t \in \mathbb{R},$$

avec

$$c_n = c_n(f) = \frac{\langle f, \varphi_n \rangle_w}{\|\varphi_n\|_w^2}, \quad n = 0, \dots, N.$$

En conclusion f_N vérifie

$$\|f - f_N\|_w = \min\{\|f - p\|_w, p \in V_N\}.$$

Enfin, en fonction de a , b et w on peut construire différentes familles de polynômes orthogonaux permettant d'approcher de manière très précise les fonctions appartenant à l'espace $\mathcal{H}_w$. À titre d'exemple on peut citer :

- $]a, b[=]-1, 1[, w(x) = 1, \varphi_n = \text{polynôme de Legendre}^3,$
- $]a, b[=]-1, 1[, w(x) = \frac{1}{\sqrt{1-x^2}}, \varphi_n = \text{polynôme de Tchebychev}.$

Dans le cas où $]a, b[$ est non borné un travail similaire peut être réalisé et on a les exemples célèbres suivants :

- $]a, b[=]0, +\infty[, w(x) = \exp(-x), \varphi_n = \text{polynôme de Laguerre},$
- $]a, b[=]-\infty, +\infty[, w(x) = \exp(-x^2), \varphi_n = \text{polynôme de Hermite}.$

Il est important de retenir que l'idée d'approcher des fonctions périodiques (objets complexes) par des polynômes (objets simples) se généralise naturellement et facilement à d'autres espaces fonctionnels (espaces vectoriels contenant des fonctions). Cette idée d'approcher des objets complexes par des objets plus simples est à la base de nombreuses méthodes en analyse numérique (mais aussi dans d'autres branches des mathématiques) afin d'estimer par ordinateur la valeur d'une intégrale, d'approcher (toujours sur ordinateur) les solutions d'équations différentielles etc.

3. À voir en TP.

2.5.2 Quelques preuves des résultats de la Section 2.3.2

Démonstration. (Proposition 24)

- Supposons l'existence de ℓ et ℓ' ($\ell \neq \ell'$). Soit $\varepsilon > 0$ quelconque, par définition de la convergence de la suite $(u_n)_{n \in \mathbb{N}}$ alors il existe $n_0 \in \mathbb{N}$ tel que pour tout $n \geq n_0$ on ait

$$|u_n - \ell| < \frac{\varepsilon}{2}.$$

De même il existe $n_1 \in \mathbb{N}$ tel que pour tout $n \geq n_1$

$$|u_n - \ell'| < \frac{\varepsilon}{2}.$$

Soit à présent $n_2 = \max(n_0, n_1)$, alors pour tout $n \geq n_2$, on a

$$|\ell - \ell'| \leq |\ell - u_n| + |u_n - \ell'| < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon.$$

Le réel ε étant quelconque, on en déduit que $\ell = \ell'$.

- Si $(u_n)_{n \in \mathbb{N}}$ est convergente, de limite ℓ , alors pour $\varepsilon > 0$ fixé, il existe $n_0 \in \mathbb{N}$ tel que pour tout $n \geq n_0$ on a

$$|u_n - \ell| < \varepsilon \Rightarrow |u_n| < \varepsilon + |\ell|.$$

Ainsi on obtient que

$$|u_n| \leq \max\{|u_0|, |u_1|, \dots, |u_{n_0-1}|, \varepsilon + |\ell|\}.$$

Par conséquent la suite $(u_n)_{n \in \mathbb{N}}$ est bornée.

□

Démonstration. (Théorème 25) D'après la Proposition 24 on sait déjà qu'une suite croissante et convergente est bornée. Il suffit alors de démontrer que si $(u_n)_{n \in \mathbb{N}}$ est une suite croissante et majorée alors la suite est convergente. La suite étant majorée alors l'ensemble

$$A = \{u_n, n \in \mathbb{N}\},$$

est un ensemble non vide et majoré de $\mathbb{R}$. Cet ensemble admet donc une borne supérieure⁴ notée λ . En particulier, on a

$$u_n \leq \lambda, \quad \forall n \in \mathbb{N}.$$

Par définition de la borne supérieure, pour tout $\varepsilon > 0$, il existe $n_0 \in \mathbb{N}$ tel que

$$u_n \geq \lambda - \varepsilon.$$

La suite $(u_n)_{n \in \mathbb{N}}$ étant croissante alors pour tout $n \geq n_0$ on a $u_{n_0} \leq u_n$. Ainsi pour tout $\varepsilon > 0$, il existe $n_0 \in \mathbb{N}$ tel que pour tout $n \geq n_0$ on ait

$$\lambda - \varepsilon < u_n \leq \lambda < \lambda + \varepsilon \Rightarrow |u_n - \lambda| < \varepsilon.$$

Ce qui implique que la suite $(u_n)_{n \in \mathbb{N}}$ est convergente, elle converge vers λ .

□

4. le plus petit des majorants de A

2.5.3 Preuve du théorème de Riemann-Lebesgue

Démonstration. (Théorème 23) Supposons dans un premier temps que f est de classe C^1 et notons

$$I_n(f) = \int_a^b f(x) e^{2i\pi nx} dx, \quad \forall n \in \mathbb{Z}.$$

En utilisant une intégration par parties, on obtient

$$\begin{aligned} I_n(f) &= \left[\frac{f(x) e^{2i\pi nx}}{2i\pi n} \right]_a^b - \int_a^b \frac{f'(x) e^{2i\pi nx}}{2i\pi n} dx \\ &= \frac{f(b) e^{2i\pi nb} - f(a) e^{2i\pi na}}{2i\pi n} - \int_a^b \frac{f'(x) e^{2i\pi nx}}{2i\pi n} dx. \end{aligned}$$

Comme f est de classe C^1 , alors f est continue sur $[a, b]$ et donc bornée. De même f' est continue sur $[a, b]$ est bornée sur $[a, b]$, d'où

$$|I_n(f)| \leq \sup_{x \in [a, b]} |f(x)| \frac{1}{\pi|n|} + \sup_{x \in [a, b]} |f'(x)| \frac{b-a}{2\pi|n|} \rightarrow 0 \quad \text{quand } |n| \rightarrow +\infty.$$

Ainsi si f est C^1 la suite $(I_n(f))_{n \in \mathbb{Z}}$ converge vers 0 quand $|n| \rightarrow +\infty$. Si maintenant f est continue sur $[a, b]$, alors par densité des fonctions polynômiales (voir Théorème 38) dans $C^0([a, b]; \mathbb{R})$, pour tout $\varepsilon > 0$ il existe g_ε une fonction polynômiale telle que

$$\sup_{x \in [a, b]} |f(x) - g_\varepsilon(x)| \leq \frac{\varepsilon}{2}.$$

De plus d'après ce qui précède la suite $(I_n(g_\varepsilon))_{n \in \mathbb{Z}}$ tend vers 0 quand $|n| \rightarrow +\infty$. Ainsi il existe $n_0 \in \mathbb{N}$ tel que pour tout $|n| \geq n_0$

$$|I_n(g_\varepsilon)| < \frac{\varepsilon}{2}.$$

On a alors par linéarité de l'intégrale

$$\begin{aligned} |I_n(f)| &= \left| \int_a^b (f(x) - g_\varepsilon(x)) e^{2i\pi nx} dx + \int_a^b g_\varepsilon(x) e^{2i\pi nx} dx \right| \\ &\leq \int_a^b |f(x) - g_\varepsilon(x)| dx + |I_n(g_\varepsilon)| \\ &\leq \sup_{x \in [a, b]} |f(x) - g_\varepsilon(x)| (b-a) + |I_n(g_\varepsilon)|. \end{aligned}$$

Par conséquent pour tout $|n| \geq n_0$ on a

$$|I_n(f)| < \varepsilon.$$

Le réel $\varepsilon > 0$ étant quelconque on en déduit que la suite $(I_n(f))_{n \in \mathbb{N}}$ tend vers 0 quand $|n| \rightarrow +\infty$. $\square$

Théorème 38 (Théorème de Weierstrass). *Pour toute fonction $f \in C^0([a, b]; \mathbb{R})$ il existe une suite de fonctions polynomiales $(b_n(f))_{n \in \mathbb{N}}$ convergent uniformément vers f .*

Ce théorème traduit la densité des fonctions polynômiales dans l'ensemble $C^0([a, b]; \mathbb{R})$. Par ailleurs, quitte à utiliser le changement de variable $x = (b - a)t + b$ on peut supposer que $a = 0$ et $b = 1$. Pour démontrer le théorème de Weierstrass on introduit deux notions. Tout d'abord, pour tout $f \in C^0([0, 1]; \mathbb{R})$, on introduit le module de continuité de f

$$\omega_f(\alpha) = \sup_{|x-y|<\alpha} |f(x) - f(y)|.$$

Si $f \in C^0([0, 1]; \mathbb{R})$ alors $\lim_{\alpha \rightarrow 0} \omega_f(\alpha) = 0$. On introduit également les polynômes de Bernstein $b_n(f) : [0, 1] \rightarrow \mathbb{R}$ avec

$$b_n(f)(x) = \sum_{k=0}^n f\left(\frac{k}{n}\right) \binom{n}{k} x^k (1-x)^{n-k} = \sum_{k=0}^n f\left(\frac{k}{n}\right) \frac{n!}{k!(n-k)!} x^k (1-x)^{n-k}.$$

Le théorème de Weierstrass est alors une conséquence immédiate du résultat suivant :

Proposition 39. *Pour tout $f \in C^0([0, 1]; \mathbb{R})$ on a*

$$\sup_{x \in [0, 1]} |f(x) - b_n(f)(x)| \leq \frac{3}{2} \omega_f\left(\frac{1}{\sqrt{n}}\right).$$

En particulier la suite $(b_n(f))_{n \in \mathbb{N}}$ converge uniformément vers f .

Démonstration. Remarquons dans un premier temps que

$$b_n(1)(x) = \sum_{k=0}^n \binom{n}{k} x^k (1-x)^{n-k} = (x + 1 - x)^n = 1.$$

De plus

$$\sum_{k=0}^n k \binom{n}{k} x^k (1-x)^{n-k} = \sum_{k=0}^n nx \binom{n-1}{k-1} x^{k-1} (1-x)^{n-k} = nx b_{n-1}(1)(x) = nx.$$

On peut réécrire cette égalité sous la forme

$$b_n(x)(x) = \sum_{k=0}^n \frac{k}{n} \binom{n}{k} x^k (1-x)^{n-k} = \frac{1}{n} \sum_{k=0}^n k \binom{n}{k} x^k (1-x)^{n-k} = \frac{nx}{x} = x.$$

De même on a

$$\begin{aligned} \sum_{k=0}^n k(k-1) \binom{n}{k} x^k (1-x)^{n-k} &= \sum_{k=2}^n n(n-1) x^2 \binom{n-2}{k-2} x^{k-2} (1-x)^{n-k} \\ &= n(n-1) x^2 b_{n-2}(1)(x) = n(n-1) x^2. \end{aligned}$$

D'où

$$n^2 b_n(x^2)(x) = n(n-1)x^2 + nx \quad \Rightarrow \quad b_n(x^2)(x) = x^2 + \frac{x(1-x)}{n}.$$

On obtient en développant et en utilisant les relations précédentes

$$\begin{aligned} \sum_{k=0}^n \left(\frac{k}{n} - x\right)^2 \binom{n}{k} x^k (1-x)^{n-k} &= b_n(x^2)(x) - 2xb_n(x)(x) + x^2 b_n(1)(x) \\ &= x^2 + \frac{x(1-x)}{n} - 2x^2 + x^2 = \frac{x(1-x)}{n}. \end{aligned}$$

On applique à présent l'inégalité de Cauchy-Schwarz pour obtenir

$$\begin{aligned} \sum_{k=0}^n \left|\frac{k}{n} - x\right| \binom{n}{k} x^k (1-x)^{n-k} &\leq \sqrt{\sum_{k=0}^n \left(\frac{k}{n} - x\right)^2 \binom{n}{k} x^k (1-x)^{n-k}} \sqrt{\sum_{k=0}^n \binom{n}{k} x^k (1-x)^{n-k}} \\ &\leq \sqrt{\frac{x(1-x)}{n}}. \end{aligned}$$

De plus comme $x(1-x) \leq 1/4$ pour tout $x \in [0, 1]$, on a

$$\sum_{k=0}^n \left|\frac{k}{n} - x\right| \binom{n}{k} x^k (1-x)^{n-k} \leq \frac{1}{2\sqrt{n}}.$$

Le nombre d'intervalles de longueur $1/\sqrt{n}$ séparant x et k/n , noté I , vérifie

$$I = \left\lceil \left|\frac{k}{n} - x\right| \sqrt{n} \right\rceil \leq \left|\frac{k}{n} - x\right| \sqrt{n} + 1,$$

où $\lceil \cdot \rceil$ désigne la partie entière supérieure, i.e., pour $x \in \mathbb{R}$ est l'unique entier $\lceil x \rceil$ vérifiant $\lceil x \rceil - 1 \leq x \leq \lceil x \rceil$. On obtient alors

$$\left| f\left(\frac{k}{n}\right) - f(x) \right| \leq \left(\left|\frac{k}{n} - x\right| \sqrt{n} + 1 \right) \omega_f\left(\frac{1}{\sqrt{n}}\right).$$

On a donc

$$\begin{aligned} |b_n(f)(x) - f(x)| &= \left| \sum_{k=0}^n \left(f\left(\frac{k}{n}\right) - f(x) \right) \binom{n}{k} x^k (1-x)^{n-k} \right| \quad (\text{car } b_n(1)(x) = 1) \\ &\leq \sum_{k=0}^n \left| f\left(\frac{k}{n}\right) - f(x) \right| \binom{n}{k} x^k (1-x)^{n-k} \\ &\leq \omega_f\left(\frac{1}{\sqrt{n}}\right) \sum_{k=0}^n \left(\left|\frac{k}{n} - x\right| \sqrt{n} + 1 \right) \binom{n}{k} x^k (1-x)^{n-k} \end{aligned}$$

$$\begin{aligned}
&\leq \omega_f \left(\frac{1}{\sqrt{n}} \right) \left(\frac{\sqrt{n}}{2\sqrt{n}} + b_n(1)(x) \right) \\
&\leq \frac{3}{2} \omega_f \left(\frac{1}{\sqrt{n}} \right).
\end{aligned}$$

Ce qui conclut la preuve du résultat.

□

Chapitre 3

Transformation de Fourier discrète

Résumé. Dans ce chapitre on étudie la transformée de Fourier discrète (TFD). Cette opération va permettre d’approcher le spectre d’un signal périodique. On présente également les principales idées de l’algorithme de « Fast Fourier Transform » (FFT). Enfin on conclura ce chapitre par deux applications de la FFT.

Références. Pour ce chapitre la plupart des résultats et preuves viennent du livre et du poly suivant :

- C. Gasquet et P. Witomski. *Analyse de Fourier et applications*, Dunod (1990), chapitre 3 (disponible à la BUTC),
- T. Rey. *Méthodes spectrales*, poly de cours Université de Lille (2022).

3.1 Introduction de la TFD

Soit f un signal a -périodique dont un échantillonnage uniforme sur une période est connu. Autrement dit, on connaît les valeurs de f en un nombre $N > 0$ fini (et uniformément réparti) de valeurs. Notons

$$y_k = f\left(k \frac{a}{N}\right) \quad \forall k \in \{0, 1, \dots, N-1\}.$$

Supposons f régulière (par exemple C_p^1 ou C_p^0 avec f C^1 par morceaux). L’objectif est alors de comprendre comment approcher numériquement les coefficients de Fourier de f , i.e., comment approcher les coefficients

$$c_n(f) = \frac{1}{a} \int_0^a f(t) \exp\left(-2i\pi n \frac{t}{a}\right) dt.$$

Pour ce faire on va utiliser la méthode des rectangles à gauche (revoir TD1). En utilisant cette approche on a

$$c_n(f) \approx \frac{1}{a} \frac{a}{N} \sum_{k=0}^{N-1} f\left(k \frac{a}{N}\right) \exp\left(-\frac{2i\pi n}{a} \frac{ka}{N}\right) = \frac{1}{N} \sum_{k=0}^{N-1} y_k \exp\left(-\frac{2i\pi nk}{N}\right).$$

On obtient alors la définition suivante :

Définition 17. Pour $N \geq 1$ notons $\omega_N \in \mathbb{C}$ l'élément donné par

$$\omega_N = \exp\left(2i\frac{\pi}{N}\right).$$

On appelle alors transformée de Fourier discrète (TFD) d'ordre N l'application notée $\mathcal{F}_N : \mathbb{C}^N \rightarrow \mathbb{C}^N$ qui à un vecteur $y \in \mathbb{C}^N$ associe le vecteur $Y \in \mathbb{C}^N$ défini par

$$Y_n = \frac{1}{N} \sum_{k=0}^{N-1} y_k \omega_N^{-n k}, \quad \text{pour tout } n \in \{0, \dots, N-1\}.$$

Autrement dit

$$(\mathcal{F}_N y)_n = Y_n, \quad \text{pour tout } n \in \{0, \dots, N-1\}.$$

Exemple 27. Donnons un premier exemple d'illustration. Soit $y = (y_0, y_1, y_2)^\top \in \mathbb{C}^3$ et calculons explicitement le vecteur $Y = \mathcal{F}_3 y \in \mathbb{C}^3$. On a pour Y_0

$$Y_0 = \frac{1}{3} \sum_{k=0}^2 y_k \omega_3^{-0 \times k} = \frac{1}{3} (y_0 \omega_3^0 + y_1 \omega_3^0 + y_2 \omega_3^0) = \frac{1}{3} (y_0 + y_1 + y_2).$$

Pour Y_1 et Y_2 on obtient

$$\begin{aligned} Y_1 &= \frac{1}{3} \sum_{k=0}^2 y_k \omega_3^{-1 \times k} = \frac{1}{3} (y_0 + y_1 e^{-2i\frac{\pi}{3}} + y_2 e^{-4i\frac{\pi}{3}}) \\ Y_2 &= \frac{1}{3} \sum_{k=0}^2 y_k \omega_3^{-2 \times k} = \frac{1}{3} (y_0 + y_1 e^{-4i\frac{\pi}{3}} + y_2 e^{-8i\frac{\pi}{3}}). \end{aligned}$$

On voit ainsi apparaître une structure matricielle. Plus précisément, en rappelant que $\omega_3 = e^{2i\frac{\pi}{3}}$ et donc $\bar{\omega}_3 = e^{-2i\frac{\pi}{3}}$, on a

$$Y = \frac{1}{3} \begin{pmatrix} 1 & 1 & 1 \\ 1 & \bar{\omega}_3 & \bar{\omega}_3^2 \\ 1 & \bar{\omega}_3^2 & \bar{\omega}_3^4 \end{pmatrix} y.$$

L'exemple précédent motive le résultat suivant :

Proposition 40. L'application $\mathcal{F}_N$ est linéaire et sa matrice dans la base canonique de $\mathbb{C}^N$ est donnée par

$$S_N = \frac{1}{N} \bar{\Omega}_N,$$

avec

$$\Omega_N = \begin{pmatrix} 1 & 1 & 1 & \dots & 1 \\ 1 & \omega_N & \omega_N^2 & \dots & \omega_N^{N-1} \\ 1 & \omega_N^2 & \omega_N^4 & \dots & \omega_N^{2(N-1)} \\ \vdots & \vdots & \vdots & & \vdots \\ 1 & \omega_N^{N-1} & \omega_N^{2(N-1)} & \dots & \omega_N^{(N-1)^2} \end{pmatrix}.$$

Démonstration. Commençons par démontrer que l'application est linéaire. Pour ce faire soient x et $y \in \mathbb{C}^N$ et notons $z = x + y \in \mathbb{C}^N$. Dans ce cas le vecteur associé $Z \in \mathbb{C}^N$ est défini, pour tout $n \in \{0, \dots, N-1\}$, par

$$\begin{aligned} Z_n &= \frac{1}{N} \sum_{k=0}^{N-1} z_k \omega_N^{-nk} = \frac{1}{N} \sum_{k=0}^{N-1} (x_k + y_k) \omega_N^{-nk} \\ &= \frac{1}{N} \sum_{k=0}^{N-1} x_k \omega_N^{-nk} + \frac{1}{N} \sum_{k=0}^{N-1} y_k \omega_N^{-nk} = X_n + Y_n. \end{aligned}$$

Autrement dit on a $(\mathcal{F}_N z)_n = (\mathcal{F}_N x)_n + (\mathcal{F}_N y)_n$ pour tout $n \in \{0, \dots, N-1\}$. On montre de manière similaire que pour tout $y \in \mathbb{C}^N$ et $\alpha \in \mathbb{C}$ alors $(\mathcal{F}_N \alpha y)_n = \alpha (\mathcal{F}_N y)_n$ pour tout $n \in \{0, \dots, N-1\}$. Donc $\mathcal{F}_N$ est une application linéaire.

Soit $X = (x_n)_{0 \leq n \leq N-1} \in \mathbb{C}^N$, alors par définition $\mathcal{F}_N X = Y \in \mathbb{C}^N$ avec $Y = (Y_n)_{0 \leq n \leq N-1}$. Le but est d'écrire que $Y = S_N X$. Pour ce faire on remarque que

$$Y_0 = (\mathcal{F}_N X)_0 = \frac{1}{N} \sum_{k=0}^{N-1} x_k \omega_N^{-0 \times k} = \frac{1}{N} \sum_{k=0}^{N-1} x_k = \frac{1}{N} (x_0 + x_1 + \dots + x_{N-1}).$$

De même

$$Y_1 = (\mathcal{F}_N X)_1 = \frac{1}{N} \sum_{k=0}^{N-1} x_k \omega_N^{-k} = \frac{1}{N} (x_0 + x_1 \omega_N^{-1} + \dots + x_{N-1} \omega_N^{-(N-1)}).$$

On obtient ainsi que pour $2 \leq n \leq N-1$ on a

$$Y_n = (\mathcal{F}_N X)_n = \frac{1}{N} \sum_{k=0}^{N-1} x_k \omega_N^{-nk} = \frac{1}{N} (x_0 + x_1 \omega_N^{-n} + x_2 \omega_N^{-2n} + \dots + x_{N-1} \omega_N^{-n(N-1)}).$$

Ainsi on obtient

$$\begin{aligned} Y &= \begin{pmatrix} Y_0 \\ Y_1 \\ Y_2 \\ \vdots \\ Y_{N-1} \end{pmatrix} = \frac{1}{N} \begin{pmatrix} x_0 + x_1 + x_2 + \dots + x_{N-1} \\ x_0 + x_1 \omega_N^{-1} + x_2 \omega_N^{-2} + \dots + x_{N-1} \omega_N^{-(N-1)} \\ x_0 + x_1 \omega_N^{-2} + x_2 \omega_N^{-4} + \dots + x_{N-1} \omega_N^{-2(N-1)} \\ \vdots \\ x_0 + x_1 \omega_N^{-(N-1)} + x_2 \omega_N^{-2(N-1)} + \dots + x_{N-1} \omega_N^{-(N-1)^2} \end{pmatrix} \\ &= \frac{1}{N} \begin{pmatrix} 1 & 1 & 1 & \dots & 1 \\ 1 & \omega_N^{-1} & \omega_N^{-2} & \dots & \omega_N^{-(N-1)} \\ 1 & \omega_N^{-2} & \omega_N^{-4} & \dots & \omega_N^{-2(N-1)} \\ \vdots & \vdots & \vdots & & \vdots \\ 1 & \omega_N^{-(N-1)} & \omega_N^{-2(N-1)} & \dots & \omega_N^{-(N-1)^2} \end{pmatrix} \begin{pmatrix} x_0 \\ x_1 \\ x_2 \\ \vdots \\ x_{N-1} \end{pmatrix} \end{aligned}$$

Donc

$$Y = \frac{1}{N} \bar{\Omega}_N X.$$

Ce qui conclut la preuve du résultat. $\square$

Remarque 41 (Complexité de l'implémentation de la TFD naïve). *Afin d'implémenter la TFD d'un vecteur $y = (y_0, \dots, y_{N-1})^\top$ il est naturel de vouloir implémenter la matrice S_N et ensuite de calculer*

$$Y = S_N y.$$

En faisant cela N^2 multiplications et $N(N-1)$ additions sont nécessaires, i.e., la complexité est de l'ordre de $N^2 + N$ opérations, ce qui est énorme et peu souhaitable. On verra à la fin de ce chapitre qu'un algorithme beaucoup plus efficace existe afin de calculer la TFD d'un vecteur. Via cet algorithme, appelé FFT, la complexité est de l'ordre de $N \log_2 N + N$ (pour rappel $\log_2 x = a \Leftrightarrow x = 2^a$).

Proposition 42. *L'application linéaire $\mathcal{F}_N$ est inversible. On note $\mathcal{F}_N^{-1}$ la transformée de Fourier discrète inverse (TFDI). De plus l'application linéaire $\mathcal{F}_N^{-1}$ a pour matrice dans la base canonique de $\mathbb{C}^N$*

$$S_N^{-1} = \Omega_N.$$

Démonstration. Il suffit de remarquer que

$$S_N \Omega_N = \frac{1}{N} \bar{\Omega}_N \Omega_N = \frac{1}{N} N I_N,$$

où I_N est la matrice identité de taille $N \times N$. En effet il convient de remarquer que pour $k \in \{0, \dots, N-1\}$ les éléments ω_N^k sont les racines N -ième de l'unité. En particulier on a

$$\sum_{k=0}^{N-1} \omega_N^k = 1 + e^{2i\frac{\pi}{N}} + (e^{2i\frac{\pi}{N}})^2 + \dots + (e^{2i\frac{\pi}{N}})^{N-1} = \frac{1 - (e^{2i\frac{\pi}{N}})^N}{1 - e^{2i\frac{\pi}{N}}} = 0.$$

Ce qui conclut la preuve. $\square$

3.1.1 Calcul de l'erreur d'approximation

La TFD permet de calculer de manière approchée les coefficients de Fourier d'une fonction. Une question naturelle est d'étudier l'erreur commise lors de cette approximation. Voici plusieurs remarques sur le sujet :

- Dans un premier temps on sait d'après le théorème de Riemann-Lebesgue que si f est continue par morceaux alors

$$c_n(f) \rightarrow 0, \quad \text{quand } |n| \rightarrow +\infty.$$

Ainsi on peut penser que pour $n \ll \text{grand}$ les coefficients de Fourier de f sont « négligeables » et que l'essentiel de l'information sur f est stockée dans les premiers coefficients de f . Ce qui permet d'espérer que pour obtenir une approximation relativement fidèle de f (par exemple aux points ka/N) on peut se restreindre à approcher un nombre fini de coefficients de f .

- On sait en fait un peu mieux que ça. En effet d'après le TD sur le chapitre 2, si f est C^1 (par morceaux) alors il existe une constante $M > 0$ telle que

$$|c_n(f)| \leq \frac{M}{|n|}, \quad \forall n \in \mathbb{Z} \setminus \{0\}.$$

Ce qui fournit une information plus précise que le théorème de Riemann-Lebesgue. De plus toujours d'après le TD sur le chapitre 2 si f est de classe C^p (par morceaux) pour $p \geq 1$ alors

$$|c_n(f)| \leq \frac{K}{|n|^p}, \quad \forall n \in \mathbb{Z} \setminus \{0\},$$

où K est une constante. Ainsi plus f est régulière et plus ses coefficients tendent rapidement vers 0. On peut donc légitimement penser que plus la fonction est régulière et plus la TFD donne une bonne approximation de f .

Essayons à présent de quantifier l'erreur d'approximation commise. Pour ce faire pour f une fonction donnée, on note $c_n(f)$ le coefficient de Fourier d'ordre n de f et $c_n^N(f)$ le coefficient de Fourier approché de f d'ordre n , i.e.,

$$c_n^N(f) = \frac{1}{N} \sum_{k=0}^{N-1} f\left(k \frac{a}{N}\right) \omega_N^{-nk}, \quad \forall n \in \{0, \dots, N-1\}.$$

Si f est de classe C^1 alors les coefficients $c_n^N(f)$ obtenus par la méthode des rectangles à gauche on peut s'attendre, d'après le premier TD (exo 2), à une inégalité de la forme

$$|c_n(f) - c_n^N(f)| \leq \frac{C}{N}, \quad \forall n \in \{0, \dots, N-1\},$$

où C est une constante. En effet on a

$$\begin{aligned} c_n(f) - c_n^N(f) &= \frac{1}{a} \int_0^a f(x) e^{-2i\pi n \frac{x}{a}} dx - \frac{1}{N} \sum_{k=0}^{N-1} f\left(k \frac{a}{N}\right) e^{-2i\pi n \frac{k}{N}} \\ &= \frac{1}{a} \left(\int_0^a f(x) \overline{e_n}(x) dx - \frac{a}{N} \sum_{k=0}^{N-1} f\left(k \frac{a}{N}\right) \overline{e_n}\left(k \frac{a}{N}\right) \right), \end{aligned}$$

où $e_n(x) = e^{2i\pi nx/a}$. En posant $g(x) = f(x) \overline{e_n}(x)$ alors on a

$$c_n(f) - c_n^N(f) = \frac{1}{a} \left(\int_0^a g(x) dx - \frac{a}{N} \sum_{k=0}^{N-1} g\left(k \frac{a}{N}\right) \right).$$

Or g est C^1 comme produit de deux fonctions C^1 , on en déduit alors du TD 1 l'existence d'une constante (dépendant de f , f' et a) vérifiant

$$|c_n(f) - c_n^N(f)| \leq \frac{C}{N}.$$

Ainsi si f est C^1 l'erreur commise est assez mauvaise (N doit être très grand pour avoir une approximation « raisonnable » des $c_n(f)$). Cependant si f est plus régulière on a le résultat (admis) suivant :

Théorème 43. *Soit f une fonction a -périodique sur $\mathbb{R}$ et de classe C^p avec $p \in \mathbb{N}$ et $p \geq 2$. Alors il existe une constante $C = C(p, f)$ telle que*

$$|c_n(f) - c_n^N(f)| \leq \frac{C}{N^p}$$

3.2 Quelques explications sur la FFT

On suppose dans un premier temps que N est pair, de la forme $N = 2M$ pour $M \geq 1$ un entier. On considère $y = (y_0, \dots, y_{N-1})^\top \in \mathbb{C}^N$ un vecteur donné. Alors en remarquant que

$$\omega_N = e^{2i\frac{\pi}{N}} = e^{i\frac{\pi}{M}} = \omega_M^{1/2} \quad \Rightarrow \quad \omega_M = \omega_N^2,$$

on va regrouper les composantes paires et impaires du vecteur y lors du calcul des vecteurs $Y_n = (\mathcal{F}_N y)_n$ pour $n \in \{0, \dots, N-1\}$. Plus précisément on a

$$\begin{aligned} Y_n &= \frac{1}{N} \sum_{k=0}^{N-1} y_k \omega_N^{-nk} = \frac{1}{2M} \sum_{k=0}^{2M-1} y_k \omega_N^{-nk} \\ &= \frac{1}{2M} \left[y_0 + y_2 \underbrace{\omega_N^{-2n}}_{=\omega_M^{-n}} + \dots + y_{2M-2} \underbrace{\omega_N^{-(2M-2)n}}_{=\omega_M^{-(M-1)n}} + y_1 \omega_N^{-n} + y_3 \omega_N^{-3n} + \dots + y_{2M-1} \omega_N^{-(2M-1)n} \right] \\ &= \frac{1}{2M} \left[y_0 + y_2 \omega_M^{-n} + \dots + y_{2M-2} \omega_M^{-(M-1)n} + \omega_N^{-n} \left(y_1 + y_3 \underbrace{\omega_N^{-2n}}_{=\omega_M^{-n}} + \dots + y_{2M-1} \underbrace{\omega_N^{-(2M-2)n}}_{=\omega_M^{-(M-1)n}} \right) \right]. \end{aligned}$$

Ainsi on obtient pour tout $n \in \{0, \dots, 2M-1\}$

$$\begin{aligned} (\mathcal{F}_N y)_n = Y_n &= \frac{1}{2} \left[\frac{1}{M} \sum_{k=0}^{M-1} y_{2k} \omega_M^{-nk} + \frac{\omega_N^{-n}}{M} \sum_{k=0}^{M-1} y_{2k+1} \omega_M^{-nk} \right] \\ &= \frac{1}{2} \left[(\mathcal{F}_M y^{\text{pair}})_n + \omega_N^{-n} (\mathcal{F}_M y^{\text{impair}})_n \right], \end{aligned}$$

où

$$y^{\text{pair}} = (y_0, y_2, \dots, y_{2M-2})^\top, \quad y^{\text{impair}} = (y_1, y_3, \dots, y_{2M-1})^\top.$$

On remarque de plus que pour tout $n \in \{0, \dots, M-1\}$ on a

$$\omega_N^{-(n+M)} = e^{-2i\pi \frac{n+M}{N}} = e^{-2i\pi \frac{n}{N}} e^{-i\pi \frac{2M}{N}} = e^{-2i\pi \frac{n}{N}} e^{-i\pi} = -\omega_N^{-n}.$$

Ainsi pour $n \in \{0, \dots, M-1\}$ on a

$$\begin{aligned} Y_{n+M} &= \frac{1}{2} \left[\frac{1}{M} \sum_{k=0}^{M-1} y_{2k} \omega_M^{-nk} \underbrace{\omega_M^{-Mk}}_{=1} + \frac{\omega_N^{-(n+M)}}{M} \sum_{k=0}^{M-1} y_{2k+1} \omega_M^{-nk} \underbrace{\omega_M^{-Mk}}_{=1} \right] \\ &= \frac{1}{2} \left[\frac{1}{M} \sum_{k=0}^{M-1} y_{2k} \omega_M^{-nk} - \frac{\omega_N^{-n}}{M} \sum_{k=0}^{M-1} y_{2k+1} \omega_M^{-nk} \right] = \frac{1}{2} \left[(\mathcal{F}_M y^{\text{pair}})_n - \omega_N^{-n} (\mathcal{F}_M y^{\text{impair}})_n \right]. \end{aligned}$$

En résumé pour calculer les vecteurs $Y_M, \dots, Y_{2M-1}$ il suffit de calculer les vecteurs

$$(\mathcal{F}_M y^{\text{pair}})_0, \dots, (\mathcal{F}_M y^{\text{pair}})_{M-1} \quad \text{et} \quad (\mathcal{F}_M y^{\text{impair}})_0, \dots, (\mathcal{F}_M y^{\text{impair}})_{M-1}.$$

On peut ainsi procéder de la manière suivante pour accélérer le calcul de la TFD d'un vecteur $y \in \mathbb{C}^N$ (avec $N = 2M$) :

1. on calcule $(\mathcal{F}_M y^{\text{pair}})_n$ et $\omega_N^{-n} (\mathcal{F}_M y^{\text{impair}})_n$ pour $n \in \{0, \dots, M-1\}$,
2. on obtient alors Y_n pour $n \in \{0, \dots, M-1\}$,
3. on calcule ensuite $Y_{n+M} = ((\mathcal{F}_M y^{\text{pair}})_n - \omega_N^{-n} (\mathcal{F}_M y^{\text{impair}})_n)/2$ pour $n \in \{0, \dots, M-1\}$.

Par cette approche on effectue $2(M-1)^2 + M-1$ multiplications lors de l'étape (1), les multiplications des étapes (2) et (3) ne coûtent rien. Ainsi cette stratégie permet de calculer la TFD de y via $\approx 2M^2 = N^2/2$ multiplications. On a donc divisé le nombre de multiplications par 2 par rapport à une implémentation naïve de la TFD.

Le gain précédent n'est pas négligeable mais on peut aller plus loin. En effet on peut répéter (en regroupant les composantes paires et impaires) l'opération précédente aux vecteurs $\mathcal{F}_M^{\text{pair}}$ et $\mathcal{F}_M^{\text{impair}}$. Comme ces vecteurs appartiennent à $\mathbb{C}^M$ pour répéter ces opérations il faut supposer que M est pair. Pour cette raison le bon cadre est de supposer que $N = 2^p$ pour $p \geq 2$. On itère alors les opérations p fois pour se ramener au calcul de la TFD d'un vecteur à 2 composantes. Or ce calcul est particulièrement simple. En effet notons $x = (x_0, x_1)^\top$ les composantes de ce vecteur alors

$$\begin{aligned} (\mathcal{F}_2 x)_0 &= \frac{1}{2}(x_0 + x_1 \omega_2^{-0}) = \frac{1}{2}(x_0 + x_1), \\ (\mathcal{F}_2 x)_1 &= \frac{1}{2}(x_0 + x_1 \omega_2^{-1}) = \frac{1}{2}(x_0 - x_1). \end{aligned}$$

On peut alors montrer que cet algorithme est d'une complexité de l'ordre de $N \log_2 N$. À titre d'exemple si $N = 1024$ alors via l'approche naïve par TFD il faut ≈ 2000000 de multiplications et d'additions contre ≈ 14000 additions et multiplications par FFT!

3.3 Applications de la TFD/FFT

Dans cette section on va présenter deux exemples d'application de l'algorithme de FFT.

3.3.1 FFT, convolution discrète et débruitage d'un signal

Introduisons dans un premier temps la notion de vecteur périodique :

Définition 18. Soit $y \in \mathbb{C}^{\mathbb{Z}}$ un vecteur de taille « infinie ». On dit que y est N -périodique si

$$y_{k+\ell N} = y_k \quad \forall \ell \in \mathbb{Z}, k \in \{0, \dots, N-1\}.$$

Exemple 28. Soit f une fonction a -périodique et y le vecteur de $\mathbb{C}^N$ obtenu via la procédure d'échantillonnage uniforme sur une période, i.e.,

$$y_k = f\left(k \frac{a}{N}\right), \quad \forall k \in \{0, \dots, N-1\}.$$

La fonction étant a -périodique alors le vecteur toujours noté y « périodisé » de $\mathbb{C}^{\mathbb{Z}}$ est N -périodique.

Nous pouvons maintenant introduire la notion de convolution discrète

Définition 19. Soient x et y deux vecteurs N -périodiques. On appelle convolution périodique (d'ordre N) de x et y le vecteur, noté $z = x *_N y$, avec pour composantes

$$z_k = (x *_N y)_k = \sum_{\ell=0}^{N-1} x_{\ell} y_{k-\ell} \quad \forall k \in \{0, \dots, N-1\}.$$

Dans une première approche le calcul de la convolution périodique d'ordre $N \geq 1$ est obtenu comme multiplication d'une matrice de taille $N \times N$, dite de Toeplitz (voir section suivante), et d'un vecteur. Ainsi pour deux vecteurs de tailles N il faut N^2 multiplications pour calculer la convolution périodique de deux vecteurs. Cette approche « directe » est trop coûteuse en temps calculs. En pratique on utilise la TFD (et en particulier l'algorithme de FFT). En effet, on a le résultat suivant :

Proposition 44. Soient $x = (x_n)$ et $y = (y_n)$ deux vecteurs N -périodiques et $X = (X_n)$ et $Y = (Y_n)$ les vecteurs obtenus par TFD, i.e., $X_n = (\mathcal{F}_N x)_n$ et $Y_n = (\mathcal{F}_N y)_n$. Alors en notant $z = (z_n)$ le vecteur $z_n = (x *_N y)_n$ et $Z = (Z_n)$ avec $Z_n = (\mathcal{F}_N z)_n$ on a

$$Z_n = N X_n Y_n \quad \forall n \in \{0, \dots, N-1\}.$$

Démonstration. Admise. □

On peut réécrire l'égalité précédente sous la forme

$$(\mathcal{F}_N(x *_N y))_n = N (\mathcal{F}_N x)_n (\mathcal{F}_N y)_n \quad \forall n \in \{0, \dots, N-1\}.$$

Ainsi la TFD transforme le produit de convolution en produit.

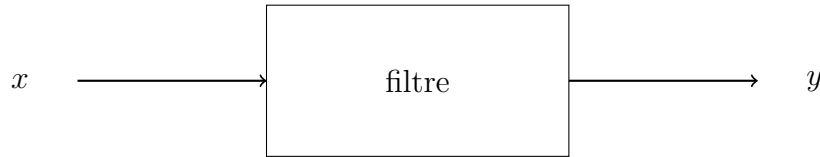
Remarque 45. Pour l'instant nous avons défini la DFT pour un vecteur de taille finie. Or, dans la proposition précédente, on applique la DFT à des vecteurs périodiques (donc « infinis »). Y a-t-il alors une ambiguïté lorsque dans le résultat précédent on parle de X et Y les vecteurs obtenus par TFD de x et y ? La réponse est oui et non ! En effet, dans le cas d'un signal N périodique discret x , le vecteur $X \in \mathbb{C}^{\mathbb{Z}}$ dont les composantes sont données par la formule

$$X_n = \frac{1}{N} \sum_{k=0}^{N-1} x_k \omega_N^{-nk}, \quad \forall n \in \mathbb{Z},$$

est appelé Transformée de Fourier du signal périodique à temps discret x . Cependant, on montre facilement que le vecteur $X \in \mathbb{C}^{\mathbb{Z}}$ obtenu via la formule précédente est N -périodique (voir TD). Ainsi dans le cas d'un signal discret N périodique le spectre du signal est entièrement contenu dans les N premiers coefficients du vecteur X . Par abus de langage on peut alors parler du vecteur X comme étant la DFT de x .

Application. Imaginons vouloir enregistrer durant un concert un morceau de musique. Lors de l'enregistrement la salle de concert est remplie d'auditeurs provoquant des interférences et rendant le signal sonore « pollué » ou bruité. Que faire pour traiter ce signal afin de le « nettoyer » au maximum des interférences et être en mesure de restituer au mieux le morceau de musique originel ?

On suppose enregistrer un morceau durant une période T . On effectue un enregistrement de N échantillons (répartis uniformément) durant la période T . On peut représenter mathématiquement ce signal (discret) d'entrée par le vecteur $x = (x_n)_{0 \leq n \leq N-1}$. Comme expliqué précédemment ce signal est pollué ou bruité, on veut traiter le signal discret x pour obtenir un signal de sortie, noté $y = (y_n)_{0 \leq n \leq N-1}$, aussi propre que possible (i.e. se rapprochant le plus possible du morceau de musique sans bruits parasites). Pour ce faire on va appliquer un filtre au signal x , schématiquement on a



Appliquer un filtre revient à effectuer une opération mathématique sur les composantes de x et les composantes de y vont être obtenues via

$$y_n = \sum_{\ell=0}^{N-1} x_{n-\ell} h_{\ell} = (x *_N h)_n, \quad \forall n \in \{0, \dots, N-1\}, \quad (3.1)$$

où le vecteur (N -périodique) $h = (h_n)$ correspond au filtre. Ce filtre permet de traiter le signal, sa définition dépend du traitement souhaité (par exemple un filtre passe bas pour atténuer l'effet du « souffle » corrompant une musique).

En notant $x = (x_0, x_1, \dots, x_{N-1})^\top$ et $y = (y_0, y_1, \dots, y_{N-1})^\top$, remarquons que (3.1) peut s'écrire sous la forme matricielle

$$y = \mathcal{M}_h x,$$

où $\mathcal{M}_h$ est la matrice (dite de Toeplitz)

$$\mathcal{M}_h = \begin{pmatrix} h_0 & h_{N-1} & h_{N-2} & \dots & h_1 \\ h_1 & h_0 & h_{N-1} & \dots & h_2 \\ h_2 & h_1 & h_0 & \dots & h_3 \\ \vdots & \vdots & \vdots & & \vdots \\ h_{N-1} & h_{N-2} & h_{N-3} & \dots & h_0 \end{pmatrix}.$$

Ainsi pour calculer les coefficients de Y il faut N^2 multiplications et $N(N-1)$ additions, ce qui est trop lourd pour un traitement informatique efficace. La Proposition 44 suggère un autre moyen beaucoup plus efficace de calculer les coordonnées du vecteur y . En effet, via l'algorithme de transformée de Fourier discrète (FFT) ainsi que son inverse (IFFT), on passe du domaine temporel au domaine fréquentiel en calculant $X = (X_n)$ et $H = (H_n)$ avec

$$X_n = (\mathcal{F}_N x)_n \quad \text{et} \quad H_n = (\mathcal{F}_N h)_n, \quad \forall n \in \{0, \dots, N-1\} \quad [\approx N \log_2 N]$$

puis

$$Y_n = N X_n H_n, \quad \forall n \in \{0, \dots, N-1\}, \quad [\approx N]$$

et enfin on repasse au domaine temporel par le biais de

$$y_n = (\mathcal{F}_N^{-1} Y)_n, \quad \forall n \in \{0, \dots, N-1\} \quad [\approx N \log_2 N]$$

Via cette méthode la complexité est de l'ordre de $N \log_2 N + N$ opérations ! On en déduit que l'on peut appliquer un filtre ou traiter un signal de manière très efficace.

On peut en fait inclure l'étude précédente dans un cadre beaucoup plus général via la théorie des Systèmes Linéaires et Invariant en Temps (SLIT). Cette théorie s'intéresse à l'étude d'un système (physique, mécanique, biologique, chimique etc) noté $\mathcal{S}$ prenant un signal d'entrée x et renvoyant un signal de sortie y :

$$y(t) = \mathcal{S}[x(t)] \quad t \in \mathbb{R}.$$

Le système $\mathcal{S}$ est un SLIT si les propriétés suivantes sont vérifiées

Linéarité. Soient x_1, x_2 deux signaux et $a \in \mathbb{R}$ alors

$$\mathcal{S}[ax_1(t) + x_2(t)] = a\mathcal{S}[x_1(t)] + \mathcal{S}[x_2(t)] \quad t \in \mathbb{R}$$

Invariance en temps. Pour $t_0 \in \mathbb{R}$ on a

$$x(t) = \mathcal{S}[y(t)] \Leftrightarrow x(t - t_0) = \mathcal{S}[y(t - t_0)] \quad t \in \mathbb{R}.$$

Exemple 29 (Circuit RC). Soit $x(t)$ une tension d'entrée et une sortie $v(t)$ représentant la tension aux bornes d'un condensateur de charge Q . Le système est ici gouverné par une équation différentielle de la forme

$$RC v'(t) + v(t) = x(t) \quad t \in \mathbb{R}.$$

De façon générale un SLIT est gouverné par une équation différentielle. Cependant un résultat mathématique affirme qu'un SLIT peut toujours s'écrire sous la forme d'un produit de convolution¹ entre le signal d'entrée x et une certaine fonction h caractérisant le système $\mathcal{S}$. Schématiquement on a

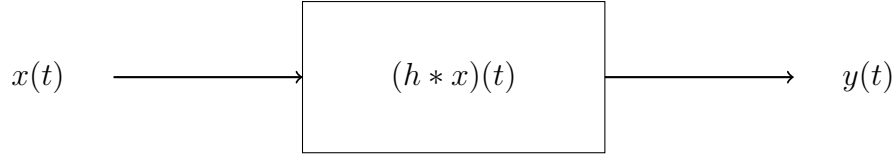


FIGURE 3.1 – Représentation graphique d'un SLIT (au niveau continu).

En pratique les signaux considérés sont obtenus via des échantillonnages du signal (continu) initial, i.e., on travaille avec des signaux discrets $x = (x_n)_{0 \leq n \leq N-1}$. Les affirmations précédentes s'adaptent au niveau discret et on obtient la représentation schématique suivante d'un SLIT numérique :

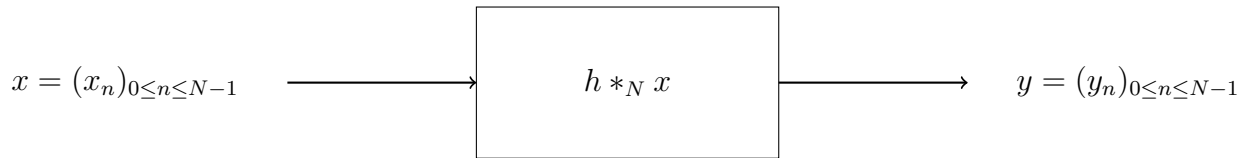


FIGURE 3.2 – Représentation graphique d'un SLIT (au niveau discret).

Remarquons que dans le cas discret (ou continu) un SLIT est entièrement déterminé par le vecteur (ou fonction) h . En particulier, pour déterminer h il suffit de déterminer la réponse du système à une impulsion dite de Dirac $x = (x_0, \dots, x_{N-1}) = (1, 0, \dots, 0)$ de telle manière à obtenir

$$h_n = \sum_{\ell=0}^{N-1} x_\ell h_{n-\ell} \quad \forall n \in \{0, \dots, N-1\}.$$

1. On définira proprement cette notion plus tard dans le cours.

Pour cette raison h est appelé réponse impulsionnelle. Par ailleurs, afin de calculer efficacement le produit de convolution $h *_N x$ on utilise la TFD et généralement la définition de la réponse impulsionnelle est donnée dans le domaine fréquentiel. En effet la TFD de la réponse impulsionnelle est appelée fonction de transfert !

3.3.2 FFT et compression d'image

On peut associer à une image en noir et blanc constituée de $N \times M$ pixels une matrice, notée A , de taille $N \times M$ où les coefficients de la matrice sont des entiers compris entre 0 (noir) et 255 (blanc). Par ailleurs, on peut généraliser la TFD en dimension 2. En effet en notant $\hat{A} \in \mathcal{M}_{N \times M}(\mathbb{C})$ la TFD de A , on a la formule suivante :

$$\hat{A}_{k,l} = \frac{1}{MN} \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} A_{n,m} e^{-2i\pi\left(\frac{kn}{N} + \frac{lm}{M}\right)}.$$

En particulier on peut appliquer des outils de la théorie de Fourier afin de faire du traitement d'image. Un exemple classique d'application de ces outils est la compression d'image. Une première idée d'algorithme de compression d'image consiste à conserver un certain pourcentage de fréquences de hautes amplitudes est de mettre brutalement à 0 les fréquences de faibles amplitudes.

On va considérer l'image de la Figure 3.3. Cette image est constituée de 512×512 pixels. On peut donc associer à cette image une matrice de taille 512×512 . En implémentant sous **Scilab** l'algorithme expliqué ci-dessus on obtient la Figure 3.4.

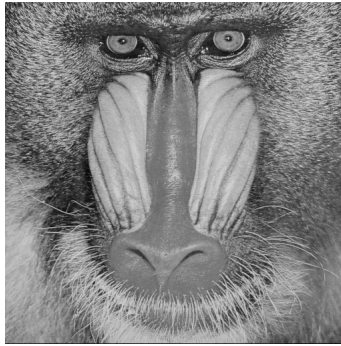


FIGURE 3.3 – Image initiale

Le code **Scilab** est disponible à la page 15. Afin de faire tourner ce code il faut dans un premier temps installer l'application *Image Processing and Computer Vision Toolbox*. Pour ce faire dans la console **Scilab**, il faut cliquer sur les onglets :

Applications → Gestionnaire de Modules - ATOMS.

Une fois le gestionnaire de modules ouvert, dans le menu de gauche cliquer sur :

Traitement d'images → Image Processing and Computer Vision Toolbox → installer.

Une fois l'installation effectuée il suffit de relancer **Scilab**. Si tout marche bien, le message suivant doit s'afficher lors de l'ouverture de la console **Scilab** :

Initialisation :

Chargement de l'environnement de travail

Start IPCV 4.1.2 for Scilab 6.0

Image Processing and Computer Vision Toolbox for Scilab

2019 - Bytecode Malaysia

Load macros

Load dependencies

Load gateways

Load help

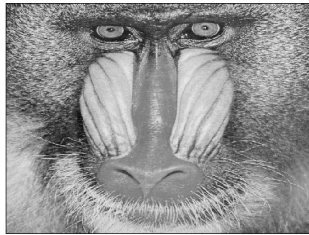
Load demos

Il suffit ensuite de copier le code !

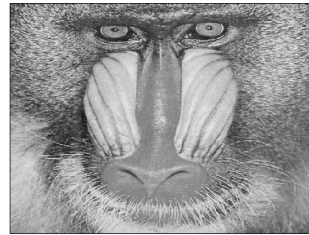
```

1  clear
2  for id=winsid() close(scf(id)) end
3
4  U=double(imread('Chemin\mandrill.bmp'));//Transforme une image en matrice
5
6  //U=rgb2gray(U);//Permet de convertir une image en noir et blanc
7
8  scf();
9  imshow(mat2gray(U))//Affiche l'image
10
11 //Graphe FFT de U
12
13 Ut=fft(U);//Calcul de la FFT 2d de la matrice U
14 Ulog=log(abs(fftshift(Ut))+1);
15
16 scf();
17 imshow(mat2gray(Ulog)) //Graphe de la DFT de U (echelle log)
18
19 //Algo compression image
20
21 t=0.01;//t*100 pourcentage de frequences preservees
22
23 Utsort=gsort(abs(Ut(:)), 'g', 'i');//Transforme abs(Ut) en un vecteur
    croissant

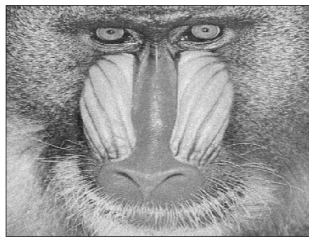
```



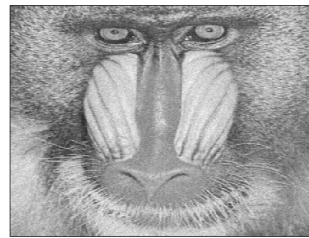
(a) 90%



(b) 50%



(c) 20%



(d) 10%



(e) 1%



(f) 0,1%

FIGURE 3.4 – Compression de l'image initiale (on indique le pourcentage de pixels préservés).

```
24 thresh=Utsort(floor((1-t)*length(Utsort)));
25 ind=abs(Ut)>thresh;//Definition du filtre
26 Utrunc=Ut.*ind;//Permet de filtrer les frequences trop basses
27
28 Utrunclog=log(abs(fftshift(Utrunc))+1);
29 scf();
30 imshow(mat2gray(Utrunclog))
31
32 //Retour dans le domaine spatial
33
34 Ulow=ifft(Utrunc);//Permet d'obtenir l'image compressee
35
36 scf();
37 imshow(mat2gray(Ulow))
```

Chapitre 4

Intégrale de Lebesgue

Résumé. Dans ce chapitre on prépare l'introduction de la notion de transformée de Fourier. Cependant afin d'introduire ce concept il faut dans un premier temps généraliser la notion d'intégrale de Riemann. C'est le but de l'intégrale de Lebesgue. Notre principal objectif est donc de présenter des concepts et outils venant de la théorie d'intégration de Lebesgue mais pas de présenter de façon précise et rigoureuse cette théorie. Nous allons adopter une vision utilitariste de l'intégrale de Lebesgue.

Références. Pour ce chapitre la plupart des résultats et preuves viennent

- C. Gasquet et P. Witomski. *Analyse de Fourier et applications*, Dunod (1990), chapitre 3 (disponible à la BUTC),
- J.-M. Bony. *Cours d'analyse, théorie des distributions et analyse de Fourier*, les éditions de l'école Polytechniques (2001).

4.1 Un argument en défaveur de l'intégrale de Riemann

Dans la suite du cours on va étudier la transformation de Fourier, cette transformation associe à une fonction une nouvelle fonction dépendant d'une intégrale sur $\mathbb{R}$, i.e., on va vouloir définir une transformation de la forme

$$f \xrightarrow{\text{Transformation}} \widehat{f}(\xi) = \int_{\mathbb{R}} f(x) e^{-2i\pi\xi x} dx \quad (\xi \in \mathbb{R}).$$

Si f est nulle en dehors d'un intervalle borné de $\mathbb{R}$ de la forme $[a, b]$ alors on a envie d'écrire que

$$\widehat{f}(\xi) = \int_a^b f(x) e^{-2i\pi\xi x} dx, \quad \xi \in \mathbb{R}.$$

Dans ce cas la théorie d'intégration de Riemann permet de donner un sens à $\widehat{f}(\xi)$. En effet si $x \mapsto f(x) e^{-2i\pi\xi x}$ est Riemann intégrable pour $\xi \in \mathbb{R}$ alors $\widehat{f}(\xi)$ est bien définie. Si par

contre f est non nulle sur un intervalle de longueur « infinie » alors la théorie de Riemann ne nous permet pas de définir $\widehat{f}(\xi)$.

Pourtant si par exemple on considère la fonction f donnée par

$$f(x) = \begin{cases} e^{-ax} & \text{si } x \geq 0, \\ 0 & \text{si } x < 0, \end{cases} \quad (4.1)$$

avec a un réel strictement positif. Alors dans ce cas, pour $\xi \in \mathbb{R}$ on a

$$\widehat{f}(\xi) = \int_0^{+\infty} e^{-(a+2i\pi\xi)x} dx.$$

Intuitivement on veut écrire que pour tout $\xi \in \mathbb{R}$ alors

$$\widehat{f}(\xi) = \lim_{n \rightarrow +\infty} \left[\frac{e^{-(a+2i\pi\xi)x}}{-(a+2i\pi\xi)} \right]_0^n = \frac{1}{a+2i\pi\xi}. \quad (4.2)$$

Ainsi $\widehat{f}(\xi) \in \mathbb{R}$ possède bien un sens pour tout $\xi \in \mathbb{R}$ même si f est non nulle sur $[0, +\infty[$

Grossièrement on vient de voir l'existence de fonctions définies sur un intervalle I non bornée de $\mathbb{R}$, donc pas Riemann intégrables, qui admettent une intégrale finie. On aimerait donc trouver une notion d'intégrable qui permette de justifier le calcul (4.2) (on veut s'affranchir de la condition I borné).

Bien entendu il existe d'autres arguments (plus ou moins théoriques) qui imposent de devoir considérer une notion d'intégrale plus générale que celle de Riemann. On peut par exemple citer l'existence de fonctions bornées qui ne sont pas Riemann intégrables, par exemple

$$\mathbf{1}_{\mathbb{Q} \cap [0,1]} = \begin{cases} 1 & \text{si } x \in \mathbb{Q} \cap [0,1], \\ 0 & \text{si } x \in (\mathbb{R} \setminus \mathbb{Q}) \cap [0,1], \end{cases} \quad (4.3)$$

l'existence de suites de fonctions Riemann intégrables qui ne convergent pas (en un certain sens) vers une fonction Riemann intégrable etc. Dans ce cours on gardera en tête la nécessité de devoir généraliser l'intégrale de Riemann afin d'être en mesure de donner un sens à l'intégrale de fonctions définies sur $\mathbb{R}$.

4.2 L'intégrale de Lebesgue

Pour présenter l'intégrale de Lebesgue on va dès le début admettre de nombreux résultats difficiles à démontrer mais qui se comprennent facilement¹. En particulier, on va d'abord considérer le cas des fonctions à valeurs dans $\overline{\mathbb{R}}_+ = [0, +\infty]$, i.e., des fonctions positives valant éventuellement $+\infty$. Par exemple la fonction

$$f(x) = \begin{cases} x^{-1} & \text{si } x \geq 0, \\ 0 & \text{si } x < 0, \end{cases}$$

1. Je renvoie à la Section 4.5.1 pour la présentation de quelques idées clefs de la construction de l'intégrale de Lebesgue des fonctions positives

est définie sur $\mathbb{R}$ à valeurs dans $\overline{\mathbb{R}}_+$ ($f(0) = +\infty$). Dans $\overline{\mathbb{R}}_+$ les opérations arithmétiques usuelles (addition et multiplication) sont inchangées avec les conventions suivantes : $a + \infty = +\infty$ ($a \in \overline{\mathbb{R}}_+$), $a \times (+\infty) = +\infty$ ($a > 0$) et $0 \times (+\infty) = 0$.

4.2.1 Le cas des fonctions positives

Pour toute fonction $f : \mathbb{R} \rightarrow \overline{\mathbb{R}}_+$ on admet l'existence d'une application (dite intégrale de Lebesgue)

$$f \mapsto \int_{\mathbb{R}} f(x) dx \in \overline{\mathbb{R}}_+,$$

vérifiant les propriétés suivantes.

1. L'application est linéaire, pour tout $f, g : \mathbb{R} \rightarrow \overline{\mathbb{R}}_+$ et $\alpha, \beta \in]0, +\infty[$

$$\int_{\mathbb{R}} (\alpha f(x) + \beta g(x)) dx = \alpha \int_{\mathbb{R}} f(x) dx + \beta \int_{\mathbb{R}} g(x) dx.$$

2. L'application est monotone, i.e., si $f(x) \leq g(x)$ pour tout $x \in \mathbb{R}$ alors

$$\int_{\mathbb{R}} f(x) dx \leq \int_{\mathbb{R}} g(x) dx.$$

3. En notant $\mathbf{1}_{[a,b]}$ ($a < b$) la fonction valant 1 sur $[a, b]$ et 0 en dehors alors

$$\int_{\mathbb{R}} \mathbf{1}_{[a,b]}(x) dx = b - a.$$

4. Soit A un sous ensemble de $\mathbb{R}$ et f une fonction de $\mathbb{R}$ dans $\overline{\mathbb{R}}_+$ alors

$$\int_{\mathbb{R}} f(x) \mathbf{1}_A(x) dx = \int_A f(x) dx.$$

5. Enfin si $(f_j)_{j \in \mathbb{N}}$ est une suite croissante de fonctions définies sur $\mathbb{R}$ et à valeurs dans $\overline{\mathbb{R}}_+$ alors

$$\int_{\mathbb{R}} \lim_{j \rightarrow +\infty} f_j(x) dx = \lim_{j \rightarrow +\infty} \int_{\mathbb{R}} f_j(x) dx.$$

Remarque 46. Le point (5) est le théorème de Beppo-Levi ou de convergence monotone. Ce résultat implique en particulier que si (u_j) est une suite de fonctions positives (prenant éventuellement la valeur $+\infty$) alors on a toujours

$$\int_{\mathbb{R}} \sum_{j=1}^{+\infty} u_j(x) dx = \sum_{j=1}^{+\infty} \int_{\mathbb{R}} u_j(x) dx.$$

En effet il suffit d'appliquer le théorème de Beppo-Levi à la suite de fonctions $S_N(x) = \sum_{j=1}^N u_j(x)$.

Revenons sur l'exemple de la fonction (4.1). On considère pour tout $n \geq 1$ la suite

$$f_n(x) = \begin{cases} e^{-ax} & \text{si } x \in [0, n], \\ 0 & \text{sinon.} \end{cases}$$

Autrement dit on peut écrire que $f_n(x) = \mathbf{1}_{[0, n]}(x)f(x)$ pour tout $x \in \mathbb{R}$. On applique alors la propriété (4) de l'intégrale de Lebesgue et on a

$$\int_{\mathbb{R}} f_n(x) dx = \int_0^n f(x) dx = \left[-\frac{e^{-ax}}{a} \right]_0^n = 1 - \frac{e^{-an}}{a}.$$

Ainsi comme $f(x) = \lim_{n \rightarrow +\infty} f_n(x)$ pour tout $x \in \mathbb{R}$ alors d'après le théorème de Beppo-Levi on obtient

$$\int_{\mathbb{R}} f(x) dx = \int_{\mathbb{R}} \lim_{n \rightarrow +\infty} f_n(x) dx = \lim_{n \rightarrow +\infty} \int_{\mathbb{R}} f_n(x) dx = \lim_{n \rightarrow +\infty} \left(1 - \frac{e^{-an}}{a} \right) = 1.$$

Des manipulations similaires permettent de justifier les égalités (4.2). À partir de maintenant on va directement écrire, par exemple,

$$\int_{\mathbb{R}} e^{-ax} \mathbf{1}_{[0, +\infty[} dx = \int_0^{+\infty} e^{-ax} dx = \left[-\frac{e^{-ax}}{a} \right]_0^{+\infty} = \frac{1}{a}.$$

Définissons à présent la mesure (de Lebesgue) d'un ensemble de $\mathbb{R}$.

Définition 20. Soit A un sous ensemble de $\mathbb{R}$, on note $\mathbf{1}_A$ la fonction caractéristique de l'ensemble A (valant 1 sur A et 0 ailleurs). On appelle alors mesure de A

$$\mu(A) = \int_{\mathbb{R}} \mathbf{1}_A(x) dx = \int_A dx \in \overline{\mathbb{R}}_+.$$

L'application μ est appelée mesure de Lebesgue, elle donne à chaque parties de $\mathbb{R}$ une valeur de $\overline{\mathbb{R}}_+$. L'application μ généralise la notion de longueur dans $\mathbb{R}$ et vérifie :

- si $A \subset B$ alors $\mu(A) \leq \mu(B)$ (conséquence de la monotonie de l'intégrale),
- Soit la suite $(A_j)_{j \in \mathbb{N}}$ de sous-ensembles de $\mathbb{R}$ deux à deux disjoints alors

$$\mu(\cup_{j=0}^{+\infty} A_j) = \sum_{j=0}^{+\infty} \mu(A_j),$$

c'est une conséquence du théorème de Beppo-Levi. Si les ensembles de la suite $(A_j)_{j \in \mathbb{N}}$ ne sont pas disjoints deux à deux alors

$$\mu(\cup_{j=0}^{+\infty} A_j) \leq \sum_{j=0}^{+\infty} \mu(A_j).$$

Définition 21 (Ensembles négligeables). *On dit qu'un ensemble A de $\mathbb{R}$ est négligeable si sa mesure de Lebesgue est nulle.*

Exemple 30. *Voici quelques exemples d'ensembles négligeables de $\mathbb{R}$:*

- si $A = \{a\}$ alors $\mu(A) = 0$. En effet, pour $j \geq 1$, en posant $A_j =]a - 1/j, a + 1/j[$ alors $A \subset A_j$ et donc $\mu(A) \leq \mu(A_j) = 2/j$ pour tout $j \geq 1$. Donc $\mu(A) = 0$,
- si $A = \{a_1, \dots, a_n\}$ alors $\mu(A) = \sum_{j=1}^n \mu(a_j) = 0$,
- si $A = \bigcup_{j=1}^{+\infty} A_j$ avec $\mu(A_j) = 0$ pour tout $j \geq 1$ et $A_j \cap A_i = \emptyset$ pour tout $i \neq j$. Alors $\mu(A) = \sum_{j=1}^{+\infty} \mu(A_j) = 0$. En particulier les ensembles $\mathbb{N}$, $\mathbb{Z}$ et $\mathbb{Q}$ sont négligeables. Grossièrement on dit que la mesure de Lebesgue « ne charge pas les points ».

Définition 22 (Propriété vraie presque partout). *On dit qu'une propriété $P(x)$ dépendant d'un point x est vraie presque partout (noté p.p.) si et seulement si l'ensemble pour lequel $P(x)$ n'est pas vérifiée est de mesure nulle.*

Exemple 31. *Voici trois exemples.*

- la fonction (4.1) est continue presque partout (elle est continue sauf en $x = 0$).
- la fonction (4.2) est nulle presque partout. En effet cette fonction n'est pas nulle sur $\mathbb{Q} \cap [0, 1] \subset \mathbb{Q}$ et $\mu(\mathbb{Q} \cap [0, 1]) \leq \mu(\mathbb{Q}) = 0$. Ainsi

$$\int_{\mathbb{R}} \mathbf{1}_{\mathbb{Q} \cap [0, 1]}(x) dx = 0.$$

- Soit f une fonction de $\mathbb{R}$ dans $\overline{\mathbb{R}}_+$ alors on a l'équivalence

$$\int_{\mathbb{R}} f(x) dx = 0 \quad \Leftrightarrow \quad f(x) = 0 \text{ p.p.}$$

On peut bien entendu généraliser les notions précédentes dans $\mathbb{R}^n$ pour $n \geq 1$.

4.2.2 Fonctions Lebesgue intégrables

Nous sommes à présent en mesure de définir la notion de fonction intégrable au sens de Lebesgue.

Définition 23. *Soit $f : \mathbb{R} \rightarrow \mathbb{C}$ une fonction. On dit que f est Lebesgue intégrable (ou plus simplement intégrable) sur $\mathbb{R}$ si*

$$\int_{\mathbb{R}} |f(x)| dx < +\infty.$$

On va noter $L^1(\mathbb{R})$ l'ensemble suivant

$$L^1(\mathbb{R}) = \left\{ f : \mathbb{R} \rightarrow \mathbb{C} : \int_{\mathbb{R}} |f(x)| dx < +\infty \right\}.$$

Soit A une partie de $\mathbb{R}$ et $f : A \rightarrow \mathbb{C}$. On dit que f est Lebesgue intégrable sur A si

$$\int_A |f(x)| dx < +\infty,$$

et on note l'ensemble des fonctions Lebesgue intégrable sur A par $L^1(A)$.

L'idée derrière cette définition est la suivante. Soit $f : \mathbb{R} \rightarrow \mathbb{R}$ une fonction de $L^1(\mathbb{R})$. On peut écrire cette fonction sous la forme $f = f_+ - f_-$ avec $f_+(x) = \sup\{f(x), 0\}$ et $f_-(x) = \sup\{-f(x), 0\}$. Les fonctions f_+ et f_- sont positives et $|f| = f_+ + f_-$. Ainsi si f est intégrable alors par monotonie de l'intégrale de Lebesgue pour les fonctions positives on a

$$\begin{aligned} f_+(x) \leq |f(x)| \text{ p.p.} &\Rightarrow \int_{\mathbb{R}} f_+(x) dx \leq \int_{\mathbb{R}} |f(x)| dx < +\infty, \\ f_-(x) \leq |f(x)| \text{ p.p.} &\Rightarrow \int_{\mathbb{R}} f_-(x) dx \leq \int_{\mathbb{R}} |f(x)| dx < +\infty. \end{aligned}$$

On définit ensuite l'intégrale de f via la formule

$$\int_{\mathbb{R}} f(x) dx = \int_{\mathbb{R}} f_+(x) dx - \int_{\mathbb{R}} f_-(x) dx.$$

De même si $f : \mathbb{R} \rightarrow \mathbb{C}$ est intégrable alors on définit son intégrale par

$$\int_{\mathbb{R}} f(x) dx = \int_{\mathbb{R}} \operatorname{Re}(f(x)) dx + i \int_{\mathbb{R}} \operatorname{Im}(f(x)) dx.$$

Exemple 32. En pratique pour montrer qu'une fonction est Lebesgue intégrable il suffit de montrer qu'elle est majorée en module par une fonction positive d'intégrale finie. Donnons plusieurs exemples.

— La fonction $\mathbf{1}_{\mathbb{Q} \cap [0,1]}$ est Lebesgue intégrable, car

$$\mathbf{1}_{\mathbb{Q} \cap [0,1]}(x) \leq \mathbf{1}_{\mathbb{Q}}(x) \quad \forall x \in \mathbb{R},$$

et $\mathbb{Q}$ est négligeable.

— Une fonction f continue définie sur $\mathbb{R}$ et non nulle sur un intervalle fermé et borné $[a, b]$ de $\mathbb{R}$ est Lebesgue intégrable. En effet on a

$$\int_{\mathbb{R}} |f(x)| dx = \int_{\mathbb{R}} |f(x)| \mathbf{1}_{[a,b]}(x) dx = \int_a^b |f(x)| dx \leq \sup_{x \in [a,b]} |f(x)| (b - a) < +\infty.$$

- la fonction $f : x \mapsto x^2$ n'est pas Lebesgue intégrable sur $\mathbb{R}$, ni sur $\mathbb{R}_+$, ni sur $\mathbb{R}_-$. Par contre sa restriction à tout intervalle bornée I de $\mathbb{R}$ est intégrable, i.e., $f \in L^1(I)$.
- Pour tout $\xi \in \mathbb{R}$, la fonction $f : x \mapsto e^{-(a+2i\pi\xi)x} \mathbf{1}_{[0,+\infty[}(x)$ est dans $L^1(\mathbb{R})$, car pour tout $x \in \mathbb{R}$ on a

$$|e^{-(a+2i\pi\xi)x} \mathbf{1}_{[0,+\infty[}(x)| \leq e^{-ax} \mathbf{1}_{[0,+\infty[}(x).$$

Or on a déjà vu que la fonction dans le membre de droite est intégrable sur $\mathbb{R}$.

Voici une proposition qui donne un exemple de fonction Lebesgue intégrable :

Proposition 47. *Soit f une fonction continue (ou continue par morceaux) sur un sous-ensemble non vide et borné I de $\mathbb{R}$, alors f est une fonction Lebesgue intégrable sur I , i.e., $f \in L^1(I)$.*

On va admettre que cette « nouvelle » notion d'intégrale vérifie les propriétés suivantes :

Proposition 48. *L'intégrale de Lebesgue vérifie :*

1. Pour tout $f, g \in L^1(\mathbb{R})$ et $\alpha, \beta \in \mathbb{C}$

$$\int_{\mathbb{R}} (\alpha f(x) + \beta g(x)) dx = \alpha \int_{\mathbb{R}} f(x) dx + \beta \int_{\mathbb{R}} g(x) dx.$$

2. Si f et g appartiennent à $L^1(\mathbb{R})$ avec f et g à valeurs réelles vérifiant $f(x) \leq g(x)$ p.p. alors

$$\int_{\mathbb{R}} f(x) dx \leq \int_{\mathbb{R}} g(x) dx$$

3. Si $f \in L^1(\mathbb{R})$ alors

$$\left| \int_{\mathbb{R}} f(x) dx \right| \leq \int_{\mathbb{R}} |f(x)| dx,$$

et si f et $g \in L^1(\mathbb{R})$ alors

$$\int_{\mathbb{R}} |f(x) + g(x)| dx \leq \int_{\mathbb{R}} |f(x)| dx + \int_{\mathbb{R}} |g(x)| dx.$$

4. Si f et $g \in L^1(\mathbb{R})$ avec $f(x) = g(x)$ pour presque tout $x \in \mathbb{R}$ alors

$$\int_{\mathbb{R}} f(x) dx = \int_{\mathbb{R}} g(x) dx.$$

Remarque importante 2. *Il est important de retenir le point (4) de la proposition précédente. En particulier, on ne change pas la valeur de l'intégrale d'une fonction de L^1 en modifiant cette fonction en un nombre au plus dénombrable de points. Par exemple la fonction identiquement nulle sur $\mathbb{R}$ est égale presque partout à la fonction $\mathbf{1}_{\mathbb{Q}}$. Or on peut munir l'espace vectoriel $L^1(\mathbb{R})$ (ou $L^1(I)$ pour I un intervalle non vide de $\mathbb{R}$) de la « norme »*

$$\|f\|_{L^1(\mathbb{R})} = \int_{\mathbb{R}} |f(x)| dx, \quad \forall f \in L^1(\mathbb{R}).$$

Cette application est en fait une semi-norme car elle vérifie tous les axiomes d'une norme sauf que

$$\|f\|_{L^1(\mathbb{R})} = 0 \quad \Rightarrow \quad f(x) = 0 \quad \text{pour presque tout } x \in \mathbb{R}.$$

Ainsi pour que l'application $\|\cdot\|_{L^1(\mathbb{R})}$ soit une « vraie » norme il faut identifier les fonctions égales presque partout. Nous ne considérerons pas ce détail technique dans la suite et on parlera de norme. Il faut cependant retenir que pour une fonction $f \in L^1(\mathbb{R})$, on ne peut pas a priori donner un sens à $f(x)$ pour $x \in \mathbb{R}$ (sauf si par exemple f est continue). En particulier en pratique si $f, g \in L^1(\mathbb{R})$ on dira par exemple que

$$|f(x)| \leq g(x) \quad \text{pour presque tout } x \in \mathbb{R}.$$

Pour conclure cette section énonçons deux résultats (admis) permettant de faire le lien entre intégrale de Riemann et intégrale de Lebesgue.

Théorème 49. *On a les propriétés suivantes :*

- Soit f une fonction Riemann intégrable sur un intervalle borné et non vide I de $\mathbb{R}$, alors f est Lebesgue intégrable sur I , i.e., $f \in L^1(I)$. De plus les deux intégrales coïncident.
- Soit f une fonction bornée sur un intervalle borné et non vide I de $\mathbb{R}$ alors f est Riemann intégrable sur $I \Leftrightarrow f$ est continue presque partout sur I .

4.2.3 Les théorèmes utiles venant la théorie de l'intégration de Lebesgue

La théorie d'intégration de Lebesgue permet de démontrer plusieurs résultats très utiles en pratique. On va énoncer et utiliser les trois résultats suivants :

- Théorème de convergence dominée (de Lebesgue),
- Théorème de continuité d'une intégrale à paramètre,
- Théorème de dérivation d'une intégrale à paramètre.

Ces résultats seront admis.

Théorème 50 (Convergence dominée). *Soit $(f_n)_{n \in \mathbb{N}}$ une suite de fonctions avec $f_n : I \rightarrow \mathbb{R}$ (ou $\mathbb{C}$). On suppose que cette suite vérifie*

1. $(f_n)_{n \in \mathbb{N}}$ converge (simplement) presque partout vers une fonction $f : I \rightarrow \mathbb{R}$,
2. il existe $g \in L^1(I)$ telle que pour tout $n \in \mathbb{N}$ on ait $|f_n(x)| \leq g(x)$ pour presque tout $x \in I$.

Alors $f \in L^1(I)$ et on a

$$\lim_{n \rightarrow +\infty} \int_I f_n(x) dx = \int_I f(x) dx.$$

Exemple 33. On considère la suite de fonctions $(f_n)_{n \in \mathbb{N}^*}$ avec $f_n : \mathbb{R} \rightarrow \mathbb{R}$ où

$$f_n(x) = \frac{\cos(nx)}{1 + n^2 x^2}, \quad \forall x \in \mathbb{R}.$$

Ici on remarque dans un premier temps que

$$\lim_{n \rightarrow +\infty} f_n(0) = \lim_{n \rightarrow +\infty} 1 = 1.$$

De plus si $x \in \mathbb{R}_*$ alors

$$|f_n(x)| \leq \frac{1}{1+n^2x^2} \rightarrow 0 \quad \text{quand } n \rightarrow +\infty.$$

Ainsi la suite de fonctions $(f_n)_{n \in \mathbb{N}^*}$ converge (simplement) presque partout vers la fonction f sur $\mathbb{R}$ avec

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

Il faut à présent majorer uniformément en n les fonctions f_n par une fonction $g \in L^1(\mathbb{R})$. Or pour tout $n \in \mathbb{N}^*$ et $x \in \mathbb{R}$ on a

$$|f_n(x)| \leq \frac{1}{1+x^2} = g(x).$$

Donc pour tout $n \in \mathbb{N}^*$ on a $|f_n(x)| \leq g(x)$ pour tout $x \in \mathbb{R}$ (et donc a fortiori pour presque tout $x \in \mathbb{R}$). De plus on remarque que

$$\int_{\mathbb{R}} |g(x)| dx = 2 \int_0^{+\infty} \frac{1}{1+x^2} dx = 2 [\arctan(x)]_0^{+\infty} = \pi < +\infty.$$

Donc $g \in L^1(\mathbb{R})$ et on peut appliquer la théorème de convergence dominée pour obtenir

$$\lim_{n \rightarrow +\infty} \int_{\mathbb{R}} f_n(x) dx = \int_{\mathbb{R}} f(x) dx = 0.$$

Ce qui conclut l'exemple. □

Exemple 34. Soit $f \in L^1(\mathbb{R})$, montrons que

$$\lim_{n \rightarrow +\infty} \int_{-n}^n f(x) dx = \int_{\mathbb{R}} f(x) dx. \quad (4.4)$$

Pour ce faire on considère la suite de fonctions $(f_n)_{n \geq 1}$ avec

$$f_n(x) = f(x) \mathbf{1}_{[-n,n]}(x) \quad x \in \mathbb{R}.$$

Alors pour presque tout $x \in \mathbb{R}$ on a $\lim_{n \rightarrow +\infty} f_n(x) = f(x)$ et également

$$|f_n(x)| \leq |f(x)| \quad \forall n \geq 1.$$

Comme f est supposée être dans l'espace $L^1(\mathbb{R})$ alors d'après le théorème de convergence dominée on obtient (4.4). □

On considère dans la suite une intégrale à paramètre, i.e., un objet de la forme

$$F(t) = \int_I f(t, x) dx,$$

où I est un intervalle de $\mathbb{R}$ et f est une fonction définie sur $J \times I$ (J intervalle de $\mathbb{R}$). Dans l'expression précédente il faut voir t comme un paramètre. Si $f \in L^1(\mathbb{R})$, on reprend l'exemple de la transformée de Fourier de f , i.e.,

$$\widehat{f}(\xi) = \int_{\mathbb{R}} f(x) e^{-2i\pi\xi x} dx \quad \xi \in \mathbb{R}.$$

La transformée de Fourier de f , notée $\widehat{f}$, est une intégrale à paramètre, le paramètre étant $\xi \in \mathbb{R}$ (la fréquence). Commençons par remarquer que si $f \in L^1(\mathbb{R})$ alors pour tout $\xi \in \mathbb{R}$ on a

$$\int_{\mathbb{R}} |f(x) e^{-2i\pi\xi x}| dx = \int_{\mathbb{R}} |f(x)| dx < +\infty,$$

et donc dans ce cas la transformée de Fourier $\widehat{f}$ de f est bien définie. Dans la suite on veut établir des conditions (simples) permettant d'affirmer que la fonction $\widehat{f}$ est continue. On a le résultat suivant :

Théorème 51 (Continuité intégrale à paramètre). *Soit F l'intégrale à paramètre*

$$F(t) = \int_I f(t, x) dx,$$

où $f : J \times I \rightarrow \mathbb{C}$ avec J et I des intervalles de $\mathbb{R}$. Si les conditions suivantes sont vérifiées :

- la fonction $t \mapsto f(t, x)$ est continue sur J pour presque tout $x \in I$,
- Il existe $g \in L^1(I)$ telle que pour tout $t \in J$ on ait

$$|f(t, x)| \leq g(x) \quad \text{pour presque tout } x \in I.$$

Alors la fonction F est continue sur J .

Exemple 35. On reprend l'exemple sur la transformée de Fourier d'une fonction de $L^1(\mathbb{R})$. Soit $f \in L^1(\mathbb{R})$ et $\widehat{f}$ sa transformée de Fourier avec

$$\widehat{f}(\xi) = \int_{\mathbb{R}} f(x) e^{-2i\pi\xi x} dx \quad \forall \xi \in \mathbb{R}.$$

On remarque dans un premier temps que pour presque tout $x \in \mathbb{R}$ la fonction

$$\xi \mapsto f(x) e^{-2i\pi\xi x},$$

est continue sur $\mathbb{R}$. De plus pour tout $\xi \in \mathbb{R}$ et presque tout $x \in \mathbb{R}$ on a

$$|f(x) e^{-2i\pi\xi x}| \leq |f(x)| \in L^1(\mathbb{R}).$$

Ainsi d'après le théorème de continuité des intégrales à paramètre, la fonction $\widehat{f}$ est continue sur $\mathbb{R}$ (même si f ne l'est pas!). $\square$

Théorème 52 (Dérivabilité intégrale à paramètre). *Soit F l'intégrale à paramètre*

$$F(t) = \int_I f(t, x) dx,$$

où $f : J \times I \rightarrow \mathbb{C}$ avec J et I des intervalles de $\mathbb{R}$. Si les conditions suivantes sont vérifiées :

- pour tout $t \in J$, la fonction $x \mapsto f(t, x) \in L^1(\mathbb{R})$,
- pour presque tout $x \in I$, la fonction $t \mapsto f(t, x)$ est dérivable sur J ,
- il existe $g \in L^1(I)$ telle que pour tout $t \in J$ on ait

$$\left| \frac{\partial f}{\partial t}(t, x) \right| \leq g(x) \quad \text{pour presque tout } x \in I.$$

Alors la fonction F est dérivable sur J et

$$\frac{dF}{dt}(t) = \int_I \frac{\partial f}{\partial t}(t, x) dx \quad \forall t \in J.$$

Exemple 36. On reprend l'exemple sur la transformée de Fourier d'une fonction $f \in L^1(\mathbb{R})$ et on suppose que $x \mapsto xf(x) \in L^1(\mathbb{R})$. La transformée de Fourier de f est donnée par

$$\widehat{f}(\xi) = \int_{\mathbb{R}} f(x) e^{-2i\pi\xi x} dx \quad \forall \xi \in \mathbb{R}.$$

On a déjà vu que si $f \in L^1(\mathbb{R})$ alors la fonction $x \mapsto f(x)e^{-2i\pi\xi x} \in L^1(\mathbb{R})$ pour tout $\xi \in \mathbb{R}$. De plus pour presque tout $x \in \mathbb{R}$, la fonction $\xi \mapsto f(x)e^{-2i\pi\xi x}$ est dérivable sur $\mathbb{R}$. Enfin pour tout $\xi \in \mathbb{R}$ et presque tout $x \in \mathbb{R}$ on a

$$\left| \frac{\partial}{\partial \xi} (f(x)e^{-2i\pi\xi x}) \right| = |-2i\pi x f(x)e^{-2i\pi\xi x}| = 2\pi |xf(x)| = g(x).$$

Par hypothèse la fonction $g \in L^1(\mathbb{R})$. Le théorème de dérivation d'une intégrale à paramètre permet alors d'affirmer que la fonction $\widehat{f}$ est dérivable sur $\mathbb{R}$ avec

$$\frac{d\widehat{f}}{d\xi}(\xi) = -2i\pi \int_{\mathbb{R}} xf(x)e^{-2i\pi\xi x} dx, \quad \forall \xi \in \mathbb{R}.$$

4.3 Les espaces L^p

Précédemment nous avons introduit l'espace vectoriel $L^1(\mathbb{R})$ (ou $L^1(I)$ pour I un intervalle non vide de $\mathbb{R}$) avec

$$L^1(\mathbb{R}) = \left\{ f : \mathbb{R} \rightarrow \mathbb{C} : \int_{\mathbb{R}} |f(x)| dx < +\infty \right\}.$$

On peut munir l'espace $L^1(\mathbb{R})$ de la norme :

$$\|f\|_{L^1(\mathbb{R})} = \int_{\mathbb{R}} |f(x)| dx, \quad f \in L^1(\mathbb{R}).$$

De manière similaire on définit l'espace vectoriel $L^2(\mathbb{R})$ par

$$L^2(\mathbb{R}) = \left\{ f : \mathbb{R} \rightarrow \mathbb{C} : \left(\int_{\mathbb{R}} |f(x)|^2 dx \right)^{1/2} < +\infty \right\},$$

avec pour norme

$$\|f\|_{L^2(\mathbb{R})} = \left(\int_{\mathbb{R}} |f(x)|^2 dx \right)^{1/2}, \quad f \in L^2(\mathbb{R}).$$

On a bien entendu une définition analogue si on remplace $\mathbb{R}$ par un intervalle I non vide de $\mathbb{R}$. Il est important de remarquer que la norme $L^2(\mathbb{R})$ est issue du produit scalaire (déjà rencontré) suivant

$$\langle f, g \rangle_{L^2(\mathbb{R})} = \int_{\mathbb{R}} f(x) \overline{g(x)} dx, \quad \forall f, g \in L^2(\mathbb{R}).$$

On a bien

$$\|f\|_{L^2(\mathbb{R})} = \left(\int_{\mathbb{R}} f(x) \overline{f(x)} dx \right)^{1/2} = \langle f, f \rangle_{L^2(\mathbb{R})}^{1/2}, \quad \forall f \in L^2(\mathbb{R}).$$

Exemple 37. Donnons deux exemples de fonctions dans L^2 .

- Soit $f : [a, b] \rightarrow \mathbb{R}$ ($a, b \in \mathbb{R}$ avec $a < b$) une fonction continue alors $f \in L^2(a, b)$. En effet on a

$$\int_a^b |f(x)|^2 dx \leq \max_{x \in]a, b[} |f(x)|^2 (b - a) < +\infty.$$

- Soit $f : x \in]1, +\infty[\mapsto 1/x$ alors $f \in L^2(1, +\infty)$. En effet

$$\int_1^{+\infty} |f(x)|^2 dx = \int_1^{+\infty} \frac{1}{x^2} dx = \left[-\frac{1}{x} \right]_1^{+\infty} = 1.$$

Remarquons que la fonction $f : x \in]1, +\infty[\mapsto 1/x$ vérifie $f \in L^2(1, +\infty)$ mais $f \notin L^1(1, +\infty)$ car

$$\int_1^{+\infty} |f(x)| dx = \int_1^{+\infty} \frac{1}{x} dx = [\log(x)]_1^{+\infty} = +\infty.$$

Ainsi il faut retenir que si I est un intervalle non borné de $\mathbb{R}$ alors l'espace $L^2(I)$ n'est pas un sous-espace vectoriel de $L^1(I)$, i.e., $L^2(I) \not\subset L^1(I)$. Dans le cas où I est un intervalle borné de $\mathbb{R}$ on a :

Proposition 53. *Si $I \subset \mathbb{R}$ est un intervalle borné alors $L^2(I) \subset L^1(I)$.*

Démonstration. On suppose que $I = [a, b]$ avec $-\infty < a < b < +\infty$ et on considère $f \in L^2(I)$ quelconque. Alors d'après l'inégalité de Cauchy-Schwarz on obtient

$$\int_a^b |f(x)| dx = \int_a^b |f(x)| \times 1 dx \leq \left(\int_a^b |f(x)|^2 dx \right)^{1/2} \left(\int_a^b dx \right)^{1/2}.$$

Donc

$$\int_a^b |f(x)| dx \leq (b-a)^{1/2} \|f\|_{L^2(a,b)} < +\infty.$$

Ce qui conclut la preuve. $\square$

En fait pour tout $1 \leq p < +\infty$ on définit l'espace $L^p(\mathbb{R})$ (ou plus généralement $L^p(I)$ avec I un intervalle de $\mathbb{R}$) par

$$L^p(\mathbb{R}) = \left\{ f : \mathbb{R} \rightarrow \mathbb{C} : \int_{\mathbb{R}} |f(x)|^p dx < +\infty \right\}.$$

On peut munir l'espace $L^p(\mathbb{R})$ de la norme :

$$\|f\|_{L^p(\mathbb{R})} = \left(\int_{\mathbb{R}} |f(x)|^p dx \right)^{1/p}, \quad f \in L^p(\mathbb{R}).$$

Enfin on définit l'espace $L^\infty(\mathbb{R})$ (ou $L^\infty(I)$) par

$$L^\infty(\mathbb{R}) = \{ f : \mathbb{R} \rightarrow \mathbb{C} : \|f\|_{L^\infty(\mathbb{R})} < +\infty \},$$

où $\|\cdot\|_{L^\infty(\mathbb{R})}$ est la norme :

$$\|f\|_{L^\infty(\mathbb{R})} = \inf_{M>0} \{ |f(x)| \leq M \text{ p.p. } x \in \mathbb{R} \}.$$

Proposition 54. *Si $I \subset \mathbb{R}$ est un intervalle **borné** alors $L^\infty(I) \subset L^p(I)$ pour tout $1 \leq p < +\infty$.*

Démonstration. Soit f une fonction quelconque dans $L^\infty(I)$. Soit également $1 \leq p < +\infty$ alors on a

$$\int_I |f(x)|^p dx \leq \|f\|_{L^\infty(I)}^p \mu(I),$$

où $\mu(I)$ est la mesure de Lebesgue de I , i.e., la longueur de I . Ainsi on obtient que

$$\left(\int_I |f(x)|^p dx \right)^{1/p} \leq \|f\|_{L^\infty(I)} \mu(I)^{1/p} < +\infty,$$

car la longueur de I est finie par hypothèse. Ce qui conclut la preuve du résultat. $\square$

4.4 To be or not to be in L^p ($p = 1$ or 2)

Dans cette section nous établissons un catalogue de fonctions de références appartenant ou non à $L^1(I)$ ou $L^2(I)$ ($I \subseteq \mathbb{R}$). En particulier dans la suite du cours on pourra utiliser ce catalogue pour proclamer qu'une fonction est dans L^1 ou L^2 sans démonstration. Comme il est impossible de faire une liste exhaustive de fonctions, **il est particulièrement important de comprendre les preuves ci-dessous afin de pouvoir les adapter au cas par cas**. Afin de clarifier au mieux la situation on sépare le cas I non borné du cas I borné.

4.4.1 Le cas I non borné

Ici on considère le cas où $I = \mathbb{R}$, $I = \mathbb{R}_+$, $I = \mathbb{R}_-$ ou $I = [1, +\infty)$ etc.

1. Soit f une fonction constante sur I définie par $f(x) = c$ ($c \in \mathbb{R}$). Alors f n'appartient ni à $L^1(I)$ ni à $L^2(I)$. En effet on a

$$\int_I |f(x)| dx = \int_I |c| dx = |c| \int_I dx = |c| \mu(I) = +\infty,$$

où $\mu(I)$ désigne la mesure de Lebesgue de I , i.e., dans notre cas la longueur de I . De même on a

$$\int_I |f(x)|^2 dx = c^2 \mu(I) = +\infty.$$

2. Soit f la fonction définie sur $\mathbb{R}_+$ par $f(x) = e^{ax}$ avec $a > 0$. Alors cette fonction n'appartient pas à $L^1(\mathbb{R}_+)$. En effet on a

$$\int_0^{+\infty} |f(x)| dx = \int_0^{+\infty} e^{ax} dx = \left[\frac{e^{ax}}{a} \right]_0^{+\infty} = +\infty$$

Ici le problème d'intégrabilité vient « du point » $+\infty$, on dit plus simplement que f n'est pas intégrable en $+\infty$

3. Soit f la fonction définie sur $\mathbb{R}_+$ par $f(x) = e^{-ax}$ avec $a > 0$. Alors cette fonction appartient à $L^1(\mathbb{R}_+)$. en effet on a

$$\int_0^{+\infty} |f(x)| dx = \int_0^{+\infty} e^{-ax} dx = \left[-\frac{e^{-ax}}{a} \right]_0^{+\infty} = \frac{1}{a} < +\infty.$$

4. Soit f la fonction définie sur $[1, +\infty)$ par $f(x) = 1/x^\alpha$ avec $\alpha > 1$. Alors cette fonction appartient à $L^1(1, +\infty)$ et à $L^2(1, +\infty)$. En effet on a

$$\int_1^{+\infty} \left| \frac{1}{x^\alpha} \right| dx = \int_1^{+\infty} x^{-\alpha} dx = \left[\frac{x^{-\alpha+1}}{-\alpha+1} \right]_1^{+\infty} = \left[-\frac{1}{(\alpha-1)x^{\alpha-1}} \right]_1^{+\infty} = \frac{1}{\alpha-1} < +\infty,$$

où on utilise le fait que $\lim_{x \rightarrow +\infty} 1/x^{\alpha-1} = 0$ car $\alpha > 1$. De même on remarque que

$$\int_1^{+\infty} \left| \frac{1}{x^\alpha} \right|^2 dx = \int_1^{+\infty} x^{-2\alpha} dx = \left[\frac{x^{-2\alpha+1}}{-2\alpha+1} \right]_1^{+\infty} = \left[-\frac{1}{(2\alpha-1)x^{2\alpha-1}} \right]_1^{+\infty} = \frac{1}{2\alpha-1} < +\infty.$$

On en déduit (par changement de variable) que la fonction $x \mapsto 1/x^\alpha$ appartient à $L^1(-\infty, 1)$ et $L^2(-\infty, 1)$ si $\alpha > 1$. En fait, plus généralement, pour tout intervalle I de la forme $[\varepsilon, +\infty)$ ou $(-\infty, \varepsilon]$ avec $\varepsilon > 0$ la fonction $x \mapsto 1/x^\alpha$ appartient à $L^1(I)$ et $L^2(I)$ si $\alpha > 1$.

5. soit f la fonction définie sur $[1, +\infty)$ par $f(x) = 1/x$. Alors cette fonction n'appartient pas à $L^1(1, +\infty)$. En effet on a

$$\int_1^{+\infty} |f(x)| dx = \int_1^{+\infty} \frac{1}{x} dx = [\log(x)]_1^{+\infty} = +\infty$$

Ici la fonction n'est pas intégrable en $+\infty$. Par contre f appartient à $L^2(1, +\infty)$ car

$$\int_1^{+\infty} |f(x)|^2 dx = \int_1^{+\infty} x^{-2} dx = \left[\frac{x^{-1}}{-1} \right]_1^{+\infty} = \left[-\frac{1}{x} \right]_1^{+\infty} = 1 < +\infty.$$

Ici la fonction est intégrable en $+\infty$.

6. Soit f la fonction définie sur $[1, +\infty)$ par $f(x) = 1/x^\alpha$ avec $0 < \alpha < 1$. Alors cette fonction n'appartient pas à $L^1(1, +\infty)$. En effet on a

$$\int_1^{+\infty} |f(x)| dx = \int_1^{+\infty} x^{-\alpha} dx = \left[\frac{x^{-\alpha+1}}{-\alpha+1} \right]_1^{+\infty} = \left[-\frac{x^{1-\alpha}}{(\alpha-1)} \right]_1^{+\infty} = +\infty,$$

où ici on utilise le fait que $\lim_{x \rightarrow +\infty} x^{1-\alpha} = +\infty$ car $\alpha < 1$ (donc $1 - \alpha > 0$). Ici la fonction n'est pas intégrable en $+\infty$.

4.4.2 Le cas I borné

Ici I est de longueur finie comme par exemple $I = [-1, 1]$, $I = [0, 10]$ etc.

1. Une fonction continue et bornée sur I est dans $L^\infty(I)$ et donc dans $L^1(I)$ et $L^2(I)$ d'après la Proposition 54.
2. Soit f la fonction définie sur $[0, 1]$ (ou $]0, 1[$) par $f(x) = 1/x^\alpha$ avec $0 < \alpha < 1$. Alors cette fonction appartient à $L^1(0, 1)$. En effet on a

$$\int_0^1 |f(x)| dx = \int_0^1 x^{-\alpha} dx = \left[\frac{x^{-\alpha+1}}{-\alpha+1} \right]_0^1 = \left[\frac{x^{1-\alpha}}{1-\alpha} \right]_0^1 = \frac{1}{1-\alpha} < +\infty,$$

ici on a utilisé le fait que $\lim_{x \rightarrow 0} x^{1-\alpha} = 0$ car $1 - \alpha > 0$. La fonction f est intégrable en 0.

3. Soit f la fonction définie sur $[0, 1]$ par $f(x) = 1/x$. Alors cette fonction n'appartient pas à $L^1(0, 1)$ car

$$\int_0^1 |f(x)| dx = \int_0^1 \frac{1}{x} dx = [\log(x)]_0^1 = +\infty.$$

Ici la fonction n'est pas intégrable en 0.

4. Soit f la fonction définie sur $[0, 1]$ par $f(x) = 1/x^\alpha$ avec $\alpha > 1$. Alors f n'appartient pas à $L^1(0, 1)$ (ni $L^2(0, 1)$). En effet on a

$$\int_0^1 |f(x)| dx = \int_0^1 x^{-\alpha} dx = \left[\frac{x^{-\alpha+1}}{-\alpha+1} \right]_0^1 = \left[-\frac{1}{(\alpha-1)x^{\alpha-1}} \right]_0^1 = +\infty,$$

car ici comme $\alpha-1 > 0$ alors $\lim_{x \rightarrow 0} 1/x^{\alpha-1} = +\infty$. La fonction f n'est pas intégrable en 0.

4.4.3 Résumé et exemples

Donnons à présent quelques exemples d'applications pour comprendre comment utiliser ce catalogue dans le cas d'une fonction ne rentrant pas exactement dans la Table 4.1. Voici plusieurs exemples :

- Soit f la fonction définie sur $\mathbb{R}_+$ par $f(x) = 1/(1+x)^\alpha$ avec $\alpha > 1$. Alors cette fonction appartient à $L^1(\mathbb{R}_+)$ et $L^2(\mathbb{R}_+)$. En effet la fonction $x \mapsto 1/(1+x)^\alpha$ est **continue et bornée** sur $[0, 1]$ donc dans $L^\infty(0, 1)$. D'après la Proposition 54 alors cette fonction appartient à $L^1(0, 1)$ et $L^2(0, 1)$. Par ailleurs un calcul similaire au cas (4) permet de montrer facilement que la fonction $x \mapsto 1/(1+x)^\alpha$ appartient à $L^1(1, +\infty)$ (et de même pour L^2). En conclusion la fonction $f \in L^1(\mathbb{R}_+)$ et $L^2(\mathbb{R}_+)$.
- On montre de la même manière que la fonction $x \mapsto 1/(1+|x|)^\alpha$ appartient à $L^1(\mathbb{R})$ et $L^2(\mathbb{R})$ si $\alpha > 1$.
- Par contre la fonction $f : x \mapsto 1/(1+x)^\alpha$ n'appartient pas à $L^1(\mathbb{R})$ ou $L^2(\mathbb{R})$ si $\alpha > 1$. En effet ici le problème d'intégrabilité vient du point -1 où le dénominateur de f s'annule. On se ramène ainsi au problème d'intégrabilité de la fonction $x \mapsto 1/x^\alpha$ en 0 avec $\alpha > 1$. Ainsi en reprenant un calcul déjà effectué on a par exemple

$$\int_{-1}^0 |f(x)| dx = \int_{-1}^0 (1+x)^{-\alpha} dx = \left[\frac{(1+x)^{-\alpha+1}}{-\alpha+1} \right]_{-1}^0 = \lim_{x \rightarrow -1} \frac{1}{(\alpha-1)(1+x)^{\alpha-1}} = +\infty,$$

car $\alpha > 1$. On effectue un calcul similaire pour montrer que $f \notin L^2(\mathbb{R})$.

- Soit f la fonction définie sur $\mathbb{R}_+$ par $f(x) = 1/(1+x^\alpha)$ pour $\alpha > 1$. Alors $f \in L^1(\mathbb{R}_+)$. En effet ici le dénominateur ne s'annule en aucun point de $\mathbb{R}_+$, il faut étudier l'intégrabilité de f en $+\infty$. Dans un premier temps on remarque que f est bornée

I	définition f	espaces
$\mu(I) < +\infty$	c une constante	$f \in L^1(I)$ et $f \in L^2(I)$
$\mu(I) = +\infty$	$c \neq 0$ une constante	$f \notin L^1(I)$ et $f \notin L^2(I)$
$\mu(I) < +\infty$	f continue et bornée	$f \in L^1(I)$ et $f \in L^2(I)$
$\mathbb{R}_+$	e^{ax} avec $a > 0$	$f \notin L^1(I)$ et $f \notin L^2(I)$
$\mathbb{R}_+$	e^{-ax} avec $a > 0$	$f \in L^1(I)$ et $f \in L^2(I)$
$[1, +\infty)$	$1/x^\alpha$ ($\alpha > 1$)	$f \in L^1(I)$ et $f \in L^2(I)$
$[1, +\infty)$	$1/x$	$f \notin L^1(I)$ et $f \in L^2(I)$
$[1, +\infty)$	$1/x^\alpha$ ($0 < \alpha < 1$)	$f \notin L^1(I)$
$[0, 1]$	$1/x^\alpha$ ($\alpha > 1$)	$f \notin L^1(I)$ et $f \notin L^2(I)$
$[0, 1]$	$1/x$	$f \notin L^1(I)$ et $f \notin L^2(I)$
$[0, 1]$	$1/x^\alpha$ ($0 < \alpha < 1$)	$f \in L^1(I)$

TABLE 4.1 – Appartenance à l'espace $L^1(I)$ ou $L^2(I)$ pour certaines fonctions de référence. On note $\mu(I)$ la mesure de Lebesgue de I , i.e., la longueur de I .

sur $[0, 1]$ donc $f \in L^\infty(0, 1)$ et a fortiori d'après la Proposition 54 dans $L^1(0, 1)$. Par ailleurs pour $x \geq 1$ alors $1 + x^\alpha \geq x^\alpha$ donc on obtient

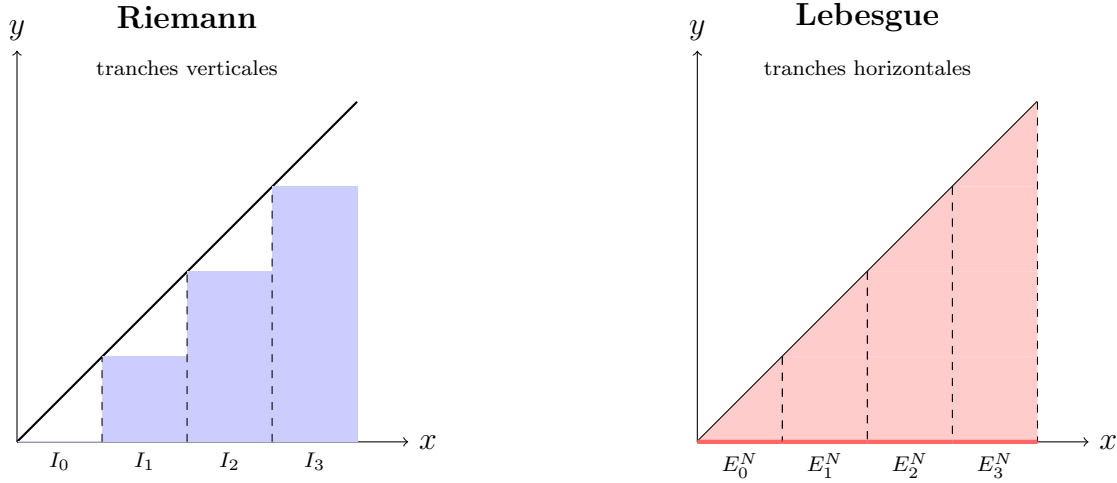
$$\int_1^{+\infty} |f(x)| dx = \int_1^{+\infty} \frac{1}{1 + x^\alpha} dx \leq \int_1^{+\infty} \frac{1}{x^\alpha} dx < +\infty,$$

car nous avons déjà vu que la fonction $x \mapsto 1/x^\alpha \in L^1(1, +\infty)$ si $\alpha > 1$.

4.5 Compléments

4.5.1 Quelques éléments de la construction de l'intégrale de Lebesgue

Pour présenter les idées principales de la théorie d'intégration de Lebesgue, nous allons considérer la fonction $f : x \in [0, 1] \mapsto x$, qui est intégrable au sens de Riemann, et comparer la méthode de construction de $\int_0^1 f(x) dx$ au sens de Riemann par « tranches verticales » et celle de Lebesgue par « tranches horizontales ». Bien entendu, comme la théorie de Lebesgue généralise celle de Riemann, si f est intégrable au sens de Riemann, elle doit être intégrable au sens de Lebesgue et la valeur de cette intégrale doit être identique.



Construction de Riemann. Une façon simple de construire l'intégrale au sens de Riemann de la fonction f est de considérer, pour un $N \geq 1$ fixé, la suite de points $x_k = k/N$ pour tout $k \in \{0, \dots, N\}$. Alors, comme déjà étudié en utilisant la méthode des rectangles à gauche, on peut dire que la valeur de $\int_0^1 f(x) dx$ est obtenue comme suit

$$\frac{1}{2} = \int_0^1 f(x) dx \approx \frac{1}{N} \sum_{k=0}^{N-1} f(x_k) = \frac{1}{N} \sum_{k=0}^{N-1} \frac{k}{N} = \frac{1}{N^2} \frac{N(N-1)}{2} = \frac{1}{2} \left(1 - \frac{1}{N}\right) \rightarrow \frac{1}{2},$$

lorsque $N \rightarrow +\infty$. Ainsi, l'idée de Riemann est de découper l'ensemble de départ de la fonction f , ici l'intervalle $[0, 1]$, en tranches verticales et de faire ensuite tendre la taille de ces tranches vers 0.

Construction de Lebesgue. À l'inverse, l'idée de Lebesgue est de découper l'ensemble d'arrivée de f , ou plus précisément l'image de f , en tranches. Pour expliquer cette idée, soit $N \geq 1$ fixé, alors pour tout $k \in \{0, \dots, 2^N - 1\}$ on considère les ensembles

$$E_k^N = \left\{ x \in [0, 1] : \frac{k}{2^N} < f(x) \leq \frac{k+1}{2^N} \right\} = f^{-1} \left(\left[\frac{k}{2^N}, \frac{k+1}{2^N} \right] \right).$$

Ici il faut faire attention, pour A un sous-ensemble de $[0, 1]$, on note $f^{-1}(A)$ la pré-image (ou image réciproque) de A par f (il ne faut pas confondre avec l'image de A par l'application réciproque f^{-1}). Pour mieux comprendre les ensembles E_k^N , considérons par exemple le cas $k = 0$, on a

$$E_0^N = \left\{ x \in [0, 1] : 0 < f(x) \leq \frac{1}{2^N} \right\} = \left\{ x \in [0, 1] : 0 < x \leq \frac{1}{2^N} \right\} = \left] 0, \frac{1}{2^N} \right],$$

et plus généralement, via la définition de f , on a

$$E_k^N = \left] \frac{k}{2^N}, \frac{k+1}{2^N} \right], \quad \forall k \in \{0, \dots, 2^N - 1\}.$$

On peut par exemple noter que, au singleton $\{0\}$ près, les ensembles E_k^N forment une partition de l'image de f . Une fois ces tranches horizontales de l'image de f définies, on introduit la suite de fonctions « simples » suivantes :

$$g_N(x) = \sum_{k=0}^{2^N-1} \frac{k}{2^N} \mathbf{1}_{E_k^N}(x), \quad \forall x \in [0, 1], \quad (4.5)$$

où

$$\mathbf{1}_{E_k^N}(x) = \begin{cases} 1 & \text{si } x \in E_k^N, \\ 0 & \text{si } x \notin E_k^N. \end{cases}$$

On remarque par définition de la fonction $\mathbf{1}_{E_k^N}$ que

$$\int_0^1 \mathbf{1}_{E_k^N}(x) dx = \int_{E_k^N} 1 dx = \ell(E_k^N) = \frac{1}{2^N}, \quad (4.6)$$

où $\ell(E_k^N)$ désigne la longueur de E_k^N . À présent, comme dans le cas de la construction de Riemann, on va approcher $\int_0^1 f(x) dx$ via la suite de fonctions simples $(g_N)_{N \geq 1}$ comme suit :

$$\int_0^1 f(x) dx \approx \int_0^1 g_N(x) dx = \int_0^1 \left(\sum_{k=0}^{2^N-1} \frac{k}{2^N} \mathbf{1}_{E_k^N}(x) \right)$$

$$\begin{aligned}
&= \sum_{k=0}^{2^N-1} \frac{k}{2^N} \int_0^1 \mathbf{1}_{E_k^N}(x) dx \\
&= \sum_{k=0}^{2^N-1} \frac{k}{2^N} \ell(E_k^N) \\
&= \frac{1}{2^{2N}} \sum_{k=0}^{2^N-1} k = \frac{1}{2} \left(1 - \frac{1}{2^N}\right) \rightarrow \frac{1}{2},
\end{aligned}$$

quand $N \rightarrow +\infty$. On dira alors que f (une fonction positive) est intégrable au sens de Lebesgue si il existe une suite de fonction simples du type (4.5) telle que la suite des valeurs de ces intégrales converge vers une valeur finie. De façon générale la construction de l'intégrale d'une fonction positive f suit le processus suivant :

- On découpe l'image de f en tranches horizontales.
- On mesure les ensemble de ce découpage (comme dans (4.6)).
- On construit des fonctions simples du type (4.5).

Ce processus est beaucoup plus souple que celui de Riemann et peut se généraliser à des fonctions définies sur des ensembles bien plus généraux que $\mathbb{R}$ (ou $\mathbb{R}^d$). Il faut cependant être capable de définir une notion de mesure, de découper l'image de la fonction considérée de manière adéquate etc.

Restons dans le cas simple des fonctions définies sur $\mathbb{R}$ (ou un sous-ensemble de $\mathbb{R}$) et admettons l'existence d'une application, notée μ , à valeurs dans $\mathbb{R}_+ \cup \{+\infty\}$, appelée mesure de Lebesgue. Cette mesure est définie sur un sous-ensemble de $\mathcal{P}(\mathbb{R})$, l'ensemble des parties de $\mathbb{R}$, noté $\mathcal{L}(\mathbb{R})$ appelé tribu de Lebesgue. Pour des raisons complexes la mesure de Lebesgue est une généralisation de la notion de longueur classique et associe à tout ensemble de $\mathcal{L}(\mathbb{R})$ une valeur dans $\mathbb{R}_+ \cup \{+\infty\}$. Par exemple

$$\mu([0, 1]) = \mu(]0, 1[) = \mu([0, 1[) = \mu([0, 1]) = 1 \quad \text{et} \quad \mu(]0, +\infty[) = +\infty.$$

Il faut retenir que la mesure de Lebesgue d'un point est nulle. Ainsi, dans le processus indiqué précédemment, pour pouvoir parler de la mesure des tranches horizontales de l'image de la fonction il faut que ces ensembles appartiennent à $\mathcal{L}(\mathbb{R})$. Si on reprend l'exemple du début de la section, il faut que $E_k^N = f^{-1}(]k 2^{-N}, (k+1) 2^{-N}[)$ appartiennent à $\mathcal{L}(\mathbb{R})$. Le fait que ces ensembles appartiennent ou non à $\mathcal{L}(\mathbb{R})$ dépend entièrement de la fonction. On arrive en fait à la notion de fonction (Lebesgue) mesurable. On dit qu'une fonction positive f est Lebesgue mesurable ou plus simplement mesurable, si pour tout $A \in \mathcal{L}(\mathbb{R})$, alors

$$f^{-1}(A) = \{x \in \mathbb{R} : f(x) \in A\} \in \mathcal{L}(\mathbb{R}).$$

En pratique, toutes les fonction rencontrées dans ce cours (et ailleurs) sont mesurables.

On comprend à présent que pour être capable de définir un découpage adéquat de l'image d'une fonction et pour donner une mesure à ces ensembles il faut se restreindre aux fonctions mesurables. Il reste alors à expliquer comment approcher des fonctions mesurables par des fonctions simples. Commençons par définir plus précisément la notion de

fonction simple dans ce contexte.

Définition. Soit $A \in \mathcal{L}(\mathbb{R})$. On dit qu'une fonction $f : A \rightarrow \mathbb{R}$ est étagée positive s'il existe un nombre **fini** de réels positifs $\alpha_1, \dots, \alpha_m$ et de sous-ensembles $A_1, \dots, A_m \in \mathcal{L}(\mathbb{R})$ de A tels que

$$f(x) = \sum_{i=1}^m \alpha_i \mathbf{1}_{A_i}(x). \quad (4.7)$$

Commençons par remarquer que si f est étagée alors f est mesurable. Par ailleurs, si f est étagée positive de la forme (4.7) alors l'intégrale de Lebesgue de f est définie par

$$\int_A f(x) dx = \sum_{i=1}^m \alpha_i \mu(A_i).$$

On dit que f est intégrable au sens de Lebesgue si $\int_A f(x) dx$ est fini.

Théorème d'approximation par des fonctions étagées. Si f est une fonction positive mesurable définie sur $\mathbb{R}$ (ou A un sous-ensemble de $\mathbb{R}$), alors il existe une suite de fonctions étagées positives (f_n) qui converge (simplement) vers f sur $\mathbb{R}$ (ou A).

Définition. Soit $A \in \mathcal{L}(\mathbb{R})$ et $f : A \rightarrow \mathbb{R}_+$ une fonction mesurable au sens de Lebesgue. On dit que f est Lebesgue intégrable sur A si le nombre

$$\sup \left\{ \int_A g(x) dx : g \text{ étagée avec } g \leq f \right\},$$

est **fini**. Dans ce cas, on appelle intégrale de Lebesgue ce nombre et on le note $\int_A f(x) dx$.

Chapitre 5

Transformée de Fourier

Résumé. Dans ce chapitre nous allons étudier la transformée de Fourier d'un signal dans L^1 ou L^2 .

Références. Pour ce chapitre la plupart des résultats et preuves viennent de
— C. Gasquet et P. Witomski. *Analyse de Fourier et applications*, Dunod (1990), chapitre 3 (disponible à la BUTC).

5.1 Transformée de Fourier dans L^1

Avant d'introduire rigoureusement la notion de transformée de Fourier donnons dans un premier temps une justification informelle. Il faut garder à l'esprit que dans ce chapitre nous allons étudier des signaux non périodiques. L'analyse de ces signaux ne peut-être effectuée via les séries de Fourier. Cependant en interprétant un signal non périodique comme un « morceau » d'un signal périodique dont la période tend vers l'infini nous allons réussir à deviner l'expression de la transformée de Fourier. Pour ce faire considérons f une fonction a -périodique et régulière vérifiant

$$f(x) = \sum_{n=-\infty}^{+\infty} c_n(f) e^{2i\pi n \frac{x}{a}} \quad \forall x \in \mathbb{R}.$$

Si à présent f désigne un signal non périodique vu comme un morceau d'un signal a -périodique avec $a \rightarrow +\infty$ alors on a

$$f(x) = \lim_{a \rightarrow +\infty} \sum_{n=-\infty}^{+\infty} \left(\frac{1}{a} \int_{-a/2}^{a/2} f(t) e^{-2i\pi \frac{n}{a} t} dt \right) e^{2i\pi \frac{n}{a} x}.$$

En notant $\Delta\xi = 1/a$ et $\xi_n = n\Delta\xi$ ($n \in \mathbb{N}$) alors on peut réécrire cette expression sous la forme

$$f(x) = \lim_{\substack{a \rightarrow +\infty \\ (\Delta\xi \rightarrow 0)}} \sum_{n=-\infty}^{+\infty} \left(\Delta\xi \int_{-a/2}^{a/2} f(t) e^{-2i\pi \xi_n t} dt \right) e^{2i\pi \xi_n x}$$

$$\begin{aligned}
&= \int_{-\infty}^{+\infty} \left(\int_{-\infty}^{+\infty} f(t) e^{-2i\pi\xi t} dt \right) e^{2i\pi\xi x} d\xi \\
&= \int_{-\infty}^{+\infty} \widehat{f}(\xi) e^{2i\pi\xi x} d\xi.
\end{aligned}$$

Ce calcul informel permet de déduire deux relations importantes :

$$\begin{aligned}
\widehat{f}(\xi) &= \int_{\mathbb{R}} f(x) e^{-2i\pi\xi x} dx \quad \text{expression de la transformée de Fourier,} \\
f(x) &= \int_{\mathbb{R}} \widehat{f}(\xi) e^{2i\pi\xi x} d\xi \quad \text{expression de la transformée de Fourier inverse.}
\end{aligned}$$

5.1.1 Définition formelle de la transformée de Fourier dans L^1

Définition 24. Pour une fonction $f \in L^1(\mathbb{R})$, on appelle transformée de Fourier de f , la fonction $\widehat{f} : \mathbb{R} \rightarrow \mathbb{C}$ définie par

$$\widehat{f}(\xi) = \int_{\mathbb{R}} f(x) e^{-2i\pi\xi x} dx, \quad \forall \xi \in \mathbb{R}. \quad (5.1)$$

On écrira (symboliquement) $\widehat{f} = \mathcal{F}f$ ou $\widehat{f}(\xi) = \mathcal{F}\{f(x)\}$.

Il est important de souligner que si $f \in L^1(\mathbb{R})$ alors la formule (5.1) possède un sens. En effet nous savons d'après le chapitre précédent que pour tout $\xi \in \mathbb{R}$, la fonction $x \mapsto f(x) e^{-2i\pi\xi x}$ appartient à $L^1(\mathbb{R})$. Ainsi d'après la Proposition 2 du Chapitre 4, on a

$$|\widehat{f}(\xi)| = \left| \int_{\mathbb{R}} f(x) e^{-2i\pi\xi x} dx \right| \leq \int_{\mathbb{R}} |f(x)| dx = \|f\|_{L^1(\mathbb{R})} \quad \forall \xi \in \mathbb{R}.$$

Exemple 38. Dans la suite du cours on va noter H la fonction de Heaviside, à savoir

$$H(x) = \begin{cases} 1 & \text{si } x > 0 \\ 0 & \text{sinon.} \end{cases}$$

Soit f la fonction définie par $f(x) = e^{-ax} H(x)$ avec $a \in \mathbb{C}$ et $\operatorname{Re}(a) > 0$. On obtient alors

$$\widehat{f}(\xi) = \int_{\mathbb{R}} f(x) e^{-2i\pi\xi x} dx = \int_0^{+\infty} e^{-x(a+2i\pi\xi)} dx = \left[-\frac{e^{-x(a+2i\pi\xi)}}{a+2i\pi\xi} \right]_0^{+\infty} = \frac{1}{a+2i\pi\xi}.$$

Justifions (une fois pour toute) le calcul de la limite $\lim_{x \rightarrow +\infty} e^{-x(a+2i\pi\xi)} / (a+2i\pi\xi) = 0$. Pour ce faire comme $a \in \mathbb{C}$ alors $a = a_1 + ia_2$ avec $a_1 > 0$. On obtient ainsi

$$e^{-x(a+2i\pi\xi)} = e^{-x(a_1+i(a_2+2\pi\xi))} = e^{-xa_1} e^{-ix(a_2+2\pi\xi)}.$$

Comme $a_1 > 0$ alors $\lim_{x \rightarrow +\infty} e^{-a_1 x} = 0$ et de plus

$$|e^{-ix(a_2+2\pi\xi)}| = 1.$$

Donc pour justifier que $\lim_{x \rightarrow +\infty} e^{-x(a+2i\pi\xi)} / (a + 2i\pi\xi) = 0$, on s'intéresse à la quantité

$$|e^{-x(a+2i\pi\xi)} / (a + 2i\pi\xi) - 0| = e^{-a_1x} \rightarrow 0 \text{ quand } x \rightarrow +\infty.$$

En résumé, pour tout $a \in \mathbb{C}$ avec $\operatorname{Re}(a) > 0$, on a

$$e^{-ax} H(x) \xrightarrow{\mathcal{F}} \frac{1}{a + 2i\pi\xi}.$$

Dans l'exemple précédent il est important de remarquer que la fonction $\hat{f} \notin L^1(\mathbb{R})$. Ainsi il est légitime de se demander à quel espace appartient $\mathcal{F}f$.

Théorème 55 (Théorème de Riemann-Lebesgue). *Soit $f \in L^1(\mathbb{R})$, alors $\hat{f}$ est une fonction continue et bornée sur $\mathbb{R}$ et de plus*

$$\lim_{|\xi| \rightarrow +\infty} |\hat{f}(\xi)| = 0. \quad (5.2)$$

Démonstration. Si $f \in L^1(\mathbb{R})$, alors on a déjà vu au chapitre précédent que $\hat{f}$ est une fonction continue d'après le théorème de continuité des intégrales à paramètres. De plus comme $f \in L^1(\mathbb{R})$ alors on a pour tout $\xi \in \mathbb{R}$

$$|\hat{f}(\xi)| \leq \int_{\mathbb{R}} |f(x)| dx = \|f\|_{L^1(\mathbb{R})},$$

et ainsi $\hat{f}$ est une fonction bornée. On admet le point (5.2). Ce qui conclut la preuve. $\square$

5.1.2 Propriétés de la transformée de Fourier

Dans cette partie nous allons énoncer et démontrer différentes règles permettant de calculer en pratique la transformée de Fourier d'une fonction de $L^1(\mathbb{R})$.

Proposition 56. *Soient f et $g \in L^1(\mathbb{R})$ et $\alpha, \beta \in \mathbb{C}$ alors*

$$\mathcal{F}\{\alpha f(x) + \beta g(x)\} = \alpha \mathcal{F}\{f(x)\} + \beta \mathcal{F}\{g(x)\}.$$

Démonstration. C'est une conséquence directe de la linéarité de l'intégrale. $\square$

Proposition 57 (Translation). *Soient $f \in L^1(\mathbb{R})$ et $c \in \mathbb{R}$, alors*

$$\mathcal{F}\{f(x - c)\} = e^{-2i\pi\xi c} \hat{f}(\xi), \quad \forall \xi \in \mathbb{R}.$$

Démonstration. On a, en utilisant un changement de variable, les égalités suivantes

$$\begin{aligned} \mathcal{F}\{f(x - c)\} &= \int_{\mathbb{R}} f(x - c) e^{-2i\pi\xi x} dx \stackrel{y=x-c}{=} \int_{\mathbb{R}} f(y) e^{-2i\pi\xi(y+c)} dy \\ &= e^{-2i\pi\xi c} \int_{\mathbb{R}} f(y) e^{-2i\pi\xi y} dy = e^{-2i\pi\xi c} \hat{f}(\xi), \end{aligned}$$

ce qui conclut la preuve. $\square$

Proposition 58 (Modulation). *Soient $f \in L^1(\mathbb{R})$ et $\xi_0 \in \mathbb{R}$, alors*

$$\mathcal{F}\{e^{2i\pi\xi_0 x} f(x)\} = \widehat{f}(\xi - \xi_0), \quad \forall \xi \in \mathbb{R}.$$

Démonstration. On a

$$\begin{aligned} \mathcal{F}\{e^{2i\pi\xi_0 x} f(x)\} &= \int_{\mathbb{R}} e^{2i\pi\xi_0 x} f(x) e^{-2i\pi\xi x} dx \\ &= \int_{\mathbb{R}} f(x) e^{-2i\pi(\xi - \xi_0)x} dx = \widehat{f}(\xi - \xi_0), \end{aligned}$$

ce qui conclut la preuve du résultat. $\square$

Proposition 59 (Changement d'échelle). *Soient $f \in L^1(\mathbb{R})$ et $c \in \mathbb{R}^*$, alors*

$$\mathcal{F}\{f(cx)\} = \frac{1}{|c|} \widehat{f}\left(\frac{\xi}{c}\right), \quad \forall \xi \in \mathbb{R}. \quad (5.3)$$

De plus, la transformée de Fourier d'une fonction paire (resp. impaire) est paire (resp. impaire).

Démonstration. Supposons dans un premier temps que $c > 0$, alors on a

$$\mathcal{F}\{f(cx)\} = \int_{-\infty}^{+\infty} f(cx) e^{-2i\pi\xi x} dx.$$

On considère alors le changement de variable $y = cx$ et on a

$$\mathcal{F}\{f(cx)\} = \int_{-\infty}^{+\infty} f(y) e^{-2i\pi\xi \frac{y}{c}} \frac{dy}{c} = \frac{1}{c} \int_{\mathbb{R}} f(y) e^{-2i\pi \frac{\xi}{c} y} dy = \frac{1}{c} \widehat{f}\left(\frac{\xi}{c}\right).$$

Si maintenant $c < 0$ alors le changement de variable $y = cx$ inverse les bornes d'intégration et on a

$$\mathcal{F}\{f(cx)\} = \int_{+\infty}^{-\infty} f(y) e^{-2i\pi\xi \frac{y}{c}} \frac{dy}{c} = -\frac{1}{c} \int_{-\infty}^{+\infty} f(y) e^{-2i\pi \frac{\xi}{c} y} dy = -\frac{1}{c} \widehat{f}\left(\frac{\xi}{c}\right).$$

Dans tous les cas on retrouve la formule (5.3). Si f est paire, en appliquant la formule (5.3) avec $c = -1$ on a

$$\mathcal{F}\{f(-x)\} = \widehat{f}(-\xi),$$

or si f est paire alors $\mathcal{F}\{f(-x)\} = \mathcal{F}\{f(x)\} = \widehat{f}(\xi)$ et donc

$$\widehat{f}(-\xi) = \widehat{f}(\xi).$$

En argumentant de la même manière si f est impaire on conclut la preuve de la proposition. $\square$

Proposition 60. Soit $f \in L^1(\mathbb{R})$ alors

$$\mathcal{F}\{\bar{f}(x)\} = \widehat{\bar{f}}(-\xi) = \overline{\int_{\mathbb{R}} f(x) e^{2i\pi\xi x} dx}, \quad \forall \xi \in \mathbb{R}.$$

Démonstration. On a directement

$$\mathcal{F}\{\bar{f}(x)\} = \int_{\mathbb{R}} \bar{f}(x) e^{-2i\pi\xi x} dx = \overline{\int_{\mathbb{R}} f(x) e^{2i\pi\xi x} dx} = \widehat{\bar{f}}(-\xi).$$

Ce qui achève la preuve de la proposition. $\square$

Proposition 61. Soit $f \in L^1(\mathbb{R})$. On suppose que f est dérivable et que $f' \in L^1(\mathbb{R})$ alors

$$\mathcal{F}\{f'(x)\} = 2i\pi\xi \mathcal{F}\{f(x)\} = 2i\pi\xi \widehat{f}(\xi), \quad \forall \xi \in \mathbb{R}.$$

De plus si f admet des dérivées jusqu'à l'ordre $n \geq 1$ avec $f^{(k)} \in L^1(\mathbb{R})$ pour tout $1 \leq k \leq n$, alors

$$\mathcal{F}\{f^{(n)}\} = (2i\pi\xi)^n \mathcal{F}\{f(x)\} = (2i\pi\xi)^n \widehat{f}(\xi), \quad \forall \xi \in \mathbb{R}.$$

Remarque 62. La Proposition (61) est très utile en pratique pour résoudre de nombreuses équations différentielles ou équations aux dérivées partielles.

Démonstration. Comme $f' \in L^1(\mathbb{R})$, on peut écrire d'après le théorème de convergence dominée de Lebesgue que

$$\mathcal{F}\{f'(x)\} = \lim_{a \rightarrow +\infty} \int_{-a}^a f'(x) e^{-2i\pi\xi x} dx.$$

Or une intégration par parties donne

$$\int_{-a}^a f'(x) e^{-2i\pi\xi x} dx = [f(x) e^{-2i\pi\xi x}]_{-a}^a + 2i\pi\xi \int_{-a}^a f(x) e^{-2i\pi\xi x} dx.$$

Supposons dans un premier temps que $f(\pm a)$ ait une limite quand $a \rightarrow +\infty$. Comme f est intégrable alors cette limite est nécessairement nulle. Ainsi on obtient

$$\mathcal{F}\{f'(x)\} = \lim_{a \rightarrow +\infty} \int_{-a}^a f'(x) e^{-2i\pi\xi x} dx = \lim_{a \rightarrow +\infty} 2i\pi\xi \int_{-a}^a f(x) e^{-2i\pi\xi x} dx = 2i\pi\xi \mathcal{F}\{f(x)\}.$$

Il reste à démontrer que la limite de $f(a)$ admet une limite lorsque $a \rightarrow +\infty$ (la limite de $f(-a)$ s'obtient de la même manière). Pour ce faire on écrit

$$f(a) = f(0) + \int_0^a f'(x) dx.$$

Or $f' \in L^1(\mathbb{R})$ donc $\lim_{a \rightarrow +\infty} \int_0^a f'(x) dx$ existe et donc $\lim_{a \rightarrow +\infty} f(a)$ existe également. En argumentant par récurrence on déduit le deuxième point de la proposition, ce qui conclut la preuve du résultat. $\square$

Proposition 63. Soit $f \in L^1(\mathbb{R})$, si f admet des dérivées jusqu'à l'ordre $n \geq 1$ qui sont dans $L^1(\mathbb{R})$, alors il existe une constante $C > 0$ telle que

$$|\widehat{f}(\xi)| \leq \frac{C}{|\xi|^n}, \quad \forall \xi \in \mathbb{R}^*.$$

Ce résultat est similaire à ce qui a été déjà observé pour les fonctions périodiques. En effet nous avons déjà vu que plus une fonction périodique est régulière plus ses coefficients de Fourier décroissent rapidement vers 0. Ici on a un comportement similaire dans le cas d'un signal non périodique. En d'autres termes, plus f est régulière et plus l'information sur le spectre de f est « concentrée » au voisinage de zéro. Ce résultat est important en pratique afin de calculer la TFD d'un signal non périodique (via l'algorithme de FFT). Nous reviendrons sur ce point plus tard dans le cours.

Démonstration. On a d'après la Proposition 61 la formule suivante

$$\mathcal{F}\{f^{(n)}(x)\} = (2i\pi\xi)^n \widehat{f}(\xi),$$

d'où

$$|\widehat{f}(\xi)| \leq \frac{1}{2\pi|\xi|^n} \int_{\mathbb{R}} |f^{(n)}(x)| dx = \frac{\|f^{(n)}\|_{L^1(\mathbb{R})}}{2\pi|\xi|^n} = \frac{C}{|\xi|^n}, \quad \forall \xi \in \mathbb{R}^*,$$

avec $C = \|f\|_{L^1(\mathbb{R})}/(2\pi)$. Ce qui conclut la preuve du résultat. $\square$

Proposition 64. Soit $f \in L^1(\mathbb{R})$, si $x \mapsto xf(x) \in L^1(\mathbb{R})$, alors $\widehat{f}$ est dérivable avec

$$\frac{d\widehat{f}}{d\xi}(\xi) = \mathcal{F}\{-2i\pi x f(x)\}.$$

De plus si $x \mapsto x^k f(x) \in L^1(\mathbb{R})$ pour tout $1 \leq k \leq n$ alors la fonction $\widehat{f}$ est n fois dérivable avec

$$\frac{d^n \widehat{f}}{d\xi^n}(\xi) = \mathcal{F}\{(-2i\pi x)^n f(x)\}.$$

Démonstration. Nous savons d'après l'Exemple 7 du Chapitre 4 que $\widehat{f}$ est dérivable sous les hypothèses de la Proposition 64 avec pour tout $\xi \in \mathbb{R}$ la formule

$$\frac{d\widehat{f}}{d\xi}(\xi) = -2i\pi \int_{\mathbb{R}} xf(x)e^{-2i\pi\xi x} dx = \mathcal{F}\{-2i\pi x f(x)\}.$$

Pour démontrer le deuxième point du résultat il suffit d'appliquer par récurrence les arguments de l'Exemple 7 du chapitre 4. Ce qui conclut la preuve de la proposition. $\square$

Exemple 39. Soit g la fonction définie par $g(x) = \frac{x^k}{k!} e^{-ax} H(x)$ pour tout $x \in \mathbb{R}$ et où $k \in \mathbb{N}^*$ et $a \in \mathbb{C}$ avec $\operatorname{Re}(a) > 0$. En utilisant la fonction f de l'Exemple 38, on peut réécrire la fonction g sous la forme

$$g(x) = \frac{1}{(-2i\pi)^k} \frac{1}{k!} (-2i\pi x)^k f(x), \quad \forall x \in \mathbb{R}.$$

Comme la fonction $x \mapsto x^k f(x) \in L^1(\mathbb{R})$ alors d'après la Proposition 64 on obtient en appliquant la transformée de Fourier à cette dernière égalité

$$\widehat{g}(\xi) = \mathcal{F}\{g(x)\} = \frac{1}{(-2i\pi)^k} \frac{1}{k!} \mathcal{F}\{(-2i\pi x)^k f(x)\} = \frac{1}{(-2i\pi)^k} \frac{1}{k!} \frac{d^k \widehat{f}}{d\xi^k}(\xi)$$

Par ailleurs comme

$$\frac{d^k \widehat{f}}{d\xi^k}(\xi) = k! \frac{(-2i\pi)^k}{(a + 2i\pi\xi)^{k+1}},$$

on en déduit que

$$\widehat{g}(\xi) = \frac{1}{(-2i\pi)^k} \frac{1}{k!} k! \frac{(-2i\pi)^k}{(a + 2i\pi\xi)^{k+1}} = \frac{1}{(a + 2i\pi\xi)^{k+1}}.$$

En résumé, pour tout $a \in \mathbb{C}$ avec $\operatorname{Re}(a) > 0$, on a

$$\frac{x^k}{k!} e^{-ax} H(x) \xrightarrow{\mathcal{F}} \frac{1}{(a + 2i\pi\xi)^{k+1}} \quad \forall k \in \mathbb{N}^*.$$

5.1.3 Transformée de Fourier inverse

Comme dans le cas des séries de Fourier, on veut savoir sous quelles conditions il est possible de reconstruire un signal non périodique $f \in L^1(\mathbb{R})$ via sa transformée de Fourier. Autrement dit, comment passer du domaine fréquentiel au domaine temporel.

Définition 25. Soit $f \in L^1(\mathbb{R})$, on appelle transformée de Fourier conjuguée de f la fonction notée $\overline{\mathcal{F}}f : \mathbb{R} \rightarrow \mathbb{C}$ définie par

$$\overline{\mathcal{F}}\{f(x)\} = \int_{\mathbb{R}} f(x) e^{2i\pi\xi x} dx.$$

On a le résultat (admis) suivant :

Théorème 65. Si f et $\widehat{f} \in L^1(\mathbb{R})$ alors

$$f(x) = \int_{\mathbb{R}} \widehat{f}(\xi) e^{2i\pi\xi x} d\xi = \overline{\mathcal{F}}\{\widehat{f}(\xi)\} = \overline{\mathcal{F}}\{\mathcal{F}\{f(x)\}\}, \quad (5.4)$$

en tout point x où f est continue.

Remarque 66. On tient à souligner deux points importants.

- On peut grossièrement interpréter la formule (5.4) sous la forme $f = \overline{\mathcal{F}}\hat{f}$ (si par exemple f est continue sur $\mathbb{R}$) ou encore

$$f = \overline{\mathcal{F}}\mathcal{F}f.$$

Autrement dit, l'opération inverse de $\mathcal{F}$ est $\overline{\mathcal{F}}$, i.e., $\mathcal{F}^{-1} = \overline{\mathcal{F}}$ (on peut faire le parallèle entre la TFD et la TFDI).

- Nous avons vu dans la section précédente que si $f \in L^1(\mathbb{R})$ alors $\hat{f}$ n'est pas nécessairement une fonction de $L^1(\mathbb{R})$. Or afin que la formule (5.4) ait un sens il est important de supposer que $\hat{f} \in L^1(\mathbb{R})$.

En pratique on va utiliser le résultat suivant afin de calculer la transformée de Fourier de certaines fonctions via la formule d'inversion :

Proposition 67. Soit $f \in L^1(\mathbb{R})$ une fonction continue, telle que $\hat{f} \in L^1(\mathbb{R})$, alors pour tout $x \in \mathbb{R}$ on a

$$\mathcal{F}\{\mathcal{F}\{f(x)\}\} = f(-x).$$

Démonstration. On pose $g = \hat{f}$ (dom freq) et on a ($\hat{g}$ dom temp)

$$\hat{g}(x) = \int_{\mathbb{R}} g(\xi) e^{-2i\pi\xi x} d\xi,$$

autrement dit

$$\hat{g}(-x) = \int_{\mathbb{R}} g(\xi) e^{2i\pi\xi x} d\xi = \int_{\mathbb{R}} \hat{f}(\xi) e^{2i\pi\xi x} d\xi \underbrace{=}_{(5.4)} f(x).$$

Ainsi on obtient l'égalité

$$f(-x) = \hat{g}(x) = \mathcal{F}\{g(\xi)\} = \mathcal{F}\{\mathcal{F}\{f(x)\}\}.$$

Ce qui conclut la preuve du résultat. □

Exemple 40. On considère la fonction g de l'Exemple 39, à savoir, $g(x) = \frac{x^k}{k!} e^{-ax} H(x)$ pour tout $x \in \mathbb{R}$ avec $k \geq 1$ et $a \in \mathbb{C}$ où $\operatorname{Re}(a) > 0$. La fonction g est continue et L^1 sur $\mathbb{R}$. De plus on a obtenu, pour tout $\xi \in \mathbb{R}$,

$$\hat{g}(\xi) = \frac{1}{(a + 2i\pi\xi)^{k+1}} \in L^1(\mathbb{R}) \quad \text{car } k \geq 1.$$

Posons $h(\xi) = \hat{g}(\xi)$. D'après la Proposition 67 on en déduit que

$$\hat{h}(x) = \mathcal{F}\{\mathcal{F}\{g(x)\}\} = g(-x) = \frac{(-x)^k}{k!} e^{ax} H(-x).$$

En résumé on a pour tout $a \in \mathbb{C}$ avec $\operatorname{Re}(a) > 0$

$$\frac{1}{(a + 2i\pi\xi)^{k+1}} \xrightarrow{\mathcal{F}} \frac{(-\xi)^k}{k!} e^{a\xi} H(-\xi) \quad k \in \mathbb{N}^*.$$

5.2 Transformée de Fourier dans L^2

L'extension de la définition de la transformée de Fourier à l'espace $L^2(\mathbb{R})$ est une construction délicate qui nécessite l'introduction de l'espace de Schwartz. Afin d'éviter cette construction élaborée on va admettre l'existence d'une application linéaire, notée $\mathcal{F}$ et appelée opérateur de Fourier, de $L^2(\mathbb{R})$ dans $L^2(\mathbb{R})$. On admet également que l'opérateur de Fourier est une isométrie de $L^2(\mathbb{R})$ d'inverse $\overline{\mathcal{F}}$, i.e.,

1. $\forall f \in L^2(\mathbb{R})$ on a $\mathcal{F}\overline{\mathcal{F}}f = \overline{\mathcal{F}}\mathcal{F}f = f$ p.p.,
2. $\forall f, g \in L^2(\mathbb{R})$ on a $\int_{\mathbb{R}} f(x)\overline{g}(x) dx = \int_{\mathbb{R}} \mathcal{F}f(\xi)\overline{\mathcal{F}g}(\xi) d\xi$,
3. $\forall f \in L^2(\mathbb{R})$ on a $\|f\|_{L^2(\mathbb{R})} = \|\mathcal{F}f\|_{L^2(\mathbb{R})}$.

Remarque 68. Le point (1) exprime l'inversibilité de $\mathcal{F}$ et les points (2) et (3) traduisent le fait que $\mathcal{F}$ est une isométrie de $L^2(\mathbb{R})$ (c'est en fait l'égalité de Plancherel-Parseval). Enfin pour distinguer la transformée de Fourier dans $L^1(\mathbb{R})$ et $L^2(\mathbb{R})$, on réservera la notation $\hat{f}$ si $f \in L^1(\mathbb{R})$.

Le résultat suivant renforce l'idée que la transformation de Fourier dans $L^2(\mathbb{R})$ est un objet relativement délicat à manipuler.

Proposition 69. On a les résultats suivants :

- La transformation de Fourier définie sur $L^1(\mathbb{R})$ et celle définie sur $L^2(\mathbb{R})$ coïncident sur $L^1(\mathbb{R}) \cap L^2(\mathbb{R})$.
- Soient $f \in L^2(\mathbb{R})$ et la suite de fonctions

$$g_n(\xi) = \int_{-n}^n f(x) e^{-2i\pi\xi x} dx.$$

Alors $\mathcal{F}f$ est la limite dans $L^2(\mathbb{R})$ de la suite de fonctions $(g_n)_{n \geq 1}$, autrement dit

$$\lim_{n \rightarrow +\infty} \int_{\mathbb{R}} |\mathcal{F}f(\xi) - g_n(\xi)|^2 d\xi = 0.$$

En résumé la transformation de Fourier dans $L^2(\mathbb{R})$ est peu maniable. Dans la pratique on utilisera le résultat suivant :

Proposition 70. On a les relations suivantes

- Si $f \in L^2(\mathbb{R})$ alors $\mathcal{F}\{\mathcal{F}\{f(x)\}\} = f(-x)$ pour presque tout $x \in \mathbb{R}$.
- Si $f \in L^1(\mathbb{R}) \cap L^2(\mathbb{R})$ alors $\mathcal{F}\{\mathcal{F}\{f(x)\}\} = f(-x)$ pour presque tout $x \in \mathbb{R}$.

Exemple 41. Reprenons la fonction f de l'Exemple 38, avec $f(x) = e^{-ax} H(x)$ pour tout $x \in \mathbb{R}$ et $a \in \mathbb{C}$ avec $\operatorname{Re}(a) > 0$. On a obtenu

$$\hat{f}(\xi) = \frac{1}{a + 2i\pi\xi}, \quad \forall \xi \in \mathbb{R}.$$

Ainsi $\hat{f} \notin L^1(\mathbb{R})$ mais par contre $\hat{f} \in L^2(\mathbb{R})$ et donc $\mathcal{F}\hat{f} \in L^2(\mathbb{R})$. Comme $f \in L^1(\mathbb{R}) \cap L^2(\mathbb{R})$ alors d'après la proposition précédente on en déduit que

$$\mathcal{F}\{\mathcal{F}\{f(x)\}\} = f(-x) = e^{ax} H(-x),$$

pour presque tout $x \in \mathbb{R}$ (car ici f est continue sur $\mathbb{R}$ sauf en $x = 0$). En résumé pour tout $a \in \mathbb{C}$ avec $\operatorname{Re}(a) > 0$ on a

$$\frac{1}{a + 2i\pi x} \xrightarrow{\mathcal{F}} e^{a\xi} H(-\xi).$$

5.3 Transformée de Fourier et convolution

Commençons par définir la notion de convolution entre fonctions.

Définition 26. On définit le produit de convolution (si il existe) entre deux fonctions f et g via la formule

$$(f * g)(x) = \int_{\mathbb{R}} f(t) g(x - t) dt = \int_{\mathbb{R}} f(x - t) g(t) dt, \quad x \in \mathbb{R}. \quad (5.5)$$

Remarque 71. Plusieurs remarques s'imposent :

- Il y a un parallèle à faire entre la formule (5.5) et la formule donnant la convolution périodique discrète (d'ordre N) entre deux vecteurs périodiques.
- Si le produit de convolution entre f et g existe alors ce produit est commutatif, i.e., $f * g = g * f$.
- Enfin si la fonction $f * g$ est bien définie on peut la voir comme une intégrale à paramètre (par exemple dans la formule (5.5) le paramètre est x).

Avant de voir sous quelles conditions $f * g$ existe, donnons un exemple.

Exemple 42. Soient $f = g = \mathbf{1}_{[0,1]}$. Alors on a

$$(f * g)(x) = \int_{\mathbb{R}} f(x - t) g(t) dt = \int_{\mathbb{R}} \mathbf{1}_{[0,1]}(x - t) \mathbf{1}_{[0,1]}(t) dt = \int_0^1 \mathbf{1}_{[0,1]}(x - t) dt.$$

Or pour $t \in \mathbb{R}$ il vient

$$\begin{aligned} \mathbf{1}_{[0,1]}(x - t) \neq 0 &\Leftrightarrow 0 \leq x - t \leq 1 \\ &\Leftrightarrow -x \leq -t \leq 1 - x \\ &\Leftrightarrow x \geq t \geq x - 1 \end{aligned}$$

Ainsi pour que $\int_0^1 \mathbf{1}_{[0,1]}(x - t) dt \neq 0$ il faut à la fois que $t \in [0, 1]$ et que $t \in [x - 1, x]$, i.e., pour que $\int_0^1 \mathbf{1}_{[0,1]}(x - t) dt \neq 0$ il faut que

$$t \in [0, 1] \cap [x - 1, x] (\neq \emptyset).$$

Ainsi en distinguant suivant le signe de x on a

$$(f * g)(x) = \begin{cases} 0 & \text{si } x \leq 0, \\ x & \text{si } 0 \leq x \leq 1, \\ 2 - x & \text{si } 1 \leq x \leq 2, \\ 0 & \text{si } x \geq 2. \end{cases}$$

en particulier ici, il faut retenir que la convolution de deux fonctions discontinues donne une fonction continue. Le produit de convolution a un effet régularisant sur les fonctions.

Proposition 72. Soient f et $g \in L^1(\mathbb{R})$. Alors

- $f * g$ existe presque partout sur $\mathbb{R}$ et $f * g \in L^1(\mathbb{R})$.
- $\|f * g\|_{L^1(\mathbb{R})} \leq \|f\|_{L^1(\mathbb{R})} \|g\|_{L^1(\mathbb{R})}$.

Proposition 73. Soient f et $g \in L^2(\mathbb{R})$ alors

- la fonction $f * g$ est bien définie sur $\mathbb{R}$ et est une fonction continue et bornée sur $\mathbb{R}$,
- $\sup_{x \in \mathbb{R}} |(f * g)(x)| \leq \|f\|_{L^2(\mathbb{R})} \|g\|_{L^2(\mathbb{R})}$.

En fait la proposition précédente peut se généraliser. On dit que p et q , deux réels positifs (éventuellement $+\infty$), sont conjugués si $\frac{1}{p} + \frac{1}{q} = 1$. On a alors le résultat suivant :

Proposition 74. Soient $f \in L^p(\mathbb{R})$ et $g \in L^q(\mathbb{R})$ avec p et q conjugués, alors

- la fonction $f * g$ est bien définie sur $\mathbb{R}$ et est une fonction continue et bornée sur $\mathbb{R}$,
- $\sup_{x \in \mathbb{R}} |(f * g)(x)| \leq \|f\|_{L^p(\mathbb{R})} \|g\|_{L^q(\mathbb{R})}$.

Proposition 75. Soient $f \in L^1(\mathbb{R})$ et $g \in L^2(\mathbb{R})$ alors

- la fonction $f * g$ existe presque partout sur $\mathbb{R}$
- $f * g \in L^2(\mathbb{R})$ et on a

$$\|f * g\|_{L^2(\mathbb{R})} \leq \|f\|_{L^1(\mathbb{R})} \|g\|_{L^2(\mathbb{R})}.$$

Remarque 76. Dans la proposition précédente 1 et 2 ne sont pas conjugués ! La Proposition 75 n'est pas un cas particulier de la Proposition 74.

Pour résumé on a le tableau suivant :

Comme déjà remarqué dans l'Exemple 42 le produit de convolution a un effet régularisant. En effet on a le résultat suivant :

Proposition 77. Soient $f \in L^1(\mathbb{R})$ et $g \in \mathcal{C}^p(\mathbb{R})$ (g est p -fois dérivable et $g^{(p)}$ est continue). On suppose que $g^{(k)}$ est bornée pour tout $k = 0, \dots, p$ alors :

- la fonction $f * g$ est une fonction $\mathcal{C}^p(\mathbb{R})$,
- $(f * g)^{(k)} = f * g^{(k)}$ pour tout $k = 1, \dots, p$.

Démonstration. Ici on va se restreindre au cas $p = 1$. On commence par remarquer que comme g et g' sont bornées par hypothèse, i.e., g et $g' \in L^\infty(\mathbb{R})$ alors $f * g$ et $f * g'$ existent. De plus d'après la Proposition 74, les fonctions $f * g$ et $f * g'$ sont continues. Pour achever la démonstration du résultat il reste à voir si il est possible d'appliquer le théorème

$L^1(\mathbb{R}) * L^1(\mathbb{R})$	$L^1(\mathbb{R})$
$L^2(\mathbb{R}) * L^2(\mathbb{R})$	bornée et continue
$L^1(\mathbb{R}) * L^\infty(\mathbb{R})$	bornée et continue
$L^1(\mathbb{R}) * L^2(\mathbb{R})$	$L^2(\mathbb{R})$

TABLE 5.1 – Existence de $f * g$.

de dérivation des intégrales à paramètre. Dans un premier temps on écrit que pour tout $x \in \mathbb{R}$

$$(f * g)(x) = \int_{\mathbb{R}} f(t) g(x - t) dt.$$

Or la fonction $x \mapsto f(t)g(x - t)$ est dérivable sur $\mathbb{R}$ pour presque tout $t \in \mathbb{R}$ (car g est dérivable sur $\mathbb{R}$). De plus pour tout $x \in \mathbb{R}$ on a

$$\frac{\partial}{\partial x}(f(t) g(x - t)) = f(t) g'(x - t) \quad \text{p.p. } t \in \mathbb{R}.$$

Comme g' est bornée sur $\mathbb{R}$ notons $M > 0$ vérifiant $\sup_{y \in \mathbb{R}} |g'(y)| \leq M$. Alors pour tout $x \in \mathbb{R}$ on a

$$\left| \frac{\partial}{\partial x}(f(t) g(x - t)) \right| = |f(t)| |g'(x - t)| \leq M |f(t)| \quad \text{p.p. } t \in \mathbb{R}.$$

Comme $f \in L^1(\mathbb{R})$ on en déduit du théorème de dérivation des intégrales à paramètre que $f * g$ est dérivable sur $\mathbb{R}$ et on a pour tout $x \in \mathbb{R}$

$$(f * g)'(x) = \int_{\mathbb{R}} f(t) g'(x - t) dt = (f * g')(x).$$

Comme la fonction $f * g'$ est continue sur $\mathbb{R}$ alors $(f * g)'$ est continue sur $\mathbb{R}$ et donc $f * g$ est $\mathcal{C}^1(\mathbb{R})$. Ce qui conclut la preuve du résultat. $\square$

Pour conclure cette section nous allons faire de nouveau un parallèle avec le cas discret. En pratique le produit de convolution n'est pas toujours simple à manipuler. Cependant en utilisant la transformée de Fourier on peut parfois « simplifier » le calcul.

Proposition 78. Soient f et $g \in L^1(\mathbb{R})$ alors

1. $f * g \in L^1(\mathbb{R})$ et on a

$$\mathcal{F}\{f * g(x)\} = \mathcal{F}\{f(x)\} \mathcal{F}\{g(x)\}. \quad (5.6)$$

2. Si $\mathcal{F}f$ et $\mathcal{F}g \in L^1(\mathbb{R})$ alors

$$\hat{f} * \hat{g}(\xi) = \mathcal{F}\{f(x)g(x)\}.$$

« Démonstration ». On donne juste l'idée de la démonstration du point (1), le calcul suivant peut être rendu rigoureux. On a par définition

$$\begin{aligned} \mathcal{F}\{f * g(x)\} &= \int_{\mathbb{R}} (f * g)(x) e^{-2i\pi\xi x} dx \\ &= \int_{\mathbb{R}} e^{-2i\pi\xi x} \left(\int_{\mathbb{R}} f(x-t) g(t) dt \right) dx \\ &= \int_{\mathbb{R}} g(t) \left(\int_{\mathbb{R}} f(x-t) e^{-2i\pi\xi x} dx \right) dt \\ &= \int_{\mathbb{R}} g(t) \mathcal{F}\{f(x-t)\} dt \quad \underbrace{\quad}_{\text{Prop. 57}} \int_{\mathbb{R}} g(t) e^{-2i\pi\xi t} \hat{f}(\xi) dt = \hat{f}(\xi) \hat{g}(\xi). \end{aligned}$$

Ce calcul permet de retrouver la formule (5.6). $\square$

Il est légitime de se demander si une formule du type (5.6) est valable si f et g sont des fonctions de $L^2(\mathbb{R})$. Cependant dans ce cas on sait d'après la Table 5.1 que $f * g$ est continue et bornée sur $\mathbb{R}$ mais rien n'indique que $f * g \in L^1(\mathbb{R})$. Ainsi a priori la transformée de Fourier de $f * g$ n'est pas bien définie. On a par contre le résultat suivant :

Proposition 79. Soient f et $g \in L^2(\mathbb{R})$ alors

$$(f * g)(x) = \overline{\mathcal{F}\{f(x)\}} \mathcal{F}\{g(x)\}, \quad \forall x \in \mathbb{R}.$$

5.4 Applications et illustrations

Dans cette section on présente deux applications de la transformée de Fourier et de la notion de convolution.

5.4.1 Résolution de l'équation de la chaleur sur un domaine non borné

On s'intéresse à l'équation de la chaleur sur domaine non borné

$$\partial_t u - D \partial_x^2 u = 0 \quad x \in \mathbb{R}, t > 0, \quad (5.7)$$

$$u(0, x) = u^0(x) \quad x \in \mathbb{R}, \quad (5.8)$$

où u représente la distribution de la chaleur dans le domaine, $D > 0$ la constante de diffusion et $u^0 \in C^2(\mathbb{R})$ la donnée initiale (u^0 représente la température à l'instant $t = 0$).

On cherche à déterminer une fonction u régulière (au moins une fois dérivable en temps et deux fois en espace) solution de (5.7)–(5.8). Pour ce faire on applique la transformée de Fourier **en espace** à l'équation (5.7). Or d'après la Proposition 61 on a

$$\widehat{\partial_x u}(\xi, t) = (-2i\pi\xi) \widehat{u}(\xi, t) \quad \text{et} \quad \widehat{\partial_x^2 u}(\xi, t) = (-4\pi^2\xi^2) \widehat{u}(\xi, t).$$

Ainsi en admettant que l'on puisse inverser intégrale et dérivée en temps on obtient que $\widehat{u}$ est solution de l'équation différentielle suivante :

$$\partial_t \widehat{u}(\xi, t) = (-4\pi^2 D \xi^2) \widehat{u}(\xi, t), \quad \xi \in \mathbb{R}, t > 0 \quad (5.9)$$

$$\widehat{u}(\xi, 0) = \widehat{u^0}(\xi), \quad \xi \in \mathbb{R}. \quad (5.10)$$

Une simple résolution de cette EDO implique que

$$\widehat{u}(\xi, t) = C e^{-4\pi^2 D \xi^2 t},$$

où C désigne une constante à déterminer. Pour cela on utilise la condition initiale (5.8) et on remarque que

$$\widehat{u}(\xi, 0) = \int_{\mathbb{R}} u(x, 0) e^{-2i\pi\xi x} dx = \int_{\mathbb{R}} u^0(x) e^{-2i\pi\xi x} dx = \widehat{u^0}(\xi).$$

Ainsi on a

$$C = \widehat{u}(\xi, 0) = \widehat{u^0}(\xi),$$

et donc $\widehat{u}$ solution de (5.9)–(5.10) est donnée par

$$\widehat{u}(\xi, t) = \widehat{u^0}(\xi) e^{-4\pi^2 D \xi^2 t}, \quad \xi \in \mathbb{R}, t \geq 0.$$

En utilisant la convolution on peut en fait réécrire cette expression. En effet, comme vu en TD la fonction g définie par

$$g(x, t) = \frac{1}{2\sqrt{\pi Dt}} e^{-\frac{x^2}{4Dt}}, \quad x \in \mathbb{R}, t \geq 0,$$

vérifie

$$\widehat{g}(\xi, t) = e^{-4\pi^2 D \xi^2 t}, \quad \xi \in \mathbb{R}, t \geq 0.$$

On en déduit que

$$\widehat{u}(\xi, t) = \widehat{u^0}(\xi) \widehat{g}(\xi, t), \quad \xi \in \mathbb{R}, t \geq 0.$$

Or d'après la Proposition 78 on a $\widehat{u^0}(\xi) \widehat{g}(\xi, t) = \widehat{u^0 * g}(\xi, t)$ et donc $\widehat{u}$ solution de (5.9)–(5.10) est donnée par

$$\widehat{u}(\xi, t) = \widehat{u^0 * g}(\xi, t), \quad \xi \in \mathbb{R}, t \geq 0.$$

En appliquant à présent la transformée de Fourier inverse on conclut que u solution de (5.7)–(5.8) est donnée par la formule

$$u(x, t) = (u^0 * g)(x, t) = \frac{1}{2\sqrt{D\pi t}} \int_{\mathbb{R}} u^0(y) e^{-\frac{(x-y)^2}{4Dt}} dy, \quad x \in \mathbb{R}, t \geq 0. \quad (5.11)$$

En particulier, via l'effet régularisant de la convolution, il est important de remarquer que la solution u de (5.7)–(5.8) est régulière même si u^0 ne l'est pas. Par exemple si on considère

$$u^0(x) = \frac{1}{2} \mathbf{1}_{[-1,1]}(x), \quad x \in \mathbb{R},$$

alors un calcul explicite (à faire en exercice) donne que la solution de (5.7)–(5.8) est donnée par

$$u(x, t) = \frac{1}{4} \left(\operatorname{erf} \left(\frac{1-x}{2\sqrt{Dt}} \right) + \operatorname{erf} \left(\frac{1+x}{2\sqrt{Dt}} \right) \right),$$

où erf est la fonction erreur définie par

$$\operatorname{erf}(x) = \frac{2}{\sqrt{\pi}} \int_0^x \exp(-z^2) dz. \quad (5.12)$$

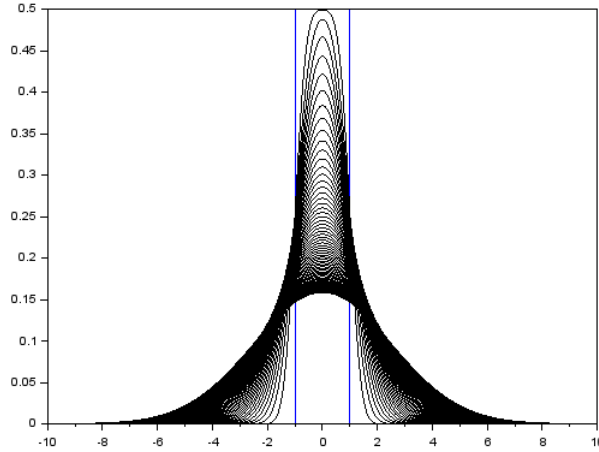


FIGURE 5.1 – Evolution de la solution (5.12) au cours du temps (u^0 en bleu).

5.4.2 Transformée de Fourier à fenêtre glissante (transformée de Gabor)

Pour un signal donné, dénoté f , on a vu au travers de ce cours qu'une manière d'obtenir des informations sur f est de considérer son spectre, i.e., si le signal f est assez régulier

on peut étudier $\hat{f}$. L'étude du spectre de f nous donne des informations dans le domaine fréquentiel mais pas dans le domaine temporel. Plus simplement si f est le résultat d'un enregistrement de quatre secondes pendant lequel une note de musique est jouée durant une seconde, alors le spectre de f permet de déterminer cette note mais ne permet pas de déterminer quand cette note a été jouée durant l'enregistrement. Ainsi en passant par le spectre on obtient une information fréquentielle mais on perd l'information temporelle. À l'inverse si on écoute f on peut déterminer très précisément l'instant où la note est jouée mais on ne peut déterminer la note jouée (sauf si votre oreille musicale est très développée). La transformée de Gabor ou transformée de Fourier à fenêtre glissante est un outil permettant dans une certaine mesure de concilier les deux. En effet cet outil permet de tracer le spectrogramme d'un signal, voici par exemple le spectrogramme d'un morceau de Beethoven.

Pour construire cette transformation de Gabor l'idée principale est de tronquer le signal initial f sur un intervalle fermé et borné de $\mathbb{R}$, i.e., on multiplie f par $\mathbf{1}_{[-A,A]}$ avec $A > 0$. Calculons la transformée de Fourier de ce nouveau signal, noté g , on a

$$\widehat{g}(\xi) = \widehat{\mathbf{1}_{[-A,A]} f}(\xi) = \widehat{\mathbf{1}_{[-A,A]}} * \widehat{f}(\xi) = s_A * \widehat{f}(\xi),$$

avec

$$s_A(\xi) = \widehat{\mathbf{1}_{[-A,A]}}(\xi) = \frac{\sin(2\pi A\xi)}{\pi\xi}.$$

Bien entendu les spectres de f et g sont différents, autrement dit en multipliant f par $\mathbf{1}_{[-A,A]}$ on commet une erreur sur le spectre de f . En observant les graphes de s_A lorsque $A \rightarrow +\infty$, voir Figure 5.2, on voit que formellement

$$s_A \rightarrow \delta_0,$$

où δ_0 désigne la distribution de Dirac en zéro avec

$$\delta_0(x) = \begin{cases} 0 & \text{si } x \neq 0, \\ +\infty & \text{si } x = 0. \end{cases}$$

Le distribution de Dirac n'est pas une fonction mais un objet qui généralise la notion de fonction. Par ailleurs il revient de remarquer que δ_0 est l'élément unité du produit de convolution. Formellement si $A \rightarrow +\infty$ on a $g = \mathbf{1}_{\mathbb{R}} f = f$ et donc

$$\widehat{f}(\xi) = \delta_0 * \widehat{f}(\xi).$$

Il faut donc retenir que plus A est grand est plus $g = \mathbf{1}_{[-A,A]} f$ « s'approche » de f , i.e., plus A est grand est plus $\widehat{g}$ et $\widehat{f}$ sont similaires. Néanmoins plus A est grand et plus les calculs sont volumineux. Ainsi en pratique on ne considère jamais le cas où A est très grand ou pire le cas où $A \rightarrow +\infty$.

Par ailleurs, la fonction s_A s'amortit très lentement et présentes des lobes importants près de l'origine. Ainsi en pratique pour les diminuer on utilise des fonctions plus régulières, appelées fenêtres d'apodisation (ou plus simplement fenêtres), concentrées autour de l'origine, voici plusieurs exemples :

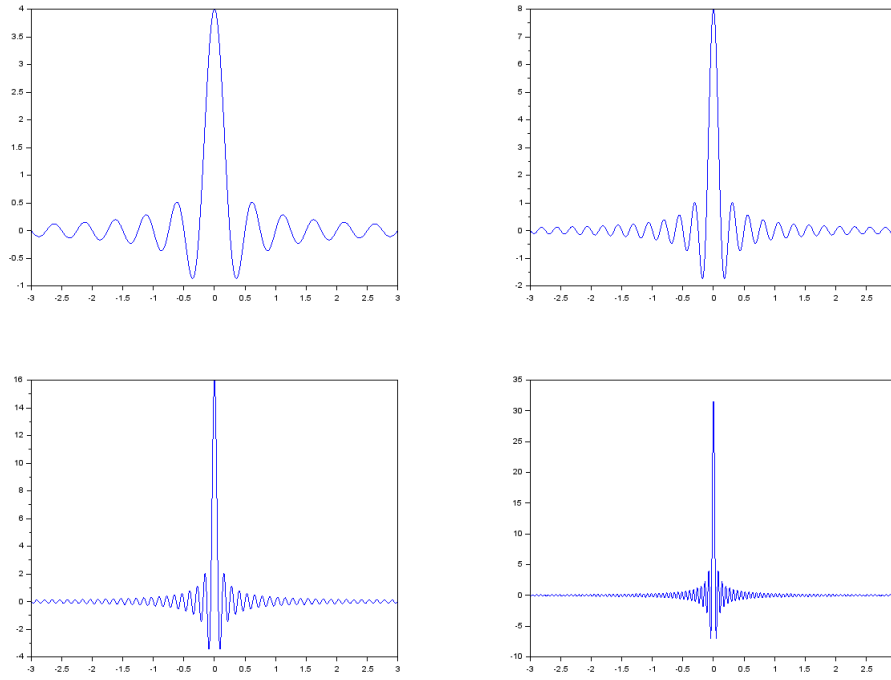


FIGURE 5.2 – Graphes de s_A pour $A = 2$ (en haut à gauche), $A = 4$ (en haut à droite), $A = 8$ (en bas à gauche) et $A = 16$ (en bas à droite).

- fenêtre de Hamming,
- fenêtre de Hanning,
- fenêtre de Blackman,
- fenêtre de Gauss.

On renvoie à la page fenêtrage de Wikipédia pour une définition formelle de ces fenêtres. Comme dans le cas des filtres il n'existe pas de fenêtre optimale, le choix dépend de l'application visée. Je conseille de suivre l'UV SY06 pour plus d'explications sur ce choix.

Dans la suite la fonction ω désigne une fenêtre d'apodisation. Pour un signal f on définit alors la transformation de Fourier à fenêtre glissante de f par la formule

$$W_f(\xi, b) = \int_{\mathbb{R}} f(t) \omega(t - b) e^{-2i\pi\xi t} dt.$$

Ici dans la formule donnant W_f on vient faire « glisser » au cours du temps la fenêtre ω sur le signal f . En fait le terme $f(t)e^{-2i\pi\xi t}$ renvoie bien à la transformée de Fourier « classique ». Autrement dit, $W_f(\xi, b)$ donne une indication sur la spectre de f autour de l'abscisse $t = b$ pour la fréquence ξ . En notant

$$F(t, \xi) = f(t) e^{-2i\pi\xi t},$$

on peut réécrire $W_f(\xi, b)$ comme un produit de convolution en temps de F et ω , i.e.,

$$W_f(\xi, b) = (F *_{\text{temps}} w)(\xi, b)$$

Comme dans le cas de la transformée de Fourier, il est légitime de se demander si la connaissance des coefficients $W_f(\xi, b)$ pour toutes valeurs de ξ et b permet de déterminer entièrement f . On a la résultat suivant :

Théorème 80. *Soit ω une fenêtre avec $\omega \in L^1(\mathbb{R}) \cap L^2(\mathbb{R})$, $\|\omega\|_{L^2(\mathbb{R})} = 1$ et $|\hat{\omega}|$ est une fonction paire. Alors les coefficients*

$$W_f(\xi, b) = \int_{\mathbb{R}} f(t) \omega(t - b) e^{-2i\pi\xi t} dt,$$

vérifient

1. *la conservation de l'énergie*

$$\int_{\mathbb{R}} \int_{\mathbb{R}} W_f(\xi, b) d\xi db = \int_{\mathbb{R}} |f(t)|^2 dt$$

2. *la formule de reconstruction*

$$f(x) = \int_{\mathbb{R}} \int_{\mathbb{R}} W_f(\xi, b) \omega(t - b) e^{2i\pi\xi t} d\xi db \quad p.p. \ x \in \mathbb{R}.$$

Plusieurs remarques s'imposent :

- Le spectrogramme d'un signal revient à tracer les coefficients $|W_f(\xi, b)|^2$ avec en abscisse le temps et en ordonnée les fréquences.
- On peut voir un autre intérêt pratique à la transformée de Fourier à fenêtre glissante. En effet en théorie afin de pouvoir calculer $\hat{f}$ il faut connaître $f(t)$ pour tout temps $t \in \mathbb{R}$. Ceci est impossible pour l'analyse en temps réel d'un signal traité au fur et à mesure de son enregistrement. En particulier dans ce cas on ne connaît pas le comportement du signal dans le futur.

Enfin pour conclure ce chapitre et à tire d'ouverture parlons de deux problèmes liés à la transformée de Fourier à fenêtre glissante :

- Le spectrogramme d'un signal souffre toujours du principe d'incertitude. Ce principe énonce l'impossibilité d'être à la fois très précis en temps et en fréquence. Par exemple pour jouer une note de musique très pure (précision fréquentielle), il faut la jouer pendant une certaine durée (donc faible précision temporelle), sans quoi elle sonne comme un choc (grande précision temporelle, mais faible précision fréquentielle). Ainsi le spectrogramme « cherche » le meilleur compromis entre la précision temporelle et fréquentielle, il est cependant toujours entaché de petites erreurs (petites imprécisions temporelles et fréquentielles). Le principe d'incertitude peut être rapproché du principe d'incertitude d'Heisenberg intervenant en mécanique quantique. En effet en mécanique quantique on ne peut connaître simultanément de manière très précise la position et la vitesse d'une particule (contrairement au cas de la mécanique classique).

- La principale contrainte de la transformée de Gabor vient du fait que la fenêtre d'apodisation ω est fixe. En particulier cette fenêtre fixe pose des problèmes dans le cas où le signal présente de fortes variations fréquentielles. Pour contourner ce problème la théorie des ondelettes a été introduite en 1983 par Jean Morlet et Alex Grossman. L'idée principale de la théorie des ondelettes est de remplacer la fenêtre fixe ω par une fenêtre variant par translation, dilatation et contraction. Ainsi même si le signal initial présente de fortes variations fréquentielles on peut « zoomer » ou « dézoomer » facilement.

5.5 Compléments

5.5.1 Transformées de Fourier usuelles

On note H la fonction de Heaviside. Voici un résumé des transformées de Fourier usuelles étudiées en cours ou en TD. On considère dans un premier temps $a \in \mathbb{C}$ avec $\operatorname{Re}(a) > 0$, on a :

$$\begin{aligned}
 e^{-ax} H(x) &\xrightarrow{\mathcal{F}} \frac{1}{a + 2i\pi\xi}, \\
 \frac{1}{a + 2i\pi x} &\xrightarrow{\mathcal{F}} e^{a\xi} H(-\xi), \\
 \frac{x^k}{k!} e^{-ax} H(x) &\xrightarrow{\mathcal{F}} \frac{1}{(a + 2i\pi\xi)^{k+1}} \quad \forall k \in \mathbb{N}^*, \\
 \frac{1}{(a + 2i\pi x)^{k+1}} &\xrightarrow{\mathcal{F}} \frac{(-\xi)^k}{k!} e^{a\xi} H(-\xi) \quad k \in \mathbb{N}^*, \\
 e^{-a|x|} &\xrightarrow{\mathcal{F}} \frac{2a}{a^2 + 4\pi^2\xi^2}, \\
 \frac{1}{a^2 + x^2} &\xrightarrow{\mathcal{F}} \frac{\pi}{a} e^{-2\pi a|\xi|}.
 \end{aligned}$$

Si maintenant $a \in \mathbb{R}$ avec $a > 0$ on a

$$\begin{aligned}
 e^{-ax^2} &\xrightarrow{\mathcal{F}} \sqrt{\frac{\pi}{a}} e^{-\frac{\pi^2}{a}\xi^2}, \\
 \sqrt{\frac{\pi}{a}} e^{-\frac{\pi^2}{a}x^2} &\xrightarrow{\mathcal{F}} e^{-a\xi^2}, \\
 \mathbf{1}_{[-a,a]}(x) &\xrightarrow{\mathcal{F}} \frac{\sin(2a\pi\xi)}{\pi\xi}.
 \end{aligned}$$

Chapitre 6

Introduction à l'échantillonnage

Résumé. Dans ce chapitre nous allons présenter de manière informelle la théorie mathématique permettant d'échantillonner un signal. Cette présentation ne sera pas complètement rigoureuse et certains calculs pourront, pour un public non averti, sembler « choquants ». Il faut garder à l'esprit que tous les calculs de ce chapitre peuvent être complètement justifiés par la théorie des distributions. Cependant, cette théorie est trop complexe pour une présentation claire en peu de temps. Toutes ces raisons justifient la présentation qui suit.

6.1 Impulsion de Dirac et peigne de Dirac

On introduit dans cette section deux objets mathématiques dont l'utilisation est courante dans de nombreuses sciences.

6.1.1 Impulsion de Dirac

Soit $a > 0$, on définit la fonction « porte » par

$$P_a(t) = \begin{cases} 1 & \text{si } |t| \leq a/2, \\ 0 & \text{sinon.} \end{cases}$$

On introduit ensuite l'impulsion de Dirac, notée δ_0 , via la formule (formelle)

$$\delta_0(t) = \lim_{a \rightarrow 0} \frac{1}{a} P_a(t).$$

Par le biais de cette formule on a la définition suivante

$$\delta_0(t) = \begin{cases} +\infty & \text{si } t = 0, \\ 0 & \text{sinon.} \end{cases}$$

Cette définition est déjà étrange au premier abord mais la suite est peut-être pire. En effet, on remarque (toujours formellement) que

$$\int_{\mathbb{R}} \delta_0(t) dt = \int_{\mathbb{R}} \lim_{a \rightarrow 0} \frac{1}{a} P_a(t) dt = \lim_{a \rightarrow 0} \int_{\mathbb{R}} \frac{1}{a} P_a(t) dt = \lim_{a \rightarrow 0} \frac{a}{a} = 1.$$

Cette formule implique que δ_0 ne peut être une fonction au sens usuel du terme. L'impulsion de Dirac est un objet mathématique appelé distribution. La notion de distribution généralise la notion de fonction classique. Si cette théorie est très riche et intéressante, elle est par contre particulièrement difficile à maîtriser. De plus, dans la pratique (pour un ingénieur par exemple) une compréhension de la notion d'impulsion de Dirac est bien suffisante.

Étudions maintenant comment utiliser l'impulsion de Dirac en pratique. On remarque dans un premier temps que si φ est une fonction alors

$$\int_{\mathbb{R}} \delta_0(t) \varphi(t) dt = \lim_{a \rightarrow 0} \frac{1}{a} \int_{\mathbb{R}} P_a(t) \varphi(t) dt = \lim_{a \rightarrow 0} \frac{1}{a} \int_{-a/2}^{a/2} \varphi(t) dt = \varphi(0).$$

De la même manière on a pour un $t_0 \in \mathbb{R}$ fixé

$$\int_{\mathbb{R}} \delta_0(t - t_0) \varphi(t) dt = \varphi(t_0). \quad (6.1)$$

On peut ainsi écrire formellement que

$$\boxed{\delta_0(t - t_0) \varphi(t) = \delta_0(t - t_0) \varphi(t_0)}.$$

Grossièrement cette formule traduit que la multiplication d'une fonction (signal) φ par une impulsion de Dirac $\delta_0(\cdot - t_0)$ permet de « filtrer » les valeurs de φ pour ne garder que sa valeur à l'instant t_0 .

Étudions à présent le spectre de l'impulsion de Dirac, on a pour tout $\xi \in \mathbb{R}$

$$\mathcal{F}\{\delta_0(t)\} = \int_{\mathbb{R}} \delta_0(t) e^{-2i\pi\xi t} dt = \exp(0),$$

et donc

$$\boxed{\mathcal{F}\{\delta_0(t)\} = 1 \quad \forall \xi \in \mathbb{R}}$$

De la même manière pour $t_0 \in \mathbb{R}$ et pour tout $\xi \in \mathbb{R}$ on obtient

$$\mathcal{F}\{\delta_0(t - t_0)\} = \int_{\mathbb{R}} \delta_0(t - t_0) e^{-2i\pi\xi t} dt.$$

Ainsi, d'après la formule (6.1) on en déduit que

$$\boxed{\mathcal{F}\{\delta_0(t - t_0)\} = e^{-2i\pi\xi t_0} \quad \forall \xi \in \mathbb{R}}.$$

6.1.2 Peigne de Dirac et principales propriétés

Soit $T_e > 0$, on définit le peigne de Dirac par la formule

$$\Pi_{T_e}(t) = \sum_{k \in \mathbb{Z}} \delta_0(t - kT_e).$$

Comme précédemment, ici on ne se pose pas la question de savoir quel sens peut avoir la somme « infinie » apparaissant dans la définition de la fonction Π_{T_e} .

En pratique, comme un peigne de Dirac est une somme d'impulsions de Dirac on va utiliser cet objet pour ne conserver que les valeurs d'un signal aux instant kT_e ($k \in \mathbb{Z}$), i.e., soit φ un signal alors

$$\varphi(t) \Pi_{T_e}(t) = \sum_{k \in \mathbb{Z}} \varphi(kT_e).$$

Nous reviendrons dans la suite sur cette formule. Etudions dans un premier temps quelques propriétés de Π_{T_e} .

Commençons par démontrer que Π_{T_e} vérifie la formule suivante :

$$\Pi_{T_e}(t) = \sum_{k \in \mathbb{Z}} \delta_0(t - kT_e) = \frac{1}{T_e} \sum_{k \in \mathbb{Z}} e^{2i\pi k \frac{t}{T_e}}. \quad (6.2)$$

Pour « démontrer » cette formule on utilise (dans le cas des distributions) la formule sommatoire de Poisson. Cette formule affirme que pour une fonction φ (vérifiant certaines hypothèses) et $a > 0$ alors

$$\sum_{k \in \mathbb{Z}} \varphi(t + ka) = \frac{1}{a} \sum_{n \in \mathbb{Z}} \widehat{\varphi}\left(\frac{n}{a}\right) e^{2i\pi n \frac{t}{a}}.$$

Ici si on considère $a = T_e$ et pour φ l'objet δ_0 alors, comme $\widehat{\delta_0}(\xi) = 1$ pour tout $\xi \in \mathbb{R}$, on obtient

$$\Pi_{T_e}(t) = \frac{1}{T_e} \sum_{n \in \mathbb{Z}} e^{-2i\pi n \frac{t}{T_e}},$$

et quitte à poser $k = -n$ on retrouve la formule (6.2).

Etudions à présent le spectre de Π_{T_e} . Pour ce faire via la formule (6.2) on obtient pour tout $\xi \in \mathbb{R}$

$$\widehat{\Pi_{T_e}}(\xi) = \frac{1}{T_e} \sum_{k \in \mathbb{Z}} \mathcal{F} \left\{ e^{2i\pi k \frac{t}{T_e}} \right\}.$$

Afin de calculer (toujours formellement) la transformée de Fourier du membre de droite on va utiliser la formule de modulation du Chapitre 5, i.e., pour toute fonction φ et pour $\xi_0 \in \mathbb{R}$ fixé on a

$$\mathcal{F} \{ e^{2i\pi \xi_0 t} \varphi(t) \} = \int_{\mathbb{R}} \varphi(t) e^{-2i\pi(\xi - \xi_0)t} dt = \widehat{\varphi}(\xi - \xi_0).$$

On applique ensuite cette formule et d'autres formules du Chapitre 5 (dans le cas des distributions) à la fonction $\varphi(t) = 1$ pour tout $t \in \mathbb{R}$ et on obtient :

$$\mathcal{F}\{1\} = \mathcal{F}\{\mathcal{F}\{\delta_0(t)\}\} = \delta_0(-t) = \delta_0(t).$$

Ainsi on en déduit que

$$\mathcal{F}\{e^{2i\pi\xi_0 t}\} = \delta_0(\xi - \xi_0),$$

et donc avec $\xi_0 = k/T_e$

$$\widehat{\Pi_{T_e}}(\xi) = \frac{1}{T_e} \sum_{k \in \mathbb{Z}} \delta_0\left(\xi - \frac{k}{T_e}\right) = \frac{1}{T_e} \Pi_{\frac{1}{T_e}}(\xi) \quad \forall \xi \in \mathbb{R}.$$

Il faut donc retenir que à un facteur $1/T_e$ près, la transformée de Fourier d'un peigne de Dirac de période T_e (dans le domaine temporel) est un peigne de Dirac de période $1/T_e$ (dans le domaine fréquentiel).

6.2 Echantillonnage d'un signal

Dans cette partie on considère un signal (une fonction) $x : \mathbb{R} \rightarrow \mathbb{R}$ supposé à bande limitée

Définition 27. On dit que un signal x est à bande limitée si il existe une fréquence $\xi_0 \in \mathbb{R}$ telle que $\widehat{x}$ soit nulle en dehors de l'intervalle $[-\xi_0, \xi_0]$.

6.2.1 Principal général

Supposons connaître les valeurs d'un signal x en des temps de la forme kT_e où $k \in \mathbb{Z}$ et $T_e > 0$ est appelée période d'échantillonnage, i.e., on suppose connaître la suite $(x(kT_e))_{k \in \mathbb{Z}}$ appelée échantillonnage du signal x . A partir de cette suite nous aimerions reconstruire le signal initial $x(t)$ pour tout $t \in \mathbb{R}$.

Intuitivement, si T_e est trop grand, l'échantillon $(x(kT_e))_{k \in \mathbb{Z}}$ ne contient pas assez d'informations pour reconstruire x . On comprend donc que pour avoir une chance de reconstruire x il faut choisir une période d'échantillonnage $T_e \ll$ assez petite $\gg$. La question qui se pose est donc de comprendre comment choisir T_e convenablement.

Pour ce faire, à partir de la suite $(x(kT_e))_{k \in \mathbb{Z}}$ on construit la « fonction » x_e de la manière suivante

$$x_e(t) = T_e \sum_{k \in \mathbb{Z}} x(kT_e) \delta_0(t - kT_e) = T_e x(t) \Pi_{T_e}(t). \quad (6.3)$$

Grossièrement, échantillonner un signal x à la période T_e revient à le multiplier par un peigne de Dirac de période T_e .

Remarque 81. La constante T_e apparaissant dans la formule (6.3) devant la somme, est une constante de renormalisation utile pour la suite.

6.2.2 Echantillonnage d'un signal à bande limitée

Pour mieux comprendre comment choisir T_e on va étudier le spectre de x_e . On a la suite d'égalités suivante :

$$\begin{aligned}
 \widehat{x}_e(\xi) &= T_e \mathcal{F}\{x(t) \Pi_{T_e}(t)\} = T_e \left(\widehat{x} * \widehat{\Pi_{T_e}} \right) (\xi) \\
 &= T_e \left(\widehat{x} * \frac{1}{T_e} \Pi_{\frac{1}{T_e}} \right) (\xi) \\
 &= \left(\widehat{x} * \Pi_{\frac{1}{T_e}} \right) (\xi) \\
 &= \int_{\mathbb{R}} \widehat{x}(\nu) \Pi_{\frac{1}{T_e}}(\xi - \nu) d\nu \\
 &= \int_{\mathbb{R}} \widehat{x}(\nu) \left(\sum_{k \in \mathbb{Z}} \delta_0 \left(\xi - \nu - \frac{k}{T_e} \right) \right) d\nu \\
 &= \sum_{k \in \mathbb{Z}} \int_{\mathbb{R}} \widehat{x}(\nu) \delta_0 \left(\xi - \nu - \frac{k}{T_e} \right) d\nu,
 \end{aligned}$$

d'où la formule importante

$$\boxed{\widehat{x}_e(\xi) = \sum_{k \in \mathbb{Z}} \widehat{x} \left(\xi - \frac{k}{T_e} \right).}$$

Cette formule est importante car elle permet de comprendre que $\widehat{x}_e$ est $1/T_e$ périodique ! En effet, on a

$$\widehat{x}_e \left(\xi + \frac{1}{T_e} \right) = \sum_{k \in \mathbb{Z}} \widehat{x} \left(\xi - \frac{(k-1)}{T_e} \right) \stackrel{k'=k-1}{=} \sum_{k' \in \mathbb{Z}} \widehat{x} \left(\xi - \frac{k'}{T_e} \right) = \widehat{x}_e(\xi).$$

Ainsi¹, même si $\widehat{x}$ n'est pas périodique (car x est supposé à bande limitée) le spectre de x_e est périodique. En particulier, si l'intervalle $[-\xi_0, \xi_0]$ (où $\widehat{x}$ est non nulle) est contenu dans l'intervalle $[-1/2T_e, 1/2T_e]$ alors les spectres de x et x_e coïncident (au moins sur l'intervalle $[-\xi_0, \xi_0]$). Autrement dit si T_e est assez petit pour avoir

$$\frac{1}{T_e} \geq 2\xi_0,$$

alors le spectre du signal échantillonné x_e est assez riche pour reconstruire exactement le signal initial x . Il y a donc deux cas possible

- $1/T_e < 2\xi_0$ et dans ce cas on observe un phénomène de recouvrement ou repliement spectral et x ne peut pas être reconstruit (sans hypothèses supplémentaires),
- $1/T_e \geq 2\xi_0$ et alors on peut reconstruire le signal.

1. Pour bien comprendre ce qui va suivre je conseille fortement de voir la Figure 6.1 !

la fréquence limite $2\xi_0$ est appelée fréquence de Nyquist. Plus précisément on a le résultat suivant :

Théorème 82. *Soit T_e une période d'échantillonnage et $(x(kT_e))_{k \in \mathbb{Z}}$ un échantillonnage d'un signal x à bande limitée dans $[-\xi_0, \xi_0]$. Alors*

1. *Si $1/T_e < 2\xi_0$ la suite $(x(kT_e))_{k \in \mathbb{Z}}$ ne permet pas de déterminer x (sans hypothèses supplémentaires).*
2. *Si $1/T_e = 2\xi_0$ alors on peut reconstruire le signal x via la formule*

$$x(t) = \sum_{k \in \mathbb{Z}} x(kT_e) \frac{\sin(2\xi_0(t - kT_e))}{2\xi_0(t - kT_e)}.$$

3. *Si $1/T_e > 2\xi_0$ il existe une infinité de formules de reconstruction de x à partir de la suite $(x(kT_e))_{k \in \mathbb{Z}}$.*

Remarque 83. *La théorie exposait précédemment ne s'applique que dans le cas d'un signal x à bande limitée. Que faire alors dans le cas où $\hat{x}$ ne peut être nulle sur tout intervalle de $\mathbb{R}$ (ce qui arrive toujours en pratique) ? Dans ce cas, il faut appliquer un filtre passe bas afin d'annuler le spectre de x à partir d'une certaine fréquence ξ_0 . Pour plus de détails je renvoie à l'UV SY06 !*

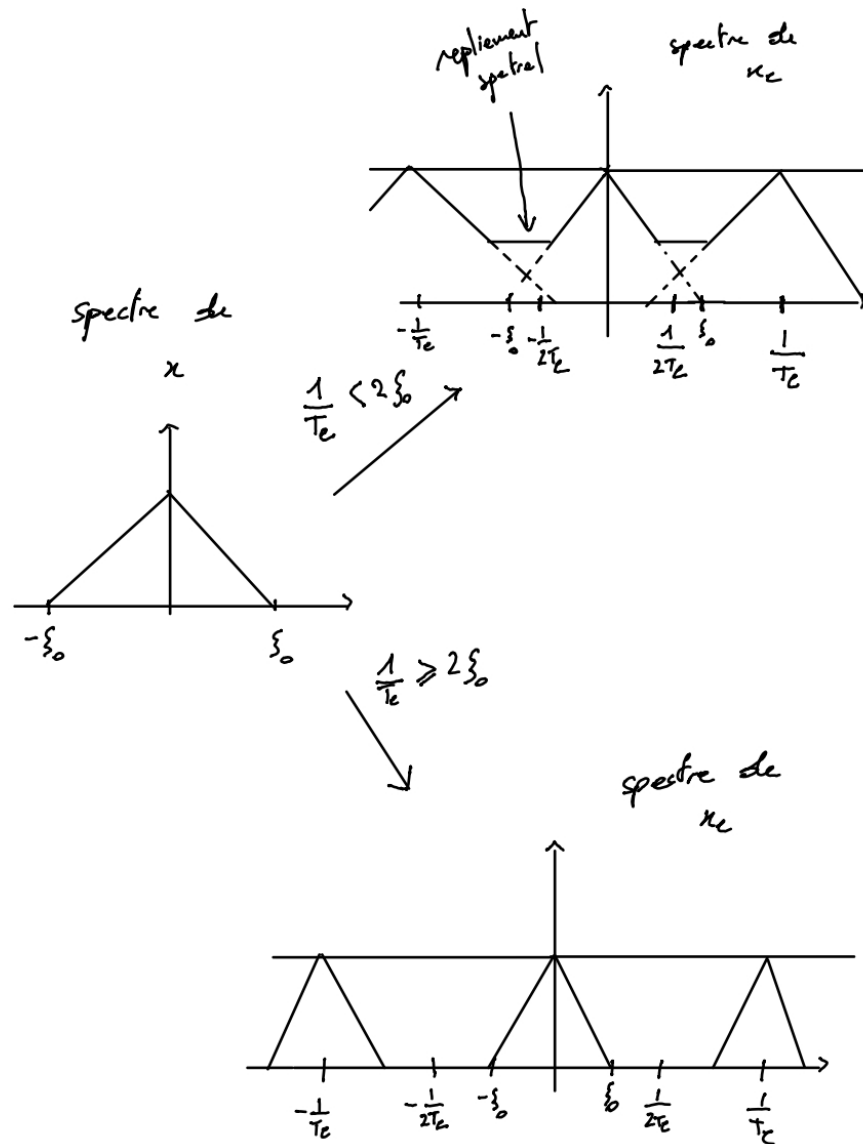


FIGURE 6.1 – Illustration du spectre du signal initial x et du signal échantillonné x_e en fonction de la période d'échantillonnage T_e .

Chapitre 7

Transformée de Laplace

Résumé. Dans ce chapitre nous allons étudier la transformée de Laplace. Cette transformation mathématique est une généralisation de la transformée de Fourier et permet en particulier de résoudre de nombreuses équations différentielles.

Références. Pour ce chapitre la plupart des résultats et preuves viennent de
— A. Lesfari. *Distributions, analyse de Fourier et transformation de Laplace*, Ellipses (2012).

7.1 Transformée de Laplace

Rentrons tout de suite dans le vif du sujet :

Définition 28. Soit $f : \mathbb{R} \rightarrow \mathbb{C}$ une fonction. On appelle transformée de Laplace de f la fonction notée $\mathcal{L}\{f(x)\}$ ou $F(p)$ de la variable complexe $p = \sigma + i\omega$ définie par

$$\mathcal{L}\{f(x)\} = F(p) = \int_0^{+\infty} f(x)e^{-px} dx = \int_0^{+\infty} f(x)e^{-(\sigma+i\omega)x} dx. \quad (7.1)$$

Cette première définition appelle plusieurs remarques et commentaires importants.

1. Dans la formule (7.1) l'intégrale ne porte pas sur tout l'intervalle $\mathbb{R}$. On va toujours considérer la transformée de Laplace des fonctions dites causales, i.e., des fonctions vérifiant $f(x) = 0$ pour $x < 0$. On peut rendre une fonction quelconque causale en multipliant cette fonction par la fonction de Heaviside (fonction notée H).
2. Bien entendu la formule (7.1) n'est pas bien définie pour toute fonction $f : \mathbb{R} \rightarrow \mathbb{C}$ causale. Cependant il est très important de comprendre, comme nous allons le voir dans les exemples, que la classe de fonctions pour laquelle il est possible de calculer la transformée de Laplace est beaucoup plus grande que dans le cas de la transformée de Fourier (grossièrement beaucoup plus grande que $L^1(\mathbb{R}_+)$). En ce sens la transformée de Laplace généralise la transformée de Fourier.

3. Il est important de comprendre le lien entre transformée de Fourier et transformée de Laplace. Soit f une fonction définie sur $\mathbb{R}$ possédant une transformée de Laplace, i.e., $\mathcal{L}\{f(x)\} < +\infty$ pour tout $p = \sigma + i\omega \in \mathbb{C}$. Pour $\sigma = 0$ alors on a

$$F(p) = F(i\omega) = \int_0^{+\infty} f(x)e^{-i\omega x} dx = \int_{\mathbb{R}} H(x)f(x)e^{-i\omega x} dx = \widehat{Hf}\left(\frac{\omega}{2\pi}\right).$$

Considérons maintenant $p = \sigma + i\omega \in \mathbb{C}$ alors

$$F(p) = \int_0^{+\infty} f(x)e^{-(\sigma+i\omega)x} dx = \int_{\mathbb{R}} H(x)f(x)e^{-\sigma x} e^{-i\omega x} dx = \mathcal{F}\{H(x)f(x)e^{-\sigma x}\}\left(\frac{\omega}{2\pi}\right).$$

Autrement dit la transformée de Laplace de f évaluée en $p = \sigma + i\omega$ est en fait la transformée de Fourier de la fonction $x \mapsto H(x)f(x)e^{-\sigma x}$ évaluée en la fréquence $\omega/2\pi$. De plus, si $\sigma > 0$, alors la fonction $x \in \mathbb{R}_+ \mapsto f(x)e^{-\sigma x}$ décroît rapidement vers 0 en $+\infty$. Grâce à ce facteur exponentielle décroissant on peut comprendre pourquoi la classe de fonction pour laquelle (7.1) possède un sens est plus grande que $L^1(\mathbb{R}_+)$ (voir l'exemple ci-dessous).

Exemple 43. On considère la fonction de Heaviside H . Ici $H \notin L^1(\mathbb{R})$ et on ne peut pas calculer sa transformée de Fourier. Cependant pour tout $p = \sigma + i\omega \in \mathbb{C}$ avec $\sigma > 0$ on a

$$F(p) = \int_0^{+\infty} H(x)e^{-px} dx = \int_0^{+\infty} e^{-px} dx = \left[-\frac{e^{-px}}{p}\right]_0^{+\infty} = \frac{1}{p}.$$

Ainsi $F(p)$ possède un sens dès que $\operatorname{Re}(p) > 0$, i.e., $F(p) < +\infty$ pour tout $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > 0$.

Proposition 84. Si l'intégrale (7.1) est finie pour un certain $p_0 \in \mathbb{C}$ avec $\operatorname{Re}(p_0) = \sigma_0$. Alors il en est de même pour tout nombre complexe $p \in \mathbb{C}$ avec $\operatorname{Re}(p) = \sigma \geq \sigma_0$.

Démonstration. On peut reformuler l'hypothèse sous la forme, il existe $p_0 \in \mathbb{C}$ tel que la fonction $x \mapsto f(x)e^{-p_0 x} dx \in L^1(\mathbb{R}_+)$ ou autrement dit

$$\int_0^{+\infty} |f(x)e^{-p_0 x}| dx < +\infty$$

En notant $p_0 = \sigma_0 + i\omega_0$ alors

$$\int_0^{+\infty} |f(x)e^{-p_0 x}| dx = \int_0^{+\infty} |f(x)| e^{-\sigma_0 x} dx < +\infty.$$

Il faut donc à présent démontrer que pour tout $p = \sigma + i\omega \in \mathbb{C}$ avec $\sigma \geq \sigma_0$ la fonction $x \mapsto f(x)e^{-px} \in L^1(\mathbb{R}_+)$. Or nous avons

$$|f(x)e^{-px}| = |f(x)|e^{-\sigma x} \leq |f(x)|e^{-\sigma_0 x}.$$

Donc en intégrant on obtient

$$\int_0^{+\infty} |f(x)e^{-px}| dx \leq \int_0^{+\infty} |f(x)|e^{-\sigma_0 x} dx < +\infty.$$

Ce qui conclut la preuve du résultat. □

Définition 29. Soit $f : \mathbb{R} \rightarrow \mathbb{C}$ une fonction causale. Le nombre

$$\sigma_0 = \inf\{\sigma \in \mathbb{R} : x \mapsto f(x)e^{-\sigma x} \in L^1(\mathbb{R}_+)\},$$

s'appelle l'abscisse de sommabilité de f .

Exemple 44. L'abscisse de sommabilité de la fonction de Heaviside est 0.

7.1.1 Propriétés de la transformée de Laplace

On énonce et démontre dans la suite des propriétés utiles pour calculer la transformée de Laplace d'une fonction.

Proposition 85 (Linéarité). Soient f et g deux fonction causales avec σ_0 et η_0 les abscisses de sommabilité de f et g et α et $\beta \in \mathbb{C}$, alors

$$\mathcal{L}\{\alpha f(x) + \beta g(x)\} = \alpha \mathcal{L}\{f(x)\} + \beta \mathcal{L}\{g(x)\} = \alpha F(p) + \beta G(p),$$

pour tout nombre $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > \max\{\sigma_0, \eta_0\}$.

Démonstration. C'est une simple conséquence de la linéarité de l'intégrale. $\square$

Proposition 86 (Translation). Soient $c > 0$ et f une fonction causale d'abscisse de sommabilité σ_0 , alors

$$\mathcal{L}\{f(x - c)\} = e^{-cp} \mathcal{L}\{f(x)\} = e^{-cp} F(p),$$

pour tout nombre $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > \sigma_0$.

Démonstration. On considère la fonction

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

La fonction g est causale et on a

$$\begin{aligned} \mathcal{L}\{g(x)\} &= \int_0^{+\infty} g(x)e^{-px} dx \\ &= \int_0^c g(x)e^{-px} dx + \int_c^{+\infty} g(x)e^{-px} dx \\ &= \int_c^{+\infty} f(x - c)e^{-px} dx \\ &\stackrel{y=x-c}{=} \int_0^{+\infty} f(y)e^{-p(y+c)} dy \\ &= e^{-cp} F(p), \end{aligned}$$

ce qui conclut la démonstration. $\square$

Proposition 87. Soient $\alpha \in \mathbb{C}$ et f une fonction causale d'abscisse de sommabilité σ_0 , alors

$$\mathcal{L}\{f(x)e^{-\alpha x}\} = F(p + \alpha), \quad \text{si } \operatorname{Re}(p + \alpha) > \sigma_0.$$

Démonstration. On a directement

$$\mathcal{L}\{f(x)e^{-\alpha x}\} = \int_0^{+\infty} f(x)e^{-\alpha x} e^{-px} dx = \int_0^{+\infty} f(x)e^{-(p+\alpha)x} dx = F(p + \alpha).$$

Ce qui achève la preuve du résultat. $\square$

Proposition 88 (Changement d'échelle). Soient $c > 0$ et f une fonction causale d'abscisse de sommabilité σ_0 , alors

$$\mathcal{L}\{f(cx)\} = \frac{1}{c} F\left(\frac{p}{c}\right), \quad \text{si } \operatorname{Re}\left(\frac{p}{c}\right) > \sigma_0.$$

Démonstration. On a directement que

$$\begin{aligned} \mathcal{L}\{f(cx)\} &= \int_0^{+\infty} f(cx)e^{-px} dx \\ &\stackrel{y=cx}{=} \int_0^{+\infty} f(y)e^{-\frac{p}{c}y} \frac{dy}{c} = \frac{1}{c} F\left(\frac{p}{c}\right). \end{aligned}$$

Ce qui termine la preuve du résultat. $\square$

Proposition 89. Soient f et g deux fonctions causales avec σ_0 et η_0 les abscisses de sommabilité de f et g , alors

$$\mathcal{L}\{f * g(x)\} = \mathcal{L}\{f(x)\} \mathcal{L}\{g(x)\} = F(p) G(p),$$

pour tout $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > \max\{\sigma_0, \eta_0\}$.

Démonstration. Admise. $\square$

La proposition suivante est très utile en pratique pour résoudre des équations différentielles.

Proposition 90. Soit f une fonction causale, d'abscisse de sommabilité σ_0 et de classe C^1 sur $[0, +\infty[$ pour laquelle il existe une constante $C > 0$ telle que

$$|f(x)| \leq Ce^{\sigma_0 x} \quad \forall x \geq \sigma_0. \quad (7.2)$$

Alors

$$\mathcal{L}\{f'(x)\} = p\mathcal{L}\{f(x)\} - f(0) = pF(p) - f(0), \quad (7.3)$$

pour tout nombre $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > \sigma_0$.

Remarque 91. Dans nos applications nous considérerons que l'hypothèse (7.2) est toujours vérifiée.

Démonstration. Soit $u \in \mathbb{R}$ avec $u > \sigma_0$. On commence par écrire que pour tout $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > \sigma_0$ on a

$$\mathcal{L}\{f'(x)\} = \int_0^{+\infty} f'(x)e^{-px} dx = \lim_{\substack{u \rightarrow +\infty \\ \varepsilon \rightarrow 0}} \int_{\varepsilon}^u f'(x)e^{-px} dx.$$

En appliquant une IPP on obtient

$$\begin{aligned} \int_{\varepsilon}^u f'(x)e^{-px} dx &= [f(x)e^{-px}]_{\varepsilon}^u + p \int_{\varepsilon}^u f(x)e^{-px} dx \\ &= f(u)e^{-pu} - f(\varepsilon)e^{-p\varepsilon} + p \int_{\varepsilon}^u f(x)e^{-px} dx. \end{aligned}$$

Ici comme $u \geq \sigma_0$ alors pour tout $p = \sigma + i\omega$ avec $\sigma > \sigma_0$ on a

$$|f(u)e^{-pu}| \leq Ce^{\sigma_0 u} |e^{-pu}| = Ce^{\sigma_0 u} e^{-\sigma u} = Ce^{(\sigma_0 - \sigma)u}.$$

Comme $\sigma > \sigma_0$ alors $\lim_{u \rightarrow +\infty} f(u)e^{-pu} = 0$. Par ailleurs par continuité de f en 0 on obtient

$$\lim_{\varepsilon \rightarrow 0} f(\varepsilon)e^{-p\varepsilon} = f(0).$$

On déduit alors du théorème de convergence dominée de Lebesgue que

$$\mathcal{L}\{f'(x)\} = p\mathcal{L}\{f(x)\} - f(0) = pF(p) - f(0),$$

ce qui conclut la preuve du résultat. □

Par récurrence on en déduit que si f est une fonction causale, d'abscisse de sommabilité σ_0 et de classe C^k avec $k \geq 1$ sur $[0, +\infty[$ (vérifiant une hypothèse du type (7.2) pour $f, f', \dots, f^{(k)}$), alors

$$\mathcal{L}\{f^{(k)}\} = p^k F(p) - p^{k-1}f(0) - p^{k-2}f'(0) - \dots - pf^{(k-2)}(0) - f^{(k-1)}(0), \quad (7.4)$$

pour tout $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > \sigma_0$. On peut réécrire cette formule sous la forme plus compacte

$$\mathcal{L}\{f^{(k)}\} = p^k F(p) - \sum_{j=0}^{k-1} p^{k-j-1} f^{(j)}(0),$$

avec la convention $f^{(0)} = f$.

7.1.2 Exemples de calculs

On donne dans la suite le détail du calcul de la transformée de Laplace de certaines fonctions (voir Section 7.3.1 pour un tableau plus complet).

1. On a déjà vu la transformée de Laplace de la fonction de Heaviside, on a obtenu pour tout $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > 0$

$$\mathcal{L}\{H(x)\} = \frac{1}{p},$$

son abscisse de sommabilité est 0.

2. Soit f_1 la fonction « rampe » définie par

$$f_1(x) = xH(x), \quad \forall x \in \mathbb{R}.$$

Soit $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > 0$ on a

$$\begin{aligned} F_1(p) &= \int_0^{+\infty} xe^{-px} dx = \left[-\frac{xe^{-px}}{p} \right]_0^{+\infty} + \int_0^{+\infty} \frac{e^{-px}}{p} dx \\ &= \left[-\frac{e^{-px}}{p^2} \right]_0^{+\infty}. \end{aligned}$$

Donc on obtient

$$F_1(p) = \frac{1}{p^2},$$

et l'abscisse de sommabilité de f_1 est 0. On en déduit directement (en effectuant n IPP) que la fonction f_n définie par $f_n(x) = x^n H(x)$ pour $n \geq 1$ vérifie

$$F_n(p) = \frac{n!}{p^{n+1}},$$

pour tout $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > 0$ (l'abscisse de sommabilité de f_n est 0).

3. Soit g la fonction exponentielle causale définie par

$$g(x) = e^{ax} H(x), \quad \forall x \in \mathbb{R}, a > 0.$$

Pour tout $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > a$ on a

$$G(p) = \int_0^{+\infty} e^{(a-p)x} dx = \left[\frac{e^{(a-p)x}}{a-p} \right]_0^{+\infty} = \frac{1}{p-a},$$

l'abscisse de sommabilité de g est a .

4. Soit r_1 la fonction sinusoidale causale définie par

$$r_1(x) = \sin(\omega x) H(x), \quad \forall x \in \mathbb{R}, \omega \in \mathbb{R}$$

On a pour tout $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > 0$

$$\begin{aligned} R_1(p) &= \int_0^{+\infty} \frac{e^{i\omega x} - e^{-i\omega x}}{2i} e^{-px} dx \\ &= \frac{1}{2i} \left(\int_0^{+\infty} e^{-(p-i\omega)x} dx - \int_0^{+\infty} e^{-(p+i\omega)x} dx \right) \\ &= \frac{1}{2i} \left(\mathcal{L}\{H(x)\}(p-i\omega) - \mathcal{L}\{H(x)\}(p+i\omega) \right) \\ &= \frac{1}{2i} \left(\frac{1}{p-i\omega} - \frac{1}{p+i\omega} \right) \\ &= \frac{1}{2i} \frac{2i\omega}{p^2 + \omega^2}. \end{aligned}$$

Ainsi on obtient pour tout $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > 0$

$$R_1(p) = \frac{\omega}{p^2 + \omega^2},$$

l'abscisse de sommabilité de r_1 est 0.

5. Soit r_2 la fonction causale définie par

$$r_2(x) = \cos(\omega x) H(x), \quad \forall x \in \mathbb{R}, \omega \in \mathbb{R}$$

On a pour tout $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > 0$

$$\begin{aligned} R_2(p) &= \int_0^{+\infty} \frac{e^{i\omega x} + e^{-i\omega x}}{2} e^{-px} dx \\ &= \frac{1}{2} \left(\int_0^{+\infty} e^{-(p-i\omega)x} dx + \int_0^{+\infty} e^{-(p+i\omega)x} dx \right) \\ &= \frac{1}{2} \left(\mathcal{L}\{H(x)\}(p-i\omega) + \mathcal{L}\{H(x)\}(p+i\omega) \right) \\ &= \frac{1}{2} \left(\frac{1}{p-i\omega} + \frac{1}{p+i\omega} \right) \\ &= \frac{1}{2} \frac{2p}{p^2 + \omega^2}. \end{aligned}$$

Ainsi on obtient pour tout $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > 0$

$$R_2(p) = \frac{p}{p^2 + \omega^2},$$

l'abscisse de sommabilité de r_1 est 0.

6. Soit r_3 la fonction sinusoidale amortie causale définie par

$$r_3(x) = e^{-ax} \sin(\omega x) H(x), \quad a > 0, \omega \in \mathbb{R}.$$

Soit $p \in \mathbb{C}$ avec $\operatorname{Re}(p) > 0$ alors

$$\begin{aligned} R_3(p) &= \int_0^{+\infty} \frac{e^{i\omega x} - e^{-i\omega x}}{2i} e^{-ax} e^{-px} dx \\ &= \frac{1}{2i} \left(\int_0^{+\infty} e^{-(a-i\omega+p)x} dx - \int_0^{+\infty} e^{-(a+i\omega+p)x} dx \right) \\ &= \frac{1}{2i} \left(\left[-\frac{e^{-(a-i\omega+p)x}}{a-i\omega+p} \right]_0^{+\infty} - \left[-\frac{e^{-(a+i\omega+p)x}}{a+i\omega+p} \right]_0^{+\infty} \right) \\ &= \frac{1}{2i} \left(\frac{1}{a-i\omega+p} - \frac{1}{a+i\omega+p} \right). \end{aligned}$$

Il reste maintenant à simplifier cette expression. On obtient

$$\begin{aligned} R_3(p) &= \frac{1}{2i} \frac{2i\omega}{(a-i\omega+p)(a+i\omega+p)} \\ &= \frac{\omega}{a^2 + ai\omega + ap - ai\omega + \omega^2 - ip\omega + ap + pi\omega + p^2} \\ &= \frac{\omega}{a^2 + 2ap + p^2 + \omega^2}. \end{aligned}$$

En conclusion on en déduit que

$$R_3(p) = \frac{\omega}{(a+p)^2 + \omega^2},$$

l'abscisse de sommabilité est 0.

7.2 Résolution d'équations différentielles

Les équations différentielles interviennent très souvent dans les sciences de l'ingénieur. Il est ainsi important de comprendre comment résoudre ces équations. Par exemple en automatique (SY14) la dynamique des systèmes linéaires est gouvernée par des équations différentielles linéaires à coefficients constants. Ces équations interviennent également par exemple en traitement du signal (SY06). Un outil indispensable à l'ingénieur pour résoudre ces équations est la transformée de Laplace. En effet la transformée de Laplace permet de réduire la résolution d'équations différentielles à la résolutions d'équations algébriques. On décrit ci-dessous la méthode générale et on l'illustre au travers de deux exemples.

On considère l'EDO **linéaire** d'ordre n à **coefficients constants** suivante

$$a_0 x^{(n)}(t) + a_1 x^{(n-1)}(t) + \dots + a_{n-1} x'(t) + a_n x(t) = y(t), \quad t > 0$$

avec

$$x(0) = x_0^0, x'(0) = x_0^1, \dots, x^{(n-1)}(0) = x_0^{n-1}.$$

Le but est de déterminer la solution de cette équation différentielle. Pour ce faire on note par $X(p)$ et $Y(p)$ les transformées de Laplace de x et y , on utilise la linéarité de la transformée de Laplace et on applique la Proposition 90 et l'égalité (7.4). On obtient alors

$$\begin{aligned} & a_0 \left(p^n X(p) - p^{n-1} x(0) - p^{n-2} x'(0) - \dots - x^{(n-1)}(0) \right) \\ & + a_1 \left(p^{n-1} X(p) - p^{n-2} x(0) - p^{n-3} x'(0) - \dots - x^{(n-2)}(0) \right) + \dots \\ & + a_{n-1} \left(p X(p) - x(0) \right) + a_n X(p) = Y(p). \end{aligned}$$

En regroupant les termes on obtient

$$X(p) = \frac{Y(p)}{a_0 p^n + a_1 p^{n-1} + \dots + a_{n-1} p + a_n} - \frac{U(p)}{a_0 p^n + a_1 p^{n-1} + \dots + a_{n-1} p + a_n}, \quad (7.5)$$

avec

$$\begin{aligned} U(p) &= a_0 \left(p^{n-1} x(0) + p^{n-2} x'(0) + \dots + x^{(n-1)}(0) \right) \\ &+ a_1 \left(p^{n-2} x(0) + p^{n-3} x'(0) + \dots + x^{(n-2)}(0) \right) + \dots + a_{n-1} x(0). \end{aligned}$$

Pour en déduire l'expression de x solution de l'EDO il suffit alors d'utiliser la transformée de Laplace inverse.

Remarques importantes. Voici deux remarques utiles.

- Comme dans le cas de la transformée de Fourier il existe une formule d'inversion permettant de passer de la transformée de Laplace d'une fonction à la fonction initiale. Cependant cette formule est relativement complexe à comprendre et nécessite des outils qui ne sont pas au programme de cette UV. Ainsi en pratique on utilise des tables. En particulier pour résoudre les exercices en liant avec les équations différentielles nous donnerons les formules permettant d'inverser la transformée de Laplace.
- Par ailleurs en pratique avant d'appliquer la transformée de Laplace inverse on utilise des décompositions en éléments simples afin de simplifier le membre de droite de (7.5) (voir l'exemple ci-dessous).

Exemple 45. On considère l'EDO suivante

$$x''(t) + 5x'(t) + 4x(t) = 0, \quad t > 0, \quad (7.6)$$

avec

$$x(0) = 2, x'(0) = -5.$$

En appliquant la méthode précédente on obtient

$$p^2 X(p) - px(0) - x'(0) + 5(pX(p) - x(0)) + 4X(p) = 0.$$

En utilisant les conditions initiales on peut alors réécrire cette équation sous la forme

$$(p^2 + 5p + 4)X(p) = 2p - 5 + 10 = 2p + 5,$$

et donc

$$X(p) = \frac{2p + 5}{p^2 + 5p + 4}.$$

Comme déjà annoncé, avant d'appliquer la transformée de Laplace inverse on va décomposer l'expression de $X(p)$ en éléments simples. Ici comme $\deg(2p+5) = 1 < 2 = \deg(p^2+5p+4)$ et que -1 et -4 sont des racines de $p^2 + 5p + 4$ on cherche alors des coefficients a et $b \in \mathbb{R}$ vérifiant

$$X(p) = \frac{2p + 5}{(p + 1)(p + 4)} = \frac{a}{p + 1} + \frac{b}{p + 4}.$$

On obtient alors facilement que $a = b = 1$ et donc

$$X(p) = \frac{1}{p + 1} + \frac{1}{p + 4}.$$

Enfin la transformée de Laplace inverse d'une fonction de la forme $p \mapsto 1/(p + c)$ est la fonction $t \mapsto e^{-ct}$. Ainsi on en déduit que

$$x(t) = e^{-t} + e^{-4t}, \quad t > 0,$$

et on vérifie facilement que $x(0) = 2$, $x'(0) = -5$ et que la fonction x vérifie (7.6).

Exemple 46. On considère l'EDO suivante

$$x''(t) - 2x'(t) + x(t) = 0, \quad t > 0, \tag{7.7}$$

avec $x(0) = 1$ et $x'(0) = -1$. En utilisant la transformée de Laplace on obtient l'équation

$$p^2 X(p) - px(0) - x'(0) - 2(pX(p) - x(0)) + X(p) = 0.$$

On en déduit la relation

$$(p^2 - 2p + 1)X(p) = p - 3,$$

autrement dit

$$X(p) = \frac{p-3}{p^2-2p+1}.$$

ici le dénominateur admet une racine double en $p=1$, on a

$$X(p) = \frac{p-3}{(p-1)^2} = \frac{p}{(p-1)^2} - 3 \frac{1}{(p-1)^2}.$$

Il reste à appliquer les formules suivantes

$$\begin{aligned} \frac{p}{(p+a)^2} &\xrightarrow{\mathcal{L}^{-1}} e^{-at}(1-at), \\ \frac{1}{(p+a)^2} &\xrightarrow{\mathcal{L}^{-1}} te^{-at}, \end{aligned}$$

avec $a = -1$. On en déduit que la solution de (7.7) vérifiant les conditions initiales $x(0) = 1$ et $x'(0) = -1$ est

$$x(t) = e^t(1+t) - 3te^t, \quad t \geq 0.$$

7.3 Compléments

7.3.1 Tableau de certaines transformée de Laplace

Dans la suite pour alléger les notations on ne note pas la fonction de Heaviside. Par exemple dans la deuxième ligne il faut comprendre que la fonction constante égale à un est multipliée par H .

Fonction causale	Transformée de Laplace
1	$\frac{1}{p}$
x^n ($n \geq 1$)	$\frac{n!}{p^{n+1}}$
e^{ax} ($a > 0$)	$\frac{1}{p - a}$
e^{-ax} ($a > 0$)	$\frac{1}{p + a}$
$\sin(\omega x)$ ($\omega \in \mathbb{R}$)	$\frac{\omega}{p^2 + \omega^2}$
$\cos(\omega x)$ ($\omega \in \mathbb{R}$)	$\frac{p}{p^2 + \omega^2}$
$e^{-at} \sin(\omega t)$ ($a > 0, \omega \in \mathbb{R}$)	$\frac{\omega}{(p + a)^2 + \omega^2}$
$e^{-at} \cos(\omega t)$ ($a > 0, \omega \in \mathbb{R}$)	$\frac{p + a}{(p + a)^2 + \omega^2}$

TABLE 7.1 – Transformée de Laplace.