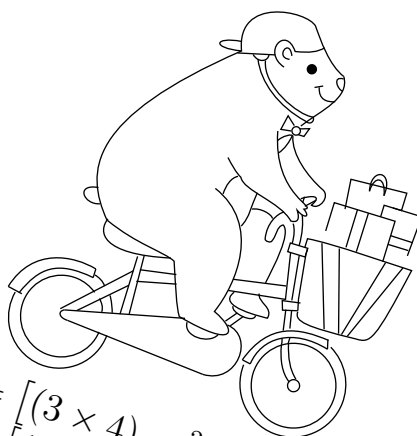


COUPLAGE INTÉGRAL POUR L'AGRÉGATION EXTERNE

- session 2026 -

—



$$\begin{aligned}
 (3x + ay)(4x - bz) + \alpha xy &= [(3 \times 4) \cdot x^2 + (4 \times a) \cdot xy + (-3 \times b) \cdot xz + (a \times (-b)) \cdot yz] + \alpha xy \\
 &= [(3 \times 4) \cdot x^2 + (4 \times a) \cdot xy + (-3 \times b) \cdot xz - (a \times b) \cdot yz] + \alpha xy \\
 &= [(3 \times 4) \cdot x^2 + (4 \times a) \cdot xy + (-3 \times b) \cdot xz - (a \times b) \cdot yz] + \alpha xy \\
 &= [(12) \cdot x^2 + (4a) \cdot xy + (-3b) \cdot xz - (ab) \cdot yz] + \alpha xy \\
 &= [(12) \cdot x^2 + (4a) \cdot xy - (3b) \cdot xz - (ab) \cdot yz] + \alpha xy \\
 &= [12x^2 + 4axy - 3bxz - abyz] + \alpha xy \\
 &= 12x^2 + 4axy - 3bxz - abyz + \alpha xy \\
 &= 12x^2 + (4a + \alpha)xy - 3bxz - abyz \\
 &= \boxed{12x^2 + (4a + \alpha)xy - 3bxz - abyz}
 \end{aligned}$$

*Exemple de développement

Prélude

Bienvenue, camarade agrégatif.ve ! Si vous vous trouvez ici, j'imagine que vous préparez l'agrégation, alors en premier lieu, laissez-moi vous souhaiter bon courage.

J'ai cherché à rassembler dans ce document la totalité des développements qui faisaient partie de mon couplage lorsque j'ai moi-même passé l'agrégation, dans l'espoir que peut-être ils puissent servir aux suivant.e.s.

Mon objectif ici a été d'essayer de rédiger mes développements avec le plus de détails possible, dans un souci de clarté mais aussi parce que leur rédaction faisait partie de ma routine pour apprendre un développement. De ce fait, il arrivera parfois qu'un développement ainsi présenté soit bien trop long pour être déroulé pendant les quinze minutes réglementaires. J'aurai parfois ajouté quelques indications sur la manière dont je les présentais à l'oral concrètement, mais pas toujours.

Vous trouverez pour chaque développement les recasages que j'estime possibles pour chaque développements. Ils ne correspondent pas aux recasages que je faisais effectivement (il y aurait bien trop de collisions !) mais devraient vous donner un peu plus de latitude pour composer votre couplage. A prendre avec des pincettes tout de même, je ne prétends pas avoir l'omniscience de l'agrégatif. A titre informatif, vous trouverez à la toute fin de ce document le couplage final que j'ai emporté à Strasbourg pour les oraux.

En suppléments, j'ai essayé (quand j'étais inspiré) d'ajouter pour chaque développement des remarques sur des applications possibles, des prolongements pour anticiper les questions et / ou des remarques sur le cadre.

Enfin, mais c'est relativement subjectif, j'ai essayé de noter la difficulté de chaque développement avec un nombre d'étoiles. Le critère principal que j'utilisais pour choisir le nombre d'étoiles était la difficulté conceptuelle du développement et le niveau académique des prérequis. Un développement avec moins d'étoiles peut avoir l'air plus dur dans le sens où il peut être difficile à faire rentrer en quinze minutes, comporter de nombreux calculs dans lesquels il est facile de se prendre les pieds, etc. Voilà à peu près la table d'étalonnage :

★	accessible en prépa / en L2
★★	accessible en L3
★★★	rien que des outils du programme, mais conceptuellement difficile
★★★★	déborde du programme tout en restant dans « son adhérence »
★★★★★	résolument hors-programme

C'est tout à fait indicatif et il est très possible que certaines indications que j'ai mises ne correspondant pas vraiment à cette table (d'autant que c'est un peu difficile de dire que quelque chose est « accessible en L3 » étant donné qu'il n'y a pas UN programme de L3 unifié...). J'espère qu'au moins ça vous donnera une petite idée a priori.

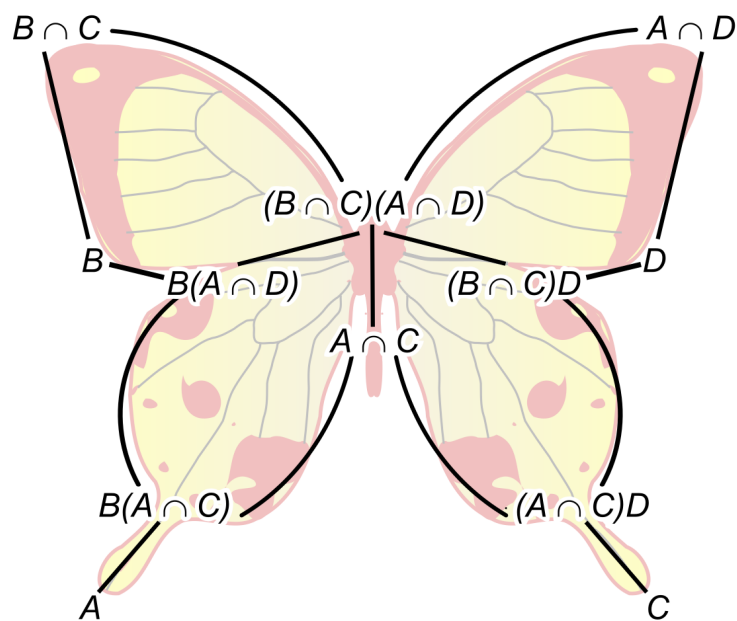
Table des matières

1	Développements d'algèbre	4
1.1	Algorithme de Berlekamp ****	5
1.2	Automorphismes de $\mathfrak{S}_n$ **	9
1.3	Autour de la cyclotomie ** / ****	13
1.4	Classification des formes quadratiques sur les corps finis **	20
1.5	Critère de Solovay-Strassen ****	23
1.6	Décomposition de Dunford-Jordan-Chevalley - une méthode algorithmique **	28
1.7	Générateurs de $O(q)$ et $SO(q)$ **	32
1.8	Norme d'un élément algébrique **** / *****	35
1.9	Par cinq points passe une conique ***	44
1.10	Probabilités sur $GL_2(\mathbb{F}_q)$ **	50
1.11	Réduction de Frobenius ***	54
1.12	Réduction de Smith - anneaux principaux *****	60
1.13	Théorème de Lie-Kolchin ***	68
1.14	Théorème de Kronecker ***	72
2	Développements d'analyse	75
2.1	Calcul d'une transformée de Fourier par le théorème de résidus **	76
2.2	Différentielle du flot d'une équation différentielle autonome *****	80
2.3	Echantillonnage de Schannon ***	89
2.4	Formule sommatoire de Poisson et inversion de Fourier dans L^2 **	92
2.5	Lemme de la grenouille *	96
2.6	Holomorphie sous l'intégrale et problème des moments ***	99
2.7	Optimisation dans un espace de Hilbert et théorème de Banach-Alaoglu ****	103
2.8	Probabilité(s) invariante(s) d'une chaîne de Markov ****	108
2.9	Quadrature de Gauss-Legendre *	113
2.10	Séries de Fourier divergentes par le théorème de Banach-Steinhaus ****	118
2.11	Théorème de Hahn-Banach - forme géométrique ***	122
2.12	Théorème de Furstenberg-Kesten ***	128
2.13	Théorème de Lévy et théorème central limite ****	133
2.14	Théorème de Liapounov ***	139
2.15	Théorème de Riesz-Fréchet-Kolmogorov *****	144
2.16	Théorème taubérien fort ***	150

3	Développements mixtes	155
3.1	Déterminant circulant et suite de polygônes *	156
3.2	Sous-groupes petits de $GL_n(\mathbb{R})$ **	160
3.3	Lemme de Morse ****	163
3.4	Nombre de cycles d'une permutation aléatoire *	168
3.5	Sous-espaces de C^0 stables par translations **	171
3.6	Théorème de Kakutani et sous-groupes compacts de $GL(E)$ ****	175
3.7	Théorème de min-max de Courant-Fischer et continuité des valeurs propres hermitiennes *	181
A	Couplage effectif	185
	Bibliographie	189

Chapitre 1

Développements d'algèbre



* « Techniques d'algèbre en arts plastiques »
Crédits image : Claudio Rocchini

1.1 Algorithme de Berlekamp ★★ ★★

Recasages possibles : 122, 123, 141, 142, 148

Lorsque K est un corps, on utilise fréquemment le fait que l'anneau $K[X]$ est factoriel. Malheureusement, calculer la décomposition d'un polynôme en produit d'irréductibles est un problème généralement compliqué. Il se trouve que dans le cas où K est un corps fini, Berlekamp a découvert un algorithme qui permet de calculer cette décomposition pour n'importe quel polynôme (de façon plus intelligente que de simplement tester tous les diviseurs potentiels).

Dans ce développement, on va donner et étudier l'algorithme de Berlekamp pour des entrées sans facteurs carrés (c'est déjà bien assez long comme ça). Vous trouverez dans [BMP05] tout ce qu'il faut pour traiter le cas général (ainsi d'ailleurs que la totalité du développement très bien présenté).

Soit p un nombre premier et q une puissance de p . On désigne par $\mathbb{F}_q$ un corps à q -éléments. Etant donné $P \in \mathbb{F}_q[X]$, on notera $\langle P \rangle$ l'idéal de $\mathbb{F}_q[X]$ engendré par P . Lorsque Q est un autre polynôme, on notera $\overline{Q}^P$ le projeté de Q dans $\mathbb{F}_q[X] / \langle P \rangle$. Notons qu'à chaque fois qu'on se donnera un élément $\overline{Q}^P \in \mathbb{F}_q[X] / \langle P \rangle$, on fixera implicitement un représentant $Q \in \mathbb{F}_q[X]$.

Un peu de contexte

Dans cette première partie, on va poser un certain nombre de notations et préciser le cadre de travail qui va nous permettre d'étudier l'algorithme.

Soit $P \in \mathbb{F}_q[X]$ un polynôme, qu'on décompose en produit d'irréductibles :

$$P = \prod_{i=1}^r P_i^{\nu_i} \quad (1)$$

Le lemme chinois fournit un isomorphisme de $\mathbb{F}_q$ -algèbres :

$$\begin{aligned} \Phi : \underbrace{\mathbb{F}_q[X] / \langle P \rangle}_{=: \mathcal{A}} &\xrightarrow{\sim} \prod_{i=1}^r \underbrace{\mathbb{F}_q[X] / \langle P_i^{\nu_i} \rangle}_{=: \mathcal{A}_i} \\ \overline{Q}^P &\mapsto \left(\overline{Q}^{P_i^{\nu_i}} \right)_{1 \leq i \leq r} \end{aligned}$$

Remarquons à ce stade qu'une $\mathbb{F}_q$ -algèbre est, par définition, un anneau A muni d'un morphisme d'anneaux dit «structural» $\mathfrak{s}_A : \mathbb{F}_q \rightarrow A$, injectif puisque $\mathbb{F}_q$ est un corps. Ceci permet d'identifier $\mathbb{F}_q$ à $\mathfrak{s}_A(\mathbb{F}_q)$, qui est un sous-corps de A , identification qui est très couramment faite dans ce contexte. Le danger ici est que l'on manipule plusieurs $\mathbb{F}_q$ -algèbres, et donc $\mathbb{F}_q$ serait assimilé à des parties de différentes algèbres qui n'ont a priori rien à voir entre elles ! Comme on va souvent tirer des éléments de $\mathfrak{s}_A(\mathbb{F}_q)$ dans $\mathbb{F}_q$ par le morphisme structural, de sorte à pouvoir les comparer avec d'autres éléments issus de l'image de $\mathbb{F}_q$ dans d'autres algèbres, on gardera toujours la trace des morphismes structuraux. Cela occasionnera une certaine lourdeur dans les notations, mais qui est justifiée par un gain de clarté non négligeable.

On remarque également que tout $\mathbb{F}_q$ -algèbre peut être munie d'un endomorphisme de $\mathbb{F}_q$ -algèbre défini par le passage à la puissance q . On notera Frob_q ce morphisme dans le cas de la $\mathbb{F}_q$ -algèbre $\mathcal{A}$.

Dans toute la suite, on utilisera les notations introduites ici.

L'algorithme

Je propose de présenter l'algorithme sans plus de contexte, d'une seule traite, afin d'avoir chacune de ses étapes en tête. Tout ceci paraîtra obscure de prime abord, et c'est bien normal, car il est sacrément astucieux ! La suite du développement consistera à en prouver la correction et la terminaison.

Algorithme : Berlekamp

Entrée : $P \in \mathbb{F}_q[X]$ un polynôme sans facteurs carrés ⁽ⁱ⁾

Sortie : La liste des facteurs irréductibles de P .

1. Calculer s la dimension du $\mathbb{F}_q$ -espace vectoriel $\text{Ker}(\text{Frob}_q - \text{Id}_{\mathcal{A}})$
2. Si $s = 1$, retourner P .
3. Sinon, prendre $\bar{V}^P \in \text{Ker}(\text{Frob}_q - \text{Id}_{\mathcal{A}}) \setminus \mathfrak{s}_{\mathcal{A}}(\mathbb{F}_q)$. Retourner :

$$\{\text{Berlekamp}(\text{pgcd}(V - \alpha, P)), \alpha \in \mathbb{F}_q \mid \text{pgcd}(V - \alpha, P) \neq 1\} \quad (2)$$

Etude de l'algorithme

On va prouver d'un seul coup que l'algorithme est correcte est termine en étudiant chacune de ses trois étapes.

Lemme 1. *L'entier s est r le nombre de facteurs irréductibles de P .*

Démonstration. Soit $Q \in \mathbb{F}_q[X]$. On observe la suite d'équivalence suivante :

$$\left[(\bar{Q}^P)^q = \bar{Q}^P \right] \iff \left[\Phi \left(\bar{Q}^P \right)^q = \Phi \left(\bar{Q}^P \right) \right] \quad (3)$$

$$\iff \left[\forall i \in \llbracket 1, r \rrbracket, \left(\bar{Q}^{P_i} \right)^q = \bar{Q}^{P_i} \right] \quad (4)$$

$$\iff \left[\forall i \in \llbracket 1, r \rrbracket, \bar{Q}^{P_i} \in \mathfrak{s}_{\mathcal{A}_i}(\mathbb{F}_q) \right] \quad (5)$$

La dernière équivalence mérite explications. Comme les P_i sont irréductibles, les $\mathbb{F}_q$ -algèbres $\mathcal{A}_i := \mathbb{F}_q[X] / \langle P_i \rangle$ sont en réalité des extensions de corps de $\mathbb{F}_q$. $\mathfrak{s}_{\mathcal{A}_i}(\mathbb{F}_q)$ est exactement l'ensemble des points fixes du morphisme de Frobenius $x \mapsto x^q$ ⁽ⁱⁱ⁾, ce qui justifie cette dernière équivalence. Comme Φ est un isomorphisme de $\mathbb{F}_q$ -algèbres, c'est en particulier un isomorphisme $\mathbb{F}_q$ -linéaire. Il induit donc un isomorphisme entre $\text{Ker}(\text{Frob}_q - \text{Id}_{\mathcal{A}})$ et $\Phi(\text{Ker}(\text{Frob}_q - \text{Id}_{\mathcal{A}}))$ (ce qui est la traduction de la première équivalence), et les dernières équivalences montrent qu'on a l'isomorphisme de $\mathbb{F}_q$ -espaces vectoriels :

$$\Phi(\text{Ker}(\text{Frob}_q - \text{Id}_{\mathcal{A}})) = \prod_{i=1}^r \mathfrak{s}_{\mathcal{A}_i}(\mathbb{F}_q) \cong \mathbb{F}_q^r \quad (6)$$

(i). C'est-à-dire que tous ses facteurs irréductibles dans $\mathbb{F}_q[X]$ sont présents avec multiplicité 1.

(ii). Tous les éléments de $\mathfrak{s}_{\mathcal{A}_i}(\mathbb{F}_q)$ sont de tels points fixes en vertu du théorème de Lagrange, et il ne peut y avoir plus de q tels points fixes puisqu'il s'agit des racines de $X^q - X$, qui est de degré q .

D'où l'espace propre $\text{Ker}(\text{Frob}_q - \text{Id}_{\mathcal{A}})$ est un $\mathbb{F}_q$ -espace vectoriel de dimension r . $\square$

Ainsi, l'étape 2 prend tout son sens ! Si l'entier s calculé est 1, P a un unique facteur irréductible, donc P est lui-même irréductible. Etudions la troisième étape.

On suppose que la troisième étape est atteinte, c'est-à-dire que $r > 1$. En conséquence, $\text{Ker}(\text{Frob}_q - \text{Id}_{\mathcal{A}}) \setminus \mathfrak{s}_{\mathcal{A}}(\mathbb{F}_q)$ est un ensemble non-vidé. Soit $\bar{V}^P$ l'un de ses éléments. On pose :

$$\forall i \in \llbracket 1, r \rrbracket, \alpha_i := \mathfrak{s}_{\mathcal{A}_i}^{-1}(\bar{V}^{P_i}) \in \mathbb{F}_q \quad (\text{iii}) \quad (7)$$

Lemme 2.

$$\prod_{\alpha \in \mathbb{F}_q} \text{pgcd}(V - \alpha, P) = P \quad (8)$$

Démonstration. Considérons $i \in \llbracket 1, r \rrbracket$, $\alpha \in \mathbb{F}_q$ et la suite d'équivalences :

$$[P_i | V - \alpha] \iff [V \equiv \alpha \pmod{P_i}] \quad (9)$$

$$\iff [\bar{V}^{P_i} = \mathfrak{s}_{\mathcal{A}_i}(\alpha)] \quad (10)$$

$$\iff [\alpha = \alpha_i] \quad (11)$$

Ainsi, on a l'égalité :

$$\prod_{\alpha \in \mathbb{F}_q} \text{pgcd}(V - \alpha, P) = \prod_{\alpha \in \mathbb{F}_q} \prod_{\substack{i=1 \\ \alpha_i = \alpha}}^r P_i \quad (12)$$

$$= \prod_{i=1}^r P_i \quad (13)$$

$$= P \quad (14)$$

$\square$

On est maintenant en passe de conclure ! Les facteurs $(\text{pgcd}(V - \alpha, P))$ sont soit 1, soit un produit des (P_i) . Par ailleurs, pour tout i , P_i est facteur d'un et d'un seul terme de la forme $\text{pgcd}(V - \alpha, P)$. De plus, tous ces termes sont des diviseurs **stricts** de P , car :

$$\text{pgcd}(V - \alpha, P) = P \iff P | V - \alpha \implies \bar{V}^P \in \mathfrak{s}_{\mathcal{A}}(\mathbb{F}_q) \quad (15)$$

donc le choix de V exclue ce cas. En d'autres termes, le nombre de facteurs irréductibles de $\text{pgcd}(V - \alpha, P)$ est inférieur strict à r , ce pour tout $\alpha \in \mathbb{F}_q$. Il suit immédiatement que l'algorithme termine, car à chaque appel récursif, le nombre de facteurs irréductibles du polynôme en entrée diminue strictement. L'algorithme est par ailleurs correct : on peut le prouver par récurrence sur r . Lorsque $r = 1$, l'algorithme s'arrête à l'étape 2 et renvoie la décomposition en irréductibles de

(iii). Par définition de $\bar{V}^P$, on a d'après le travail précédent que $\bar{V}^{P_i}$ est un élément de la copie de $\mathbb{F}_q$ dans $\mathcal{A}_i$. Comme on va vouloir comparer tous ces éléments, on les tire en arrière par le morphisme structural pour récupérer un élément de $\mathbb{F}_q$.

P . Si on suppose que $r > 1$ et que l'algorithme est correct pour toute entrée ayant au plus $r - 1$ facteurs irréductibles, alors par récurrence, l'algorithme renvoie la liste des facteurs irréductibles de tous les $(\text{pgcd } V - \alpha, P)$, qui à eux tous contiennent exactement les facteurs irréductibles de P . Ceci achève l'étude de l'algorithme de Berlekamp !

1.2 Automorphismes de $\mathfrak{S}_n \star \star$

Recasages possibles : 101, 103, 104, 105, 108

Considérons G un groupe quelconque. Tout élément de G s'identifie à un automorphisme de G via la formule :

$$\forall g \in G, \iota(g) : x \mapsto gxg^{-1} \quad (1)$$

Un tel automorphisme est qualifié d'intérieur, et nous noterons $\text{Int}(G)$ l'ensemble des tels automorphismes. Dans ce cadre, la flèche ι correspond également à l'action de G sur lui-même par conjugaison. Ainsi, quel que soit le groupe auquel on s'intéresse, on dispose d'une suite exacte^(iv) :

$$1 \longrightarrow Z(G) \hookrightarrow G \twoheadrightarrow \text{Int}(G) \longrightarrow 1 \quad (2)$$

Ainsi, dans le cas où le groupe est de centre trivial, on obtient un isomorphisme entre G et son groupe d'automorphismes intérieurs. Dans le cas du groupe symétrique $\mathfrak{S}_n$, il se passe quelque chose d'encore plus joli : tous les automorphismes de $\mathfrak{S}_n$ se réalisent comme un automorphisme intérieur, c'est-à-dire que pour $n \geq 3$, on a un isomorphisme $\mathfrak{S}_n \cong \text{Aut}(\mathfrak{S}_n)$! Tous les automorphismes, vraiment ? Eh non, il y a (de façon superbe ou catastrophique selon les gens) un phénomène exceptionnel qui se produit lorsque $n = 6$, mettant en défaut le théorème...

Dans ce développement, nous allons démontrer ce résultat d'une façon peut-être moins classique que ce qui est généralement présenté, qui repose beaucoup plus sur la théorie des groupes et moins sur la combinatoire. **En conséquence, cette preuve n'est pas présentable dans la 190.** En contrepartie, je pense qu'elle gagne un recasage impeccable dans la 103.

Il s'agit de la preuve alternative présentée par Perrin, qui à mon sens à l'énorme avantage d'être beaucoup plus naturelle (donc facile à retenir !) que la preuve combinatoire, mais elle est certainement plus exigeante en terme de théorie. J'apprécie particulièrement le fait que le cas $n = 6$ apparaisse de lui-même au cours de la preuve, en mettant en évidence un argument de structure de $\mathfrak{S}_n$, ce qui, je trouve, ne se produit pas dans la preuve combinatoire classique.

Soit $n \in \mathbb{N}^*$. On note $\mathfrak{S}_n$ le groupe symétrique et $\mathfrak{A}_n$ le groupe alterné. On notera un k -cycle de la façon compacte habituelle :

$$\sigma = (a_1 \ a_2 \ \dots \ a_k) \quad (3)$$

Etant donné g un élément d'un groupe G , on appellera *centralisateur* de g et on notera $C(g)$ le stabilisateur de g sous l'action par conjugaison de G sur lui-même, c'est-à-dire :

$$C(g) := \{h \in G \mid hgh^{-1} = g\} \quad (4)$$

En tant que stabilisateur, c'est un sous-groupe de G . Notre objectif est de démontrer le théorème suivant :

Théorème 3 (Automorphismes de $\mathfrak{S}_n$ ([Per], theo. 8.7)). *Les automorphismes de $\mathfrak{S}_n$ sont tous intérieurs, sauf peut-être pour $n = 6$* ^(v).

(iv). C'est la première (mais pas la dernière) fois que j'utilise une suite exacte dans ce développement. Snobbisme superflu dont on se passerait aisément, pourrait-on dire, ce que je ne peux pas vraiment nier. Disons que c'est la manière dont j'aime bien voir les choses et que je trouve que le formalisme des suites exactes apporte une vision nette et synthétique de ce genre de situations où apparaissent des noyaux et des quotients. Une fois qu'on a essayé, on y devient vite accroc. Mais vous pouvez complètement éviter de parler de suite exacte !

(v). Cet insolent $n = 6$, constante ridicule s'il en est, vous énerve ? Moi aussi, mais c'est la vie.

Un lemme efficace ([Per], prop. 8.8)

Toute la démonstration va reposer sur l'observation suivant :

Lemme 4. *Soit $\Phi \in \text{Aut}(\mathfrak{S}_n)$. Si l'image par Φ de toute transposition reste une transposition, alors Φ est intérieur.*

Démonstration. On suppose que pour toute transposition τ , $\Phi(\tau)$ est encore une transposition. Remarquons que les cas $n = 1$ et $n = 2$ sont triviaux, on supposera donc $n \geq 3$. Considérons la famille :

$$\forall i \in \llbracket 2, n \rrbracket, \tau_i := (1 \ i) \quad (5)$$

dont on sait qu'elle engendre $\mathfrak{S}_n$. Etant donnés $i \neq j$, τ_i et τ_j ne commutent pas car partagent 1 dans leurs supports. En conséquence, les transposition $\Phi(\tau_i)$ et $\Phi(\tau_j)$ ne commutent pas non plus, et nécessairement, leurs supports ne sont pas disjoints. Notons donc :

$$\Phi(\tau_2) = (\alpha_1 \ \alpha_2) \quad \Phi(\tau_3) = (\alpha_1 \ \alpha_3) \quad (6)$$

Soit $i \in \llbracket 2, n \rrbracket$. $\Phi(\tau_i)$ est nécessairement de la forme $(\alpha_1 \ \alpha_i)$ ou $(\alpha_2 \ \alpha_3)$. Montrons que le deuxième cas est exclu en supposant par l'absurde $\Phi(\tau_i) = (\alpha_2 \ \alpha_3)^{(vi)}$. On a alors :

$$\tau_2 \tau_3 \tau_i = (1 \ i \ 3 \ 2) \quad (7)$$

Mais en appliquant Φ :

$$\Phi(\tau_2) \Phi(\tau_3) \Phi(\tau_i) = (\alpha_1 \ \alpha_2)(\alpha_1 \ \alpha_3)(\alpha_2 \ \alpha_3) = (\alpha_1 \ \alpha_3) \quad (8)$$

Or Φ préserve l'ordre en tant qu'automorphisme, on a donc une contradiction. Ceci permet d'écrire :

$$\forall i \in \llbracket 2, n \rrbracket, \exists! \alpha_i \in \llbracket 1 \ n \rrbracket : \Phi(\tau_i) = (\alpha_1 \ \alpha_i) \quad (9)$$

On a ainsi construit une permutation $\alpha : i \mapsto \alpha_i$ et on a :

$$\forall i \in \llbracket 2, n \rrbracket, \alpha \tau_i \alpha^{-1} = \Phi(\tau_i) \quad (10)$$

c'est-à-dire que Φ et $\iota(\alpha)$ coïncident sur une partie génératrice de $\mathfrak{S}_n$: $\Phi = \iota(\alpha)$ et Φ est intérieur. $\square$

Grâce à ce lemme, il nous suffit de montrer que tout automorphisme de $\mathfrak{S}_n$ pour $n \neq 6$ envoie les transpositions sur les transpositions, ce qui constituera la seconde partie du développement.

Description des automorphismes de $\mathfrak{S}_n$

Soit τ une transposition et $\Phi \in \text{Aut}(\mathfrak{S}_n)$. Comme Φ préserve l'ordre des éléments de $\mathfrak{S}_n$, $\Phi(\tau)$ est d'ordre 2, et donc, lorsqu'on le décompose en un produit de cycles à support disjoints, c'est un produit de k transpositions^(vii). Grâce au lemme, on aura prouvé que Φ est intérieur si on montre que $k = 1$. Remarquons tout d'abord que k est impair. En effet, Φ envoie les commutateurs de $\mathfrak{S}_n$

(vi). On peut encore exclure le cas $n = 3$ car dans ce cas, τ_2 et τ_3 représentent toute la famille $(\tau_i)_i$ et il n'y a rien à montrer

(vii). On utilise ici le fait que l'ordre d'une permutation est le ppcm de la longueur des cycles dans sa décomposition en produit de cycles à support disjoints

sur les commutateurs, donc envoie le groupe dérivé $D(\mathfrak{S}_n)$ dans lui-même. Puisqu'on manipule des ensembles finis, on a donc $\Phi(D(\mathfrak{S}_n)) = D(\mathfrak{S}_n)$. Mais pour $n \geq 3$, $D(\mathfrak{S}_n) = \mathfrak{A}_n$, donc la signature de $\Phi(\tau)$ (qui est $(-1)^k$) est négative si $n \geq 3$ (et si $n \leq 2$, il va de soi que $k = 1$...). **Ainsi, cela permet de prouver d'emblée le théorème pour $n \leq 5$, et nous supposons dans toute la suite $n \geq 6$. On raisonne par contraposée, en supposant que $k \neq 1$.** Montrer le théorème revient donc à prouver que $n = 6$.

On va étudier les centralisateurs respectifs de τ et de $\Phi(\tau)$. Cela va nous donner des informations sur $\Phi(\tau)$, car Φ induit un isomorphisme $C(\tau) \cong C(\Phi(\tau))$ ^(viii).

Etude du centralisateur de τ

Notons $\tau = (a \ b)$ et $F := \llbracket 1, n \rrbracket \setminus \{a, b\}$. Alors, pour $\sigma \in \mathfrak{S}_n$:

$$\sigma \in C(\tau) \iff \sigma \tau \sigma^{-1} = (\sigma(a) \ \sigma(b)) = (a \ b) \quad (11)$$

$$\iff \sigma(\{a, b\}) = \{a, b\} \quad (12)$$

$$\iff \sigma(F) = F \quad (13)$$

Ainsi, on a un morphisme surjectif bien défini :

$$\begin{aligned} r : C(\tau) &\rightarrow \mathfrak{S}(F) \cong \mathfrak{S}_{n-2} \\ \sigma &\mapsto \sigma|_F \end{aligned}$$

Et son noyau est composé des permutations qui se comportent comme l'identité sur F , c'est-à-dire $\text{Ker}(r) = \{Id, \tau\}$. Il s'agit d'un groupe d'ordre 2 ; à isomorphismes près, on a la suite exacte :

$$1 \longrightarrow \mathbb{Z}/2\mathbb{Z} \hookrightarrow C(\tau) \twoheadrightarrow \mathfrak{S}_{n-2} \longrightarrow 1 \quad (14)$$

Etude du centralisateur de $\Phi(\tau)$

Notons, pour $i \in \llbracket 1, k \rrbracket$, $\tau_i = (a_i \ b_i)$ les k -transpositions apparaissant dans la décomposition en produit de cycles de $\Phi(\tau)$. On considère N le sous-groupe de $\mathfrak{S}_n$ engendré par les (τ_i)

Lemme 5. *N est un sous-groupe distingué de $C(\Phi(\tau))$, isomorphe à $(\mathbb{Z}/2\mathbb{Z})^k$*

Démonstration. Les τ_i sont à supports disjoints : ils commutent deux à deux et N est abélien. Comme ils sont tous d'ordre 2, tous les éléments de N sont d'ordre 1 ou 2, donc :

$$N \cong (\mathbb{Z}/2\mathbb{Z})^k \quad (15)$$

De plus, comme les (τ_i) commutent et sont d'ordre 2, ils sont tous dans $C(\Phi(\tau))$, et $N < C(\Phi(\tau))$. Enfin, soit $\sigma \in C(\Phi(\tau))$. On a alors :

$$(a_1 \ b_1) \dots (a_k \ b_k) = \sigma [(a_1 \ b_1) \dots (a_k \ b_k)] \sigma^{-1} = (\sigma(a_1) \ \sigma(b_1)) \dots (\sigma(a_k) \ \sigma(b_k)) \quad (16)$$

Or le terme de droite est encore une écriture de $\Phi(\tau)$ en produit de cycles à supports disjoints. Par unicité à l'ordre près des termes, l'automorphisme $\iota(\sigma)$ réalise une permutation des (τ_i) . En particulier, $\iota(\sigma)$ envoie N sur lui-même, c'est-à-dire que $N \triangleleft C(\Phi(\tau))$. $\square$

(viii). Même si ça ne me paraissait pas clair au début, il suffit de l'écrire pour s'en convaincre.

Conclusion de l'étude

Ainsi, par l'isomorphisme $C(\tau) \cong C(\Phi(\tau))$, on obtient que $C(\tau)$ contient un sous-groupe distingué de cardinal 2^k , noté H . Comme r est un morphisme surjectif, $r(H)$ est encore distingué dans $\mathfrak{S}_{n-2}$, et $r(H)$ est de cardinal 2^l , avec $l = k$ si $H \cap \text{Ker}(r) = \{Id\}$ ou $l = k-1$ si $\text{Ker}(r) \subset H$. Comme on a supposé $k \neq 1$, cela impose $l > 1$. Or, si n était supérieur strict à 6, on a $n-2 \geq 5$, et donc on connaîtrait les sous-groupes distingués de $\mathfrak{S}_{n-2}$. Ainsi :

$$r(H) \in \{\{Id\}, \mathfrak{A}_{n-2}, \mathfrak{S}_{n-2}\} \quad (17)$$

Or, chacun de ces cas est exclu pour des raisons de cardinal. La seule valeur possible pour n est donc 6.

Remarque : A défaut de savoir faire la preuve, il est bon d'au moins savoir que le cas $n = 6$ échappe à notre preuve pour une bonne raison : le théorème est faux dans ce cas là ! Perrin propose une preuve basée sur les isomorphisme exceptionnelles, considérations auxquelles je suis personnellement réfractaire, mais qui j'en suis sûr plaira grandement aux amoureux.ses de la théorie des groupes !. ♦

1.3 Autour de la cyclotomie ★★ / ★★★★★

Recasages possibles : 102, 120, 121, 123, 125, 127, 141

Ce développement est modulaire : il contient trop de choses pour être présenté tel quel. On peut prendre les parties les plus intéressantes pour la leçon en cours et admettre le reste, mais il s'apprend bien comme un seul bloc uni.

J'ai décomposé ce document en quatre parties, les trois dernières étant indépendantes :

1. En premier lieu, on montre plusieurs résultats préliminaires, essentiels pour manipuler les polynômes cyclotomiques, mais qui peuvent être simplement énoncés dans le plan pour gagner du temps et faire des choses plus poussées en développement.
2. On prouve ensuite l'irréductibilité des polynômes cyclotomiques dans $\mathbb{Q}[X]$. Indispensable pour la 102 et la 127 mais peut ne pas être mis en avant dans d'autres leçons, par exemple la 121, la 123 ou la 125.
3. On applique l'étude des polynômes cyclotomiques à la démonstration d'une version faible du théorème de progression arithmétique de Dirichlet. Un must have pour la 121 !
4. Dans la dernière partie, on étudie la réductibilité des polynômes cyclotomiques à coefficients dans les corps finis. Cette partie est certainement la plus délicate, mais elle est plus originale que le travail dans $\mathbb{Q}[X]$. Elle me semble parfaite pour la 123, la 125, la 141 et dans une moindre mesure, la 120.

Soit K un corps et n un entier. Dans toute la suite, on fixe L un corps de décomposition du polynôme $X^n - 1$. On pose $\mu_n(K)$ le groupe des racines n -ièmes de l'unité dans L et $\mu_n^*(K)$ l'ensemble (éventuellement vide) des racines n -ièmes primitives de l'unité, c'est-à-dire les éléments d'ordre n dans le groupe $\mu_n(K)$. On définit alors le n -ième polynôme cyclotomique à coefficients dans L :

$$\Phi_{n,K} := \prod_{\omega \in \mu_n^*(K)} (X - \omega) \in L[X] \quad (1)$$

Généralités sur les polynômes cyclotomiques ([Goz97], section VI.1)

Cette section contient un certain nombre de résultats préliminaires indispensables à connaître pour la suite mais qui peuvent être admis ou pas selon le temps et la leçon présentée.

Lemme 6 (Une relation de récurrence). *On suppose que la caractéristique de K ne divise pas n . Alors :*

$$X^n - 1 = \prod_{d|n} \Phi_{d,K} \quad (2)$$

Démonstration. Le polynôme $X^n - 1$ est séparable car il ne partage aucune de ses racines dans L avec son polynôme dérivé $nX^{n-1} \neq 0$ ^(ix). Il est donc à racines simples dans L . Par le théorème de Lagrange :

$$\mu_n(K) = \bigsqcup_{d|n} \mu_n^*(K) \quad (3)$$

(ix). Ce polynôme est non nul car n n'est pas divisé par la caractéristique de K .

et donc :

$$X^n - 1 = \prod_{\omega \in \mu_n(K)} (X - \omega) \quad (4)$$

$$= \prod_{d|n} \prod_{\omega \in \mu_d^*(K)} (X - \omega) \quad (5)$$

$$= \prod_{d|n} \Phi_{d,K} \quad (6)$$

□

Lemme 7.

$$\Phi_{n,\mathbb{Q}} \in \mathbb{Z}[X]$$

Démonstration. On procède par récurrence forte sur n . Si $n = 1$, alors $\Phi_{n,\mathbb{Q}} = X - 1 \in \mathbb{Z}[X]$. Supposons donc $n \geq 2$ et que tous les polynômes cyclotomiques d'ordre inférieur strict à n sont à coefficients entiers. Alors, d'après le lemme 6, $X^n - 1 = \Phi_{n,\mathbb{Q}} P$ où $P \in \mathbb{Z}[X]$. On peut voir cette écriture comme celle d'une division euclidienne dans $L[X]$. Mais alors, comme P est unitaire dans $\mathbb{Z}[X]$, on peut également faire la division euclidienne de $X^n - 1$ par P :

$$X^n - 1 = QP + R, \quad \deg(R) < \deg(P) \quad (7)$$

Comme il s'agit également d'une division euclidienne dans $L[X]$, on peut conclure par unicité : $R = 0$, $Q = \Phi_{n,\mathbb{Q}} \in \mathbb{Z}[X]$. □

Remarque : Par la même preuve, on prouve que $\Phi_{n,K} \in \mathbb{F}_p[X]$ si K est de caractéristique p première. ♦

Lemme 8. Si K est de caractéristique $p > 0$, et si n est premier avec p , alors :

$$\Phi_{n,K} = \overline{\Phi_{n,\mathbb{Q}}} \quad (8)$$

c'est-à-dire que $\Phi_{n,K}$ est obtenu en projetant les coefficients (entiers) de $\Phi_{n,\mathbb{Q}}$ dans $\mathbb{F}_p$.

Démonstration. On procède encore une fois par récurrence en utilisant le lemme 6. Pour $n = 1$,

c'est clair. Supposons $n \geq 2$ et que pour tout $d|n$, $d < n$ le résultat est acquis ^(x). Alors :

$$\overline{X^n - 1} = \Phi_{n,K} \prod_{\substack{d|n \\ d \neq n}} \Phi_{d,K} \quad (\text{lemme 6 dans } K) \quad (9)$$

$$= \Phi_{n,K} \prod_{\substack{d|n \\ d \neq n}} \overline{\Phi_{d,\mathbb{Q}}} \quad (\text{hypothèse de récurrence}) \quad (10)$$

$$= \overline{\Phi_{n,\mathbb{Q}} \prod_{\substack{d|n \\ d \neq n}} \Phi_{d,\mathbb{Q}}} \quad (\text{lemme 6 dans } \mathbb{Q}) \quad (11)$$

$$= \overline{\Phi_{n,\mathbb{Q}}} \prod_{\substack{d|n \\ d \neq n}} \overline{\Phi_{d,\mathbb{Q}}} \quad (12)$$

Il ne reste qu'à simplifier par $\prod_{\substack{d|n \\ d \neq n}} \overline{\Phi_{d,\mathbb{Q}}}$. □

Irréductibilité sur $\mathbb{Q}$ ([Goz97], théorème VI.11)

Nous allons montrer dans cette partie le très classique :

Théorème 9 (Irréductibilité dans $\mathbb{Q}[X]$). $\Phi_{n,\mathbb{Q}}$ est irréductible dans $\mathbb{Q}[X]$ et dans $\mathbb{Z}[X]$.

Soit f un diviseur irréductible unitaire de $\Phi_{n,\mathbb{Q}}$ dans $\mathbb{Q}[X]$. On va montrer que $f = \Phi_{n,\mathbb{Q}}$ en montrant que ces deux polynômes ont les mêmes racines complexes.

On a $f|X^n - 1$. Notons $h \in \mathbb{Q}[X]$ tel que $fh = X^n - 1$.

Lemme 10.

$$f, h \in \mathbb{Z}[X]$$

Démonstration. Comme $X^n - 1$ et f sont unitaires, h l'est également. Soit $r \in \mathbb{N}^*$ tel que $rf \in \mathbb{Z}[X]$. Notons alors r' le pgcd des coefficients de rf . Comme f est unitaire, $r'|r$ et donc $\frac{r}{r'}f \in \mathbb{Z}[X]$ est de contenu 1 avec $\frac{r'}{r} \in \mathbb{Z}$. Notons cet entier $\tilde{r}$. On fait de même avec h : il existe $\tilde{q}$ un entier tel que $\tilde{q}h$ soit à coefficients entiers et de contenu 1. On a alors :

$$\tilde{r}\tilde{q}(X^n - 1) = (\tilde{r}f)(\tilde{q}h) \quad (13)$$

En appliquant le contenu (qui est multiplicatif) à cette égalité, on obtient :

$$1 = \tilde{r}\tilde{q} \quad (14)$$

D'où $\tilde{r} = \tilde{q} = 1$ et $f \in \mathbb{Z}[X], h \in \mathbb{Z}[X]$. □

(x). Si n est premier avec p , c'est également le cas de tous ses diviseurs.

Ce résultat va nous permettre de projeter nos égalités dans $\mathbb{F}_p[X]$ par la suite.

Soit ω une racine complexe de f . Pour montrer $f = \Phi_{n,\mathbb{Q}}$, il suffit de montrer que si m est un entier premier avec n , alors ω^m est encore racine de f , car dans ce cas ces deux polynômes auront les mêmes racines. On peut encore réduire le problème en montrant que pour tout p un nombre premier ne divisant pas n , alors ω^p est racine de f . En le montrant pour toute racine ω , on aura bien que ω^m est encore racine de f si m est premier avec n . On a :

$$0 = (\omega^p)^n - 1 = f(\omega^p)h(\omega^p) \quad (15)$$

Par l'absurde, on suppose $f(\omega^p) \neq 0$. Alors $h(\omega^p) = 0$. Puisque f est irréductible et unitaire sur $\mathbb{Q}$, c'est le polynôme minimal de ω et donc on a :

$$f|h(X^p) \text{ i.e. } \exists g \in \mathbb{Q}[X] : h(X^p) = fg \quad (16)$$

Remarquons que cette égalité pouvant être vue comme une division euclidienne dans $\mathbb{Q}[X]$, on obtient immédiatement que $g \in \mathbb{Z}[X]$ en effectuant la même division euclidienne dans $\mathbb{Z}[X]$. On peut donc projeter l'égalité dans $\mathbb{F}_p[X]$:

$$\bar{f}\bar{g} = \overline{h(X^p)} = \overline{h(X)}^p \text{ (xi)} \quad (17)$$

Soit alors $\theta \in \mathbb{F}_p[X]$ un diviseur irréductible de $\bar{f}$. Alors :

$$\theta|\bar{h}^p \implies \theta|\bar{h} \quad (18)$$

Or on a :

$$\overline{X^n - 1} = \bar{f}\bar{h} \quad (19)$$

Donc :

$$\theta^2|\overline{X^n - 1} \quad (20)$$

ce qui est absurde car n étant premier avec p , $X^n - 1$ est séparable dans $\mathbb{F}_p[X]$ et ses facteurs irréductibles sont donc simples. En remontant le fil des arguments, on a bien montré que $\Phi_{n,\mathbb{Q}} = f$ et donc que les polynômes cyclotomiques sont tous irréductibles sur $\mathbb{Q}$!

Pour l'irréductibilité dans $\mathbb{Z}[X]$, il suffit d'utiliser le fait que ces polynômes sont unitaires, donc de contenu 1.

Corollaire 11. *$L/\mathbb{Q}$ est une extension de degré $\varphi(n)$, où φ désigne l'indicatrice d'Euler.*

Théorème de Dirichlet faible ([Goz97], proposition VII.13)

Les polynômes cyclotomiques vont nous permettre de démontrer une version faible du théorème de progression arithmétique de Dirichlet.

Théorème 12 (Dirichlet faible). *La suite $(\lambda n + 1)_{\lambda \in \mathbb{N}}$ contient une infinité de nombres premiers. En d'autres termes, il existe une infinité de nombres premiers congrus à 1 modulo n .*

(xi). On utilise le fait que le morphisme de Frobenius s'étend en un morphisme d'anneaux sur $\mathbb{F}_p[X]$ laissant stable les éléments de $\mathbb{F}_p$.

L'énoncé est impliqué par le suivant : pour tout entier $N > n$, il existe un nombre premier $p \geq N$ congru à 1 modulo n . La cyclotomie va fournir un candidat astucieux en exploitant le lemme suivant :

Lemme 13. *Soit $a \in \mathbb{N}$. Si p est un nombre premier divisant $\Phi_{n,\mathbb{Q}}(a)$ mais pas $\Phi_{d,\mathbb{Q}}(a)$, où d parcourt l'ensemble des diviseurs stricts de n , alors $p \equiv 1[n]$.*

Démonstration. Comme $p | \Phi_{n,\mathbb{Q}}(a)$, on a que $p | a^n - 1$. En réduisant modulo p , on obtient $a^n \equiv 1[p]$. Notons $o(a)$ l'ordre de a dans le groupe multiplicatif $\mathbb{F}_p^\times$. On vient de montrer que $o(a) | n$. Si $o(a) < n$, alors $\bar{a}^{o(a)} - 1 = 0$ dans $\mathbb{F}_p$ et donc :

$$p | a^{o(a)} - 1 = \prod_{d | o(a)} \Phi_{d,\mathbb{Q}}(a) \quad (21)$$

Il suit que p divise l'un des $\Phi_{d,\mathbb{Q}}(a)$, $d | o(a) | n$, ce qui est exclu par hypothèse. Donc $o(a) = n$. Par théorème de Lagrange sur le groupe multiplicatif $\mathbb{F}_p^\times$ (d'ordre $p - 1$), on a que $n | p - 1$, c'est-à-dire $p \equiv 1[n]$. $\square$

On propose alors p un diviseur premier de $\Phi_{n,\mathbb{Q}}(N!)$. On écrit $\Phi_{n,\mathbb{Q}}(N!) = pd$. Remarquons que $\Phi_{n,\mathbb{Q}}(0)$ est un entier, produit de racines de l'unité : c'est donc ± 1 . Ceci mène à une relation de Bézout :

$$pd = \Phi_{n,\mathbb{Q}}(N!) = N!P(N!) + \Phi_{n,\mathbb{Q}}(0) \quad (22)$$

où P est un polynôme à coefficients entiers. Ainsi, p est premier avec $N!$ et en particulier, $p > N > n$ (et notamment p est premier à n).

Si d est un diviseur strict de n tel que $p | \Phi_{d,\mathbb{Q}}(N!)$, alors en réduisant modulo p l'égalité :

$$(N!)^n - 1 = \prod_{d | n} \Phi_{d,\mathbb{Q}}(N!) \quad (23)$$

il vient que $\overline{N!}$ est une racine double de $\overline{X^n - 1}$, ce qui est absurde car n est premier avec p et donc $\overline{X^n - 1}$ est séparable. En appliquant le lemme, on a que $p \equiv 1[n]$, et on montre le théorème !

Facteurs irréductibles dans $\mathbb{F}_q[X]$ ([Esc97], exercice 14.7)

Dans cette section, on va étudier la réductibilité des polynômes cyclotomiques à coefficients dans un corps fini.

Soit q une puissance d'un nombre premier p . Soit n un entier premier avec p . On pose $s := [L : K]$.

Théorème 14. *Tous les facteurs irréductibles de $\Phi_{n,\mathbb{F}_q}$ sont d'ordre s . De plus, s est l'ordre de q dans le groupe multiplicatif $(\mathbb{Z}/n\mathbb{Z})^\times$.*

Démonstration. Soit f un facteur irréductible de $\Phi_{n,\mathbb{F}_q}$ et soit $\omega \in \mu_n^*(\mathbb{F}_q)$ une racine de f . Puisque ω engendre $\mu_n(\mathbb{F}_q)$, on a l'égalité ensembliste :

$$\mathbb{F}_q(\omega) = \mathbb{F}_q(\mu_n(\mathbb{F}_q)) = L \quad (24)$$

Comme par ailleurs $\mathbb{F}_q(\omega)$ est un corps de rupture de f , le degré de l'extension $[\mathbb{F}_q(\omega) : \mathbb{F}_q]$ est le degré de f . Par l'égalité précédente, c'est également s . Donc tous les facteurs irréductibles de $\Phi_{n, \mathbb{F}_q}$ ont même degré.

Pour montrer que q est d'ordre s dans $(\mathbb{Z}/n\mathbb{Z})^\times$, on s'intéresse à

$$\begin{aligned} \text{Frob} : L &\rightarrow L \\ x &\mapsto x^q \end{aligned}$$

On va montrer que Frob est d'ordre s dans $\text{Aut}_{\mathbb{F}_q}(L)$, le groupe des automorphismes de l'extension $L/\mathbb{F}_q$. Considérons α un générateur du groupe multiplicatif $L^\times$, d'ordre $q^s - 1$. Considérons alors $k \in \mathbb{N}$. On a :

$$\text{Frob}^{\circ k}(\alpha) = \alpha \iff \alpha^{q^k} = \alpha \quad (25)$$

$$\iff o(\alpha) \mid q^k - 1 \quad (26)$$

$$\iff q^s - 1 \mid q^k - 1 \quad (27)$$

$$(28)$$

D'où nécessairement, s est inférieur ou égal à l'ordre de Frob . Comme d'un autre côté, $\text{Frob}^{\circ s}$ est l'identité sur $\mathbb{F}_q$ par théorème de Lagrange, on obtient bien que Frob est d'ordre s . Pour terminer, prenons $\omega \in \mu_n^*(\mathbb{F}_q)$. Alors, pour $k \in \mathbb{N}$:

$$\text{Frob}^{\circ k}(\omega) = \omega \iff \text{Frob}^{\circ k} = \text{Id}_L \quad \text{car } L = \mathbb{F}_q(\omega) \quad (29)$$

$$\iff s \mid k \quad (30)$$

Mais d'autre part :

$$\text{Frob}(\omega)^{\circ k} = \omega \iff \omega^{q^s - 1} = 1 \quad (31)$$

$$\iff q^k - 1 \equiv 0[n] \text{ car } \omega \in \mu_n^*(\mathbb{F}_q) \quad (32)$$

$$\iff o(q) \mid k \quad (33)$$

Ces deux suites d'équivalences prouvent bien que l'ordre de q dans $(\mathbb{Z}/n\mathbb{Z})^\times$ est s . $\square$

Remarque : Il est bon de savoir qu'il existe une preuve plus directe pour arriver au résultat. La preuve que j'ai développé au-dessus reste tout à fait adaptée pour les leçons sur 123 et 125 car l'étude du morphisme de Frobenius est intéressante en elle-même, mais pour la 120 on pourrait procéder comme suit : un corps fini K contient une racine primitive n -ième de l'unité si, et seulement si, son groupe des inversibles $K^\times$ contient un sous-groupe d'ordre n . Ceci est équivalent au fait que n divise $|K^\times|$ car $K^\times$ est cyclique et contient donc exactement un sous-groupe cyclique d'ordre chacun de ses diviseurs. Finalement, K contient une racine n -ième primitive si, et seulement si :

$$|K| \equiv 1[n] \quad (34)$$

Dans notre cas, L est la plus petite extension (au sens du degré) qui contient une telle racine. Comme L est de cardinal q^s , s satisfait $q^s \equiv 1[n]$ et est minimal parmi les tels entiers : c'est exactement l'ordre de q dans $(\mathbb{Z}/n\mathbb{Z})^\times$! $\blacklozenge$

On a quelques conséquences intéressantes de cette étude :

Corollaire 15. Si n est tel que $(\mathbb{Z}/n\mathbb{Z})^\times$ n'est pas cyclique, alors $\Phi_{n,\mathbb{Q}}$ est réductible sur tous les corps finis de cardinal premier avec n .

Démonstration. Si q est une puissance d'un nombre premier ne divisant pas n , alors les facteurs irréductibles de $\overline{\Phi_{n,\mathbb{Q}}} = \Phi_{n,\mathbb{F}_q}$ sont tous de degré l'ordre de q dans $(\mathbb{Z}/n\mathbb{Z})^\times$. Donc si $(\mathbb{Z}/n\mathbb{Z})^\times$ n'est pas cyclique, q ne peut pas être d'ordre $\varphi(n)$ et les facteurs irréductibles de $\Phi_{n,\mathbb{F}_q}$ sont stricts. $\square$

Corollaire 16. $\Phi_{8,\mathbb{Q}}$ est un polynôme irréductible dans $\mathbb{Z}[X]$ et $\mathbb{Q}[X]$ mais réductible sur tous les corps finis.

Démonstration. On a $\Phi_{8,\mathbb{Q}} = X^4 + 1$. Soit $\mathbb{F}_q$ un corps fini de caractéristique p première.

Si $p = 2$, alors $\Phi_{8,\mathbb{F}_q} = (X + 1)^4$ est réductible.

Sinon, p est premier avec 8. Dans ce cas, il suffit de remarquer que :

$$(\mathbb{Z}/8\mathbb{Z})^\times \cong \mathbb{Z}/2\mathbb{Z} \times \mathbb{Z}/2\mathbb{Z} \quad (35)$$

En particulier, ce groupe n'est pas cyclique, et donc $\Phi_{8,\mathbb{F}_q}$ est réductible. $\square$

Enfin, on peut trouver une application pratique de ce théorème : il permet de construire à moindre coût des corps finis de caractéristique donnée qui contiennent des racines primitives, ce qui est utile pour résoudre des équations mais aussi intervient (pour ce que j'en sais) dans l'algorithme de la transformée de Fourier rapide sur les corps finis.

Application 17. $\mathbb{F}_{7^4}$ est la plus petite extension de $\mathbb{F}_{49}$ qui contient une racine primitive 10-ème de l'unité.

Démonstration. On calcule l'ordre de 49 dans $(\mathbb{Z}/10\mathbb{Z})^\times$. Commençons par calculer $\varphi(10)$:

$$\varphi(10) = \varphi(2)\varphi(5) = 4 \quad (36)$$

Comme $49 \not\equiv 1[10]$ et $49^2 \equiv 9^2 \equiv 81 \equiv 1[10]$, on trouve que 49 est d'ordre 2 dans $(\mathbb{Z}/10\mathbb{Z})^\times$. Donc toute extension de degré 2 de $\mathbb{F}_{49}$ est un corps de décomposition de $X^{10} - 1$ et contient un élément d'ordre 10.

L'avantage de cette méthode est qu'elle ne nécessite pas de calculer les facteurs irréductibles de $\Phi_{10,\mathbb{F}_{49}}$, ni même de calculer $\Phi_{10,\mathbb{F}_{49}}$ lui-même. Il suffit de trouver un irréductible de degré 4 sur $\mathbb{F}_7$ pour fabriquer l'extension qui nous intéresse ! $\square$

1.4 Classification des formes quadratiques sur les corps finis

★★

Recasages possibles : 123, 148, 170

Le groupe linéaire agit de plusieurs façons sur les ensembles de matrices, et savoir caractériser les orbites sous ces actions forme un ensemble de problèmes qu'il est intéressant de savoir résoudre. Dans le cas de l'action d'équivalence, le pivot de Gauss fournit une réponse immédiate. Pour l'action de similitude, c'est le théorème de réduction de Frobenius qui fournit les invariants. Mais pour l'action de congruence sur les matrices symétriques, il n'y a pas de réponse dans le cas général (du moins pas à ma connaissance).

On va répondre ici à ce problème dans le cas des matrices à coefficients dans un corps fini en caractérisant les classes d'isométrie des espaces quadratiques, problème qui est équivalent à déterminer les orbites de congruence des matrices symétriques. On aura en chemin besoin d'étudier le groupe des carrés sur un corps fini, en déployant pour cela des outils assez variés pour répondre à notre problème.

Soit p un nombre premier impair et q une puissance non triviale de p . On fixe $\mathbb{F}_q$ un corps fini à q -éléments ainsi que (E, φ) un $\mathbb{F}_q$ -espace quadratique de dimension d finie avec φ non nulle. Nous allons montrer :

Théorème 18 (Réduction des formes quadratiques, corps finis ([Rom], theo 16.19)). Soit $\alpha \in \mathbb{F}_q$ un élément non carré. Il existe un unique entier r et un unique $\delta \in \{1, \alpha\}$ tels que, dans une certaine base $\mathcal{B}$ de E , on ait :

$$\text{Mat}_{\mathcal{B}}(\varphi) = \begin{pmatrix} I_{r-1} & 0 & 0 \\ 0 & \delta & 0 \\ 0 & 0 & 0_{d-r} \end{pmatrix} \quad (1)$$

Cette forme réduite caractérise entièrement la classe d'isométrie de φ .

Groupe de carrés de $\mathbb{F}_q$

On va dans un premier temps donner quelques résultats sur les carrés de $\mathbb{F}_q$. On notera dans toute la suite :

$$K := \{x^2, x \in \mathbb{F}_q\} \quad K^\times = K \setminus \{0\} \quad (2)$$

Proposition 19. $K^\times$ est un sous-groupe de $\mathbb{F}_q^\times$ d'indice 2.

Démonstration. $K^\times$ est l'image de $\mathbb{F}_q^\times$ par le morphisme de groupes $f : x \mapsto x^2$. C'est donc un sous-groupe de $\mathbb{F}_q^\times$. Par ailleurs, le noyau du morphisme f est exactement l'ensemble des racines de $X^2 - 1$. Comme $\mathbb{F}_q$ est un corps de caractéristique impaire^(xii), ce polynôme a exactement deux racines données par 1 et -1 . Donc $K^\times$ est de cardinal $\frac{q-1}{2}$: il est bien d'indice 2. $\square$

(xii). En caractéristique 2, $1 = -1$ et ce morphisme est bijectif : c'est le Frobenius !

Lemme 20. Soient $(a, b) \in (\mathbb{F}_q^\times)^2$ et $c \in \mathbb{F}_q$. L'équations

$$ax^2 + by^2 = c \quad (3)$$

d'inconnues $(x, y) \in \mathbb{F}_q^2$ admet au moins une solution.

Démonstration. On commence par isoler x dans l'équation :

$$x^2 = \frac{c - by^2}{a} \quad (4)$$

Ainsi, l'ensemble des valeurs envisageables pour x^2 est donné par :

$$Z := \left\{ \frac{c - bz}{a} \mid z \in K \right\} \quad (5)$$

Pour que notre équation ait des solutions, il faut et il suffit que Z contienne un élément qui soit un carré. Nous allons montrer que $Z \cap K$ est non vide par dénombrement.

Tout d'abord, comme a et b sont inversibles dans $\mathbb{F}_q^\times$, l'application $z \mapsto a^{-1}(c - bz)$ est bijective. Z a donc même cardinal que K , c'est-à-dire $\frac{q+1}{2}$ (puisque $K^\times$ est d'indice 2 dans $\mathbb{F}_q^\times$ et que K contient exactement un élément supplémentaire). Mais la formule du crible donne :

$$|K \cup Z| = |Z| + |K| - |Z \cap K| = q + 1 - |Z \cap K| \quad (6)$$

comme $K \cup Z$ est un sous-ensemble de $\mathbb{F}_q$, son cardinal est majoré par q , d'où :

$$|K \cap Z| \geq 1 \quad (7)$$

Il suffit de prendre pour x une racine carrée d'un élément de $K \cap Z$: il existe alors $y \in \mathbb{F}_q$ tel que (x, y) est solution de l'équation par construction. $\square$

Réduction des formes quadratiques

Ces quelques résultats en poche, on va prouver le théorème de réduction. Commençons par remarquer que pour $\alpha \in \mathbb{F}_q \setminus K$ fixé, le couple $\{1, \alpha\}$ forme un système de représentants des éléments de $\mathbb{F}_q^\times$ modulo le groupe des carrés inversibles (puisque celui-ci est d'indice 2), observation qui aura son importance.

Commençons par prouver l'unicité de r et δ . La valeur de r ne fait pas mystère : il ne peut s'agir que du rang de φ . Supposons donnés $\mathcal{B}$ et $\mathcal{B}'$ deux bases de E et $\delta, \delta' \in \{1, \alpha\}^2$ tels que

$$Mat_{\mathcal{B}}(\varphi) = \begin{pmatrix} I_{r-1} & 0 & 0 \\ 0 & \delta & 0 \\ 0 & 0 & 0_{d-r} \end{pmatrix} \quad Mat_{\mathcal{B}'}(\varphi) = \begin{pmatrix} I_{r-1} & 0 & 0 \\ 0 & \delta' & 0 \\ 0 & 0 & 0_{d-r} \end{pmatrix} \quad (8)$$

Dans ce cas, δ et δ' sont deux valeurs du discriminant de φ restreint à un supplémentaire de son radical (ou noyau). Puisque celui-ci est uniquement déterminé par φ , $\delta = \delta'$.

Intéressons-nous maintenant à l'existence. On raisonne par récurrence sur d . Pour $d = 1$, le résultat provient immédiatement du fait que tout élément de $\mathbb{F}_q$ est congrus à 1 ou α modulo les carrés. Supposons $d \geq 2$ et le résultat acquis pour toute dimension inférieure. Soit $\mathcal{B} = (e_1, \dots, e_d)$ une base φ -orthogonale de E fournie par le théorème de réduction de Gauss. Quitte à permuter, on suppose que $\varphi(e_i) \neq 0$ pour $1 \leq i \leq r$.

Si $r = 1$, alors il existe $\delta \in \{1, \alpha\}$ et $x \in \mathbb{F}_q^\times$ tels que $x^2\varphi(e_1) = \delta$ par réduction modulo les carrés. Ainsi, en remplaçant e_1 par xe_1 , on obtient la forme souhaitée.

Sinon, le lemme précédent prouve qu'il existe $(x, y) \in \mathbb{F}_q^2$ tels que :

$$x^2\varphi(e_1) + y^2\varphi(e_2) = 1 \quad (9)$$

Posons $\varepsilon_1 := xe_1 + ye_2$. Par construction, on a $\varphi(\varepsilon_1) = 1$. Il suit que φ restreinte à $\text{Vect}(\varepsilon_1)$ est non dégénérée. Ainsi, en notant H l'orthogonal de $\text{Vect}(\varepsilon_1)$, H est un hyperplan de E , supplémentaire orthogonal de $\text{Vect}(\varepsilon_1)$. Complétons alors ε_1 en une base $\mathcal{C}$ de E en prenant des vecteurs $(\varepsilon_2, \dots, \varepsilon_d)$ une base φ -orthogonale de H dans laquelle $\varphi|_H$ a la forme souhaitée, obtenue par hypothèse de récurrence. Par construction, la base $\mathcal{C}$ est φ -orthogonale, donc la matrice de φ dans cette base est diagonale, et a la forme souhaitée. $\square$

Annexe - Et en pratique ?

Si ce théorème fournit, en théorie, un représentant canonique de la classe d'isométrie de φ , donnant ainsi une condition nécessaire et suffisante pour savoir si deux formes quadratiques sur $\mathbb{F}_q$ sont isométriques (ou, de façon équivalente, pour savoir si deux matrices symétriques à coefficients dans $\mathbb{F}_q$ sont congruentes), la mise sous forme réduite est en pratique un problème qui n'est pas aisé, car il nécessite de savoir réduire des éléments de $\mathbb{F}_q$ modulo les carrés. On va voir ici un algorithme pour répondre à ce problème dans la pratique.

Etant donnée φ une forme quadratique sur $\mathbb{F}_q$, on commence par calculer une base φ -orthogonale $\mathcal{B}$ grâce au théorème de réduction de Gauss (dont la preuve fournit un algorithme applicable pour la construction de telles bases). On obtient alors :

$$\text{Mat}_{\mathcal{B}}(\varphi) = \begin{pmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \ddots & 0 & \vdots \\ \vdots & 0 & \lambda_r & 0 \\ 0 & \cdots & 0 & 0 \end{pmatrix} \quad (10)$$

où les (λ_i) sont non nuls. L'entier r qui apparaît ici est le rang de φ : il s'agit du même r que celui qui apparaît dans la forme réduite. On calcule alors D le produit des λ_i , terme qui correspond au déterminant dans la base $\mathcal{B}$ tronquée de φ restreinte à un supplémentaire de son radical. D donc congrus à δ modulo les carrés. Reste à déterminer cette congruence...

Le problème se ramène en fait à savoir si D est ou non un carré de $\mathbb{F}_q$. Si $q = p$, ce problème se résout à l'aide de la loi de réciprocité quadratique (cf par exemple [Dem08]). Dans le cas général, on peut montrer qu'un élément $D \in \mathbb{F}_q^\times$ est un carré si, et seulement si, $D^{\frac{q-1}{2}}$, qui appartient à $\mathbb{F}_p$, est un carré, problème qu'on sait résoudre encore grâce à la loi de réciprocité quadratique, mais ceci est un développement à part entière^(xiii).

(xiii). Par ici !

1.5 Critère de Solovey-Strassen ★★☆☆

Recasages possibles : 120, 121, 127

Ce développement est une idée d'Alice Morinière, merci à elle pour cette proposition !

On va dans ce développement étudier un critère de primalité d'usage pratique qui ouvre la voie à des test probabilistes de primalité assez efficaces dans la pratique.

Le critère de Solovey-Strassen repose sur le symbole de Legendre et sa généralisation à $\mathbb{Z}/n\mathbb{Z}$, le symbole de Jacobi. Comme il n'est pas envisageable, en 15 minutes, de développer la théorie qui les entoure en plus du développement, il faut avoir de bonnes bases sur le sujet avant de se lancer dans le dev !

J'ai eu l'occasion de présenter ce développement en oral blanc et il a été bien apprécié : c'est une valeur sûre !

Etant donnés a et n deux entiers, avec n , on notera $a \wedge n$ leur p.g.c.d. positif. Si n est non nul, on notera $[a]_n$ le projeté de a dans $\mathbb{Z}/n\mathbb{Z}$. Si b est un second entier, on écrira :

$$a \equiv b[n] \quad (1)$$

lorsque $[a]_n = [b]_n$. Le groupe des inversible de $\mathbb{Z}/n\mathbb{Z}$ sera noté $(\mathbb{Z}/n\mathbb{Z})^\times$. On désignera par φ l'indicatrice d'Euler, dont la valeur en n est le cardinal de ce groupe.

Si p est un nombre premier impair et a un entier, on notera $\left(\frac{a}{p}\right)$ le symbole de Legendre de a modulo p , défini par :

$$\left(\frac{a}{p}\right) := \begin{cases} 0 & \text{si } p \mid a \\ 1 & \text{si } a \text{ est un carré modulo } p \\ -1 & \text{sinon} \end{cases} \quad (2)$$

Lorsque n est un entier impair positif admettant la décomposition en produit de nombres premiers $\prod_p p^{\nu_p(n)}$, on notera $\left(\frac{a}{n}\right)$ le symbole de Jacobi de a modulo n , défini par :

$$\left(\frac{a}{n}\right) := \prod_p \left(\frac{a}{p}\right)^{\nu_p(n)} \quad (3)$$

Nombres de Carmichael

Lorsque p est un nombre premier, l'anneau $\mathbb{Z}/p\mathbb{Z}$ est en fait un corps dont le groupe des inversibles est de cardinal $p-1$. Le théorème de Lagrange indique donc que pour tout a premier avec p , on a :

$$a^{p-1} \equiv 1[p] \quad (4)$$

Nous allons dans cette première partie montrer que la réciproque est fausse, et caractériser les cas pathologiques qui la mettent en défaut.

Définition 21 (Nombre de Carmichael). *Un entier $n > 2$ est dit de Carmichael s'il est composé (c'est-à-dire non premier) et si :*

$$\forall a \in \mathbb{N}, a \wedge n = 1 \implies a^{n-1} \equiv 1[n] \quad (5)$$

Théorème 22 (Caractérisation des nombres de Carmichael ([Dem08], prop. 3.33)). Soit n un entier composé. C'est un nombre de Carmichael si, et seulement si :

1. Les facteurs premiers de n sont simples.
2. Pour tout facteur premier p de n , $p - 1 \mid n - 1$.

Démonstration.

Commençons par le sens indirect, qui est plus élémentaire. On suppose donc que les facteurs premiers de n sont simples et que pour tout tel facteur premier p , $p - 1 \mid n - 1$. Fixons a un entier premier avec n (et donc en particulier premier avec tous les facteurs premiers de n). Si p est un facteur premier de n , a est donc inversible modulo p . Par théorème de Lagrange, et comme $p - 1 \mid n - 1$, il vient :

$$a^{p-1} \equiv a^{n-1} \equiv 1[p] \quad (6)$$

Ainsi, $a^{n-1} - 1$ est divisible par tous les facteurs premiers (simples) de n , donc est divisible par n . On a bien :

$$a^{n-1} \equiv 1[n] \quad (7)$$

Dans le sens direct, on suppose n de Carmichael. Par l'absurde, on suppose que n admet un facteur premier p multiple. Notons $n = p^2 m$ avec $m \in \mathbb{Z}$. On considère alors $a := pm + 1$. On a :

$$a^{n-1} = \sum_{k=0}^{n-1} \binom{n-1}{k} p^k m^k \equiv 1 + pm \equiv a[n] \quad (8)$$

Pourtant, on a une relation de Bézout évidente entre n et a : ces deux entiers sont premiers entre eux. Donc par hypothèse :

$$a^{n-1} \equiv 1[n] \quad (9)$$

Mais comme p est premier, $p \geq 2$ et donc :

$$1 < a < n \quad (10)$$

Ce qui entraîne une contradiction, car alors $[a]_n \neq [1]_n$. Les facteurs premiers de n sont donc simples. Fixons un tel facteur premier p . Le groupe $(\mathbb{Z}/p\mathbb{Z})^\times$ est cyclique (car c'est le groupe des inversibles d'un corps fini) d'ordre $p - 1$. Il contient donc un élément ω d'ordre $p - 1$. D'après le lemme des reste chinois, on a l'isomorphisme d'anneaux :

$$\begin{aligned} \Phi : \mathbb{Z}/n\mathbb{Z} &\rightarrow \mathbb{Z}/p\mathbb{Z} \times \mathbb{Z}/m\mathbb{Z} \\ [a]_n &\mapsto ([a]_p, [a]_m) \end{aligned} \quad (11)$$

où m est tel que $n = pm$. Ainsi, il existe un élément $a \in \mathbb{Z}$ tel que $[a]_p = \omega$ et $[a]_m = [1]_m$. Un tel élément a est nécessairement premier avec n , car inversible modulo p et m . On a donc par hypothèse :

$$a^{n-1} \equiv 1[n] \quad (12)$$

Et donc, en appliquant l'isomorphisme Φ :

$$\Phi([a^{n-1}]_n) = \Phi([1]_n) = (\omega^{n-1}, [1]_m^{n-1}) = ([1]_p, [1]_m) \quad (13)$$

Donc $\omega^{n-1} = [1]_p$. Comme ω est d'ordre $p - 1$, on a $p - 1 \mid n - 1$, et ceci vaut pour tous les facteurs premiers de n .

□

Critère de Solovey-Strassen

L'étude des nombres de Carmichael va nous permettre d'établir un critère pratique de primalité reposant sur les symboles de Jacobi.

Lorsque p est un nombre premier, les résultats généraux sur les résidus quadratiques indiquent :

$$\forall a \in \mathbb{Z}, a \wedge p = 1 \implies a^{\frac{p-1}{2}} \equiv \left(\frac{a}{p}\right) [p] \quad (14)$$

Nous allons ici établir la réciproque.

Théorème 23 (Solovey-Strassen, ([Dem08], prop 5.28)). *Soit n un entier impair. n est un nombre premier si, et seulement si :*

$$\forall a \in \mathbb{Z}, a \wedge n = 1 \implies \left(\frac{a}{n}\right) \equiv a^{\frac{n-1}{2}} [n] \quad (\star)$$

Démonstration. On s'intéresse uniquement au sens indirect, l'autre sens étant couvert par la théorie des résidus quadratiques dans les corps $\mathbb{F}_p$. Supposons donc que n est un entier satisfaisant $(\star)$. Par l'absurde, on suppose que n est composé. Montrons que c'est nécessairement un nombre de Carmichael. Si a est un entier premier avec n , on a alors :

$$a^{n-1} = \left(a^{\frac{n-1}{2}}\right)^2 \equiv \left(\frac{a}{n}\right)^2 \equiv 1 [n] \quad (15)$$

Ainsi, d'après la caractérisation des nombres de Carmichael, il existe $(p_1, \dots, p_r)$ des nombres premiers distincts tels que $n = \prod p_i$ et tels que pour tout $i \in \llbracket 1, r \rrbracket$, $p_i - 1 \mid n - 1$. Puisqu'on a supposé n composé, $r \geq 2$. On considère une famille $(x_i)_{1 \leq i \leq r}$ telle que :

$$\begin{cases} \forall i \in \llbracket 1, r \rrbracket, x_i \in (\mathbb{Z}/p_i\mathbb{Z})^\times \\ x_1 \text{ n'est pas un carré modulo } p_1 \\ \forall i \in \llbracket 2, r \rrbracket, x_i \text{ est un carré modulo } p_i \end{cases} \quad (16)$$

Mais d'après le lemme chinois, on a un isomorphisme d'anneaux :

$$\begin{aligned} \Phi : \mathbb{Z}/n\mathbb{Z} &\rightarrow \prod_{i=1}^r \mathbb{Z}/p_i\mathbb{Z} \\ [a]_n &\mapsto ([a]_{p_i})_{1 \leq i \leq r} \end{aligned} \quad (17)$$

Ainsi, en appliquant Φ^{-1} à la famille $(x_i)_{1 \leq i \leq r}$, on trouve un entier a premier avec n (car inversible modulo tous les (p_i)) tel que :

$$\forall i \in \llbracket 1, r \rrbracket, [a]_{p_i} = x_i \quad (18)$$

On trouve donc :

$$\left(\frac{a}{n}\right) = \prod_{i=1}^r \left(\frac{a}{p_i}\right) = -1 \quad (19)$$

Mais par réduction modulo p_r , en notant k tel que $k(p_r - 1) = n - 1$:

$$a^{\frac{n-1}{2}} \equiv \left(a^{\frac{p_r-1}{2}}\right)^k \equiv \left(\frac{a}{p_r}\right)^k \equiv 1[p_r] \quad (20)$$

Il reste ne qu'à exploiter le diagramme commutatif :

$$\begin{array}{ccc} \mathbb{Z} & \xrightarrow{\pi_{p_r}} & \mathbb{Z}/p_r\mathbb{Z} \\ \downarrow \pi_n & \nearrow & \\ \mathbb{Z}/n\mathbb{Z} & & \end{array} \quad (21)$$

et l'hypothèse $(\star)$: par projection de $\mathbb{Z}/n\mathbb{Z}$ dans $\mathbb{Z}/p_r\mathbb{Z}$, on obtient :

$$-1 \equiv \left(\frac{a}{n}\right) \equiv 1[p_r] \quad (22)$$

ce qui est absurde car p_r est impair, puisqu'on a supposé n impair. Donc n est premier. $\square$

On obtient immédiatement un corollaire pratique :

Corollaire 24. *Soit n une entier impair composé. l'ensemble des entiers a de $\llbracket 1, n-1 \rrbracket$ premiers avec n tels que :*

$$\left(\frac{a}{n}\right) \equiv a^{\frac{n-1}{2}} [n] \quad (23)$$

est de cardinal au plus $\varphi(n)/2$, où φ est l'indicatrice d'Euler.

Démonstration. Il suffit de remarquer que l'ensemble des tels éléments a , une fois projeté dans $\mathbb{Z}/n\mathbb{Z}$, est un sous-groupe strict de $(\mathbb{Z}/n\mathbb{Z})^\times$. Comme, d'après le théorème de Lagrange, son cardinal doit être un diviseur strict de $\varphi(n)$, il est majoré par $\varphi(n)/2$. Enfin, comme $\varphi(n) \leq n$, on trouve le résultat. $\square$

Annexe - Test probabiliste de Solovey-Strassen

Tout le travail précédemment effectué permet de créer un test probabiliste de primalité. Il n'y a pas le temps de le présenter pendant le développement, mais ce serait vraiment dommage de se priver de le mettre dans le plan... Vous trouverez beaucoup de remarques instructives sur l'algorithme sur le document d'Alice M.

Algorithme : Solovey-Strassen

Entrée : un entier $n > 1$ impair, un entier k

Sortie : Vrai ou Faux

Répéter k fois :

Choisir $a \in \llbracket 2, n \rrbracket$ avec probabilité uniforme

Si $a \wedge n \neq 1$:

Retourner Faux

Sinon, si $\left(\frac{a}{n}\right) \not\equiv a^{\frac{n-1}{2}} [n]$:

Retourner Faux
Retourner Vrai

Proposition 25. *Si l'algorithme renvoie Faux lorsque testé sur n , alors n est composé. Sinon, la probabilité pour que n soit composé est inférieure à $\frac{1}{2^k}$.*

Remarque :

- D'après [Dem08], le coût de calcul du symbole de Jacobi s'effectue en $O(\log_2(n)^2)$.
- Le coût de l'algorithme d'Euclide est en $O(\log_2(n)^3)$ pour le calcul du pgcd, et le calcul de $a^{\frac{n-1}{2}} \pmod n$ peut s'effectuer en utilisant l'algorithme d'exponentiation rapide avec une division euclidienne par n à chaque étape, pour un coup total de $n^2 \log_2(n)$.
- Le coût de l'algorithme est finalement en $O(k \cdot \log_2(n)^3)$, soit bien meilleur que les tests déterministes connus.
- Ce test ne permet pas de savoir avec certitude qu'un entier est premier. Toutefois, en prenant $k = 30$, la probabilité qu'un entier composé passe le test de Solovey-Strassen est inférieure à un sur un milliard.

◆

1.6 Décomposition de Dunford-Jordan-Chevalley - une méthode algorithmique ★★

Recasages possibles : 150, 152, 156

La théorème de décomposition de Dunford (a.k.a. de Jordan-Chevalley) est certainement l'un des théorèmes stars du programme d'algèbre de l'agreg. Relativement facile à démontrer, il brille surtout par la quantité d'applications pratiques qui en découlent.

On va voir ici un algorithme très efficace permettant de calculer concrètement la décomposition de Dunford, fournissant par là-même une des moyens d'application effectifs de cette décomposition.

Soient K un corps, E un K -espace vectoriel de dimension finie d et u un endomorphisme de E . On notera χ_u le polynôme caractéristique de u . On fera l'hypothèse que χ_u est scindé dans K , permettant d'écrire :

$$\chi_u = \prod_{i=1}^r (X - \lambda_i)^{\alpha_i} \quad (1)$$

où les (λ_i) sont les valeurs propres distinctes de u . Dans ces conditions, u admet une décomposition de Dunford qu'on note $u = \delta + \nu$, avec δ diagonalisable et ν nilpotent.

On introduit pour toute la suite du développement le polynôme $P \in K[X]$ défini par :

$$P := \prod_{i=1}^r (X - \lambda_i) \quad (2)$$

On va montrer le théorème suivant :

Théorème 26 (Dunford, version algorithmique ([Rom], section 20.6)). *Il existe une suite $(w_n) \in \mathcal{L}(E)^{\mathbb{N}}$ de premier terme u satisfaisant :*

1. *Pour tout $n \in \mathbb{N}$, $P(w_n)$ est nilpotent.*
2. *Pour tout $n \in \mathbb{N}$, $P'(w_n)$ est inversible.*
3. *$\forall n \in \mathbb{N}$, $w_{n+1} = w_n - P(w_n)[P'(w_n)]^{-1}$*

Par ailleurs, cette suite est stationnaire à partir du rang $\lceil \log_2(d) \rceil$, et sa « limite » est δ .

Remarque : De façon surprenant, l'intuition qui permet de trouver la suite (w_n) vient de l'analyse, alors qu'on travaille ici sur un corps quelconque où il n'est pas question d'utiliser des résultats d'ordre analytique. En effet, la partie diagonalisable de u est nécessairement annulée par le polynôme P , car elle a les mêmes valeurs propres que u et est annulée par un polynôme scindé à racines simples. On cherche donc moralement une « racine » de P dans $\mathcal{L}(E)$ qui soit « suffisamment proche de u ». Si on remplaçait ici $\mathcal{L}(E)$ par $\mathbb{R}$ ou $\mathbb{C}$, une technique usuelle pour trouver une telle racine consiste à appliquer la méthode de Newton à P partant de u . Or si l'on observe attentivement la formule de récurrence de la suite w_n , on retrouve précisément la formule donnée par la méthode de Newton, en ayant bien sûr remplacé la dérivée analytique de P par sa dérivée formelle. ♦

Démonstration. On commence par prouver par récurrence l'existence d'une telle suite. Pour cela,

on note $\mathcal{P}(n)$ l'assertion « il existe $w_0, \dots, w_n$ n -endomorphismes satisfaisant les propriétés du théorème ».

Pour $n = 0$, le théorème impose $w_0 = u$. Il s'agit donc de montrer que u satisfait (1) et (2).

Si α désigne le maximum des α_i (qui sont les multiplicités algébriques des valeurs propres de u), alors $\chi_u \mid P^\alpha$, et donc par le théorème de Cayley-Hamilton :

$$(P(u))^\alpha = P^\alpha(u) = 0 \quad (3)$$

donc $P(u)$ est nilpotent.

Comme P est scindé à racines simples sur K , il est premier avec son polynôme dérivé dans l'anneau $K[X]$. Celui-ci étant principal, on dispose du théorème de Bézout, qui fournit l'existence de polynômes A et B tels que :

$$AP + BP' = 1 \quad (B)$$

En évaluant cette égalité en u , il vient :

$$B(u)P'(u) = Id_E - A(u)P(u) \quad (4)$$

Comme $A(u)$ et $P(u)$ sont deux polynômes en u , ils commutent, et puisque $P(u)$ est nilpotent, $A(u)P(u)$ est nilpotent. Or dans un anneau commutatif, la somme d'un inversible et d'un nilpotent est toujours inversible. En appliquant ce résultat dans l'anneau $K[u]$, on trouve que $B(u)P'(u)$ est inversible, et donc $P'(u)$ est inversible. On a bien montré $\mathcal{P}(0)$.

Si l'on suppose le résultat vrai au rang n , on peut définir l'endomorphisme $w_{n+1} := w_n - P(w_n)[P'(w_n)]^{-1}$, puisque w_n satisfait (2). Pour montrer que w_{n+1} satisfait (1) et (2), on va encore s'inspirer de l'analyse et utiliser la formule de Taylor pour les polynômes :

$$P(Y) - P(X) = \sum_{k=1}^r \frac{P^{(k)}(X)}{k!} (Y - X)^k \in K[X, Y] \quad (5)$$

On évalue ce polynôme en $Y = w_{n+1}$ et $X = w_n$. Observons tout d'abord :

$$w_{n+1} - w_n = -P(w_n)[P'(w_n)]^{-1} \quad (6)$$

On obtient donc :

$$P(w_{n+1}) - P(w_n) = -P(w_n) + \sum_{k=2}^r (-1)^k \frac{P^{(k)}(w_n)}{k!} P^k(w_n)[P'(w_n)]^{-k} \quad (7)$$

D'où, en ajoutant $-P(w_n)$ de part et d'autre de l'égalité, on obtient la relation de divisibilité dans $K[u]$:

$$P(w_n)^2 \mid P(w_{n+1}) \quad (\text{xiv}) \quad (D)$$

En invoquant de nouveau les coefficients de Bézout A et B introduits en (B), on obtient encore :

$$B(w_{n+1})P'(w_{n+1}) = Id_E - A(w_{n+1})P(w_{n+1}) \quad (8)$$

et donc $P'(w_{n+1})$ est inversible comme somme d'un inversible et d'un nilpotent de $K[u]$.

(xiv). Cette relation reste vraie si $r = 1$: dans ce cas, $P(w_{n+1}) = 0$. En particulier, $P(w_{n+1})$ est nilpotent car divisible par un nilpotent d'après $\mathcal{P}(n)$.

Ainsi, par récurrence, la suite (w_n) est correctement définie et vérifie les propriétés souhaitées. Reste à montrer que la suite est ultimement stationnaire. Une récurrence immédiate sur la relation (D) montre :

$$\forall n \in \mathbb{N}, P(u)^{2^n} \mid P(w_n) \quad (9)$$

Comme $P(u)$ est nilpotent, $P(u)^d = 0$ (où d est la dimension de E). Donc pour tout $n \geq n_0 := \lceil \log_2(d) \rceil$, $P(u)^{2^n} = 0$ et $0 \mid P(w_n)$. En d'autres termes, $P(w_n) = 0$ pour n assez grand, et donc $w_{n+1} = w_n$. La suite est bien stationnaire à partir du rang n_0 .

Montrons enfin que $w_{n_0} = \delta$. Comme $P(w_{n_0}) = 0$, w_{n_0} est annulé par un polynôme scindé à racines simples et donc est diagonalisable. Or, par télescopage :

$$u - w_{n_0} = \sum_{k=0}^{n_0-1} w_k - w_{k+1} = - \sum_{k=0}^{n_0-1} \underbrace{P(w_k)[P'(w_k)]^{-1}}_{\text{nilpotent}} \quad (10)$$

C'est donc une somme de nilpotents dans l'anneau commutatif $K[u]$: c'est un nilpotent. Finalement, on a la décomposition $u = w_{n_0} + (u - w_{n_0})$ qui donne u comme somme d'un endomorphisme diagonalisable et d'un endomorphisme nilpotent commutant ensemble (puisque deux polynômes en u). Par unicité de la décomposition de Dunford, $w_{n_0} = \delta$. $\square$

Annexe - quid du polynôme P ?

J'ai vendu tout à l'heure la version algorithmique de Dunford comme permettant d'effectuer des calculs pratiques, mais il reste un problème à la réalisation de l'algorithme : c'est le calcul du polynôme P , dans l'expression duquel on voit apparaître explicitement les valeurs propres de u . Or bien évidemment, il est impossible de calculer explicitement les valeurs propres de u en toute généralité.

Proposition 27 ([BMP05], prop 5.39). *Si K est de caractéristique 0, P est donné par $\frac{\chi_u}{\chi_u \wedge \chi'_u}$.*

Démonstration. C'est un calcul assez direct : si $\chi_u = \prod_{i=1}^r (X - \lambda_i)^{\alpha_i}$, on trouve :

$$\chi'_u = \sum_{i=1}^r \alpha_i (X - \lambda_i)^{\alpha_i - 1} \prod_{j=1, j \neq i}^r (X - \lambda_j)^{\alpha_j} \quad (11)$$

Ainsi $(X - \lambda_i)^{\alpha_i - 1} \mid \chi'_u$ pour tout i et donc divise $\chi_u \wedge \chi'_u$. Par ailleurs, $(X - \lambda_i)^{\alpha_i}$ divise χ'_u si, et seulement si :

$$(X - \lambda_i) \mid \prod_{j=1, j \neq i}^r (X - \lambda_j) \quad (12)$$

ce qui est exclu en caractéristique 0 car les (λ_j) sont distincts (mais en caractéristique p , cela arrive dès que $p \mid \alpha_i$). On a donc :

$$\chi_u \wedge \chi'_u = \prod_{i=1}^r (X - \lambda_i)^{\alpha_i} \quad (13)$$

et la formule souhaitée s'ensuit. □

Je n'ai pas trouvé de façon de faire en toute généralité. Dans les corps finis, on a toujours la technique légèrement overkill d'utiliser l'algorithme de Berlekamp pour calculer les facteurs irréductibles de χ_u et en déduire le polynôme $P...$

1.7 Générateurs de $O(q)$ et $SO(q)$ ★★

Recasages possibles : 108, 148, 158, 161, 170, 171

Dans ce développement, on va trouver une partie génératrice assez simple du groupe des isométries d'une forme quadratique anisotrope. Il s'agit d'une version un peu plus générale que ce qui est rédigé dans [Per], où le corps de base n'a pas besoin d'être $\mathbb{R}$ – ce qui justifie à mon sens beaucoup plus le recasage dans la 170 – mais qui est pourtant un quasi copié-collé de la démonstration dans le cas euclidien.

Soit K un corps de caractéristique différente de 2. On considère (E, q) un espace quadratique de dimension finie d sur K . On notera $\tilde{q}$ la forme polaire de q . On s'intéresse au groupe des isométries $O(q)$, composé des isomorphismes linéaires qui laissent q invariante, et à $SO(q)$ composé des isométries de déterminant 1.

Etant donné F un sous-espace vectoriel de E^* , on notera F° son annulateur.

Définition 28 (Réflexion, renversement). Soit $u \in O(q)$. Notons E_1 l'espace propre de u associé à la valeur propre 1. On dit que u est une réflexion si E_1 est de codimension 1 (en d'autres termes, c'est un hyperplan). On dit que c'est un renversement si E_1 est de codimension 2.

Dans toute la suite, on suppose que q est anisotrope, c'est-à-dire :

$$\forall x \in E, q(x) = 0 \implies x = 0 \quad (1)$$

Générateurs de $O(q)$ ([Per], théorème VI.2.4)

On va montrer :

Théorème 29 (Générateurs de $O(q)$). Soit $u \in O(q)$. On note $p_u := d - \dim(E_1)$, où E_1 est l'espace propre associé à la valeur propre 1 de u . u est produit d'exactly p_u réflexions, et ne peut s'écrire comme un produit de strictement moins de p_u réflexions. En particulier, $O(q)$ est engendré par les réflexions.

On raisonne par récurrence sur p_u en commençant par montrer que u est produit d'au plus $p - U$ réflexions.

Si $p_u = 0$, alors $E_1 = E$ et donc u est l'identité. On conviendra alors que u , en tant qu'élément neutre de la composition, est produit de zéro réflexions.

Supposons que u n'est pas l'identité et que le résultat est acquis pour toute isométrie v telle que $p_v < p_u$. Soit $x \in E_1^\perp \setminus \{0\}$ et notons $y := u(x)$. Notons que $y \neq x$ car u étant supposée anisotrope, l'intersection d'un sous-espace et de son orthogonal est toujours réduite à 0. On remarque alors que $x - y$ et $x + y$ sont orthogonaux :

$$\tilde{q}(x + y, x - y) = \tilde{q}(x, x) - \tilde{q}(u(x), u(x)) = 0 \quad (2)$$

car u est une isométrie.

Remarquons à ce stade que, comme q est anisotrope, étant donné n'importe quel sous-espace vectoriel F de E , on a la décomposition en somme directe :

$$E = F \oplus F^\perp \quad (3)$$

En effet, on a toujours $\dim(F) + \dim(F^\perp) \geq \dim(E)$, et comme q est anisotrope, $F \cap F^\perp = \{0\}$. On peut donc considérer l'hyperplan^(xv) $H := \text{Vect}(x - y)^\perp$ et la décomposition orthogonale $E = H \oplus \text{Vect}(x - y)$. Posons alors τ la réflexion par rapport à H . On a :

$$\begin{cases} \tau(x) - \tau(y) = -x + y \\ \tau(x) + \tau(y) = x + y \end{cases} \iff \begin{cases} \tau(x) = y \\ \tau(y) = x \end{cases} \quad (4)$$

Ainsi, on a :

$$\tau u(x) = \tau(y) = x \quad (5)$$

et comme x a été choisi dans $E_1^\perp$, on a, pour tout $z \in E_1$:

$$\check{q}(x - y, z) = \check{q}(x, z) - \check{q}(u(x), u(z)) = -\check{q}(x, z) = 0 \quad (6)$$

donc $E_1 \subset H$: tous les vecteurs de E_1 sont stabilisés par τ . Il suit que $E_1 \oplus \text{Vect}(x)$ est inclus dans l'espace propre associé à 1 de τu et donc $p_{\tau u} < p_u$. Par récurrence, τu est produit d'au plus $p_{\tau u}$ réflexions, et comme une réflexion est l'inverse d'elle-même, u est produit d'au plus $p_{\tau u} + 1 \leq p_u$ réflexions.

Montrons maintenant que u ne peut pas être produit de strictement moins de τ_u réflexions. Pour cela, on montre le lemme suivant par un argument de dualité :

Lemme 30. Soit $r \in \mathbb{N}$, et soient $\tau_1, \dots, \tau_r$ réflexions d'hyperplans respectifs $H_1, \dots, H_r$. Alors :

$$\dim \left(\bigcap_{i=1}^r H_i \right) \geq d - r \quad (7)$$

Démonstration. Considérons $(\varphi_1, \dots, \varphi_r) \in E^{*r}$ des formes linéaires telles que :

$$\forall i \in \llbracket 1, r \rrbracket, H_i = \text{Ker}(\varphi_i) \quad (8)$$

On a alors :

$$\bigcap_{i=1}^r H_i = \bigcap_{i=1}^r \text{Vect}(\varphi_i)^\circ \quad (9)$$

$$= \left(\sum_{i=1}^r \text{Vect}(\varphi_i) \right)^\circ \quad (10)$$

$$= \text{Vect}(\varphi_1, \dots, \varphi_r)^\circ \quad (11)$$

d'où le résultat. $\square$

Il vient immédiatement que u n'est pas produit de strictement moins que p_u réflexions, puisque si $u = \tau_1, \dots, \tau_r$ avec les (τ_i) des réflexions, l'intersection de leurs hyperplans respectifs est contenue dans $E_1 \dots$

(xv). C'est bien un hyperplan car $x - y \neq 0$ et q est anisotrope.

Générateurs de $SO(q)$ ([Per], théorème VI.2.6)

Démontrons un résultat similaire pour $SO(q)$.

Théorème 31 (Générateurs de $SO(q)$). *Supposons $d \geq 3$. Soit $u \in SO(q)$. Alors u est produit d'au plus p_u renversements.
En particulier, $SO(q)$ est engendré par les renversements.*

Remarque : La condition $d \geq 3$ est plus que naturelle puisque les renversements n'existent tout simplement pas en dimension 0, 1 et 2. $\blacklozenge$

Puisque $SO(q) \subset O(q)$, u est produit d'exactly p_u réflexions :

$$u = \tau_1 \dots \tau_{p_u} \quad (12)$$

De plus, comme une réflexion est de déterminant -1 (qui est différent de 1 car K n'est pas de caractéristique 2!), p_u est pair.

Traitons séparément le cas $d = 3$. Dans ce cas, p_u est 0 (et u est l'identité, donc produit de zéro renversements) ou $p_u = 2$. Mais dans ce second cas, $-\tau_1$ et $-\tau_2$ sont des renversements (il suffit de considérer leurs matrices respectives dans une bonne base pour s'en convaincre), donc u est produit de deux renversements.

On ne suppose plus $d = 3$. Notons H_1 et H_2 les hyperplans stabilisés par τ_1 et τ_2 . Grâce au lemme 30, et comme $d \geq 3$, $H_1 \cap H_2$ contient un sous-espace vectoriel F de dimension $d - 3$. Comme F est stable par τ_1 et τ_2 et que ces deux réflexions sont des isométries, elles stabilisent également $F^\perp$ ^(xvi)... qui est de dimension 3! Donc il existe σ_1 et σ_2 des renversements de $F^\perp$ tels que $\sigma_1 \sigma_2 = \tau_1|_{F^\perp} \tau_2|_{F^\perp}$. En prolongeant ces renversements par l'identité sur E tout entier, on fabrique deux renversements de E dont le produit est $\tau_1 \tau_2$.

Comme on a montré que p_u est pair, il suffit de regrouper les p_u réflexions dont u est produit et d'exprimer chaque pair comme un produit de deux renversements pour obtenir le résultat.

Remarque : Cette fois-ci, p_u n'est plus le nombre minimal de renversements dont u est le produit. En effet, il suffit de prendre u un renversement : il est alors produit d'un seul renversement mais $p_u = 2$. $\blacklozenge$

(xvi). En effet, prenons $x \in F^\perp$ et $y \in F$. Alors $\tilde{q}(u(x), y) = \tilde{q}(u^{-1} \circ u(x), u^{-1}(y)) = \tilde{q}(x, u^{-1}(y)) = 0$ car u^{-1} stabilise également F , puisque u induit un endomorphisme injectif (donc bijectif) sur F .

1.8 Norme d'un élément algébrique ★★★ / ★★★★★

Recasages : 123, 125, 127, 144, 149

Ce développement est une idée de Matthias Hosten, un grand merci à lui pour cette chouette proposition :)

Ce document est très long, mais il ne faut évidemment pas tout faire. Voici un petit préambule qui explique la structure. J'ai décomposé le développement en trois parties et deux annexes de la façon suivante :

- 1. Dans une première partie, on va introduire la norme d'un élément d'une extension finie de corps, et en établir les propriétés les plus élémentaires. Cette partie est à traiter quelque soit le chemin que vous souhaitez emprunter car tout le reste en dépend.*
- 2. La seconde partie applique ces résultats pour caractériser le fait d'être un carré dans un corps fini de caractéristique impaire. Je n'ai pas de référence à fournir, mais elle n'est pas très difficile. Je la traite uniquement pour la leçon 123, mais c'est une affaire de goût, elle se fait très bien en 125 et 149 aussi.*
- 3. La troisième partie aborde la théorie des nombres : on va utiliser la norme pour caractériser les inversibles de l'anneau des entiers d'un corps de nombres. Cette partie demande plus de bagage théorique que la précédente, mais est particulièrement satisfaisante une fois maîtrisée. Elle est excellente pour la 144 (petite pépite d'utilisation des polynômes symétriques élémentaires) et la perturbante 127. Selon vos goûts, elle est également très bien pour la 125 et la 149. Attention tout de même, elle est nettement plus difficile que le reste (c'est elle qui vaut sa cinquième étoile au développement).*
- 4. Une première annexe où je montre un lemme que j'utilise sans démonstration dans la première partie : le polynôme minimal et le polynôme caractéristique d'un endomorphisme ont les mêmes facteurs irréductibles.*
- 5. Une deuxième annexe où je prouve de façon élémentaire que l'ensemble des entiers d'un corps de nombre est bel et bien un anneau, indispensable à savoir montrer pour la partie 3.*

J'ai essayé autant que possible de développer des arguments un peu différents de ceux de Matthias lorsque c'était possible, mais n'hésitez pas à aller voir son document pour choisir les techniques qui vous plaisent le plus !

Etant donnée une extension de corps L/K et $x \in L$ un élément algébrique sur K , on notera $\pi_{L/K,x}$ son polynôme minimal à coefficients dans K . Si f est un endomorphisme K -linéaire d'un K -espace vectoriel de dimension finie, on notera χ_f son polynôme caractéristique.

Pour un polynôme $P \in K[X]$, on note $\mathcal{Z}(P, L)$ l'ensemble de ses racines contenues dans L .

Norme d'un élément algébrique

Dans toute cette partie, on fixe L/K une extension **finie**. Notons $n := [L : K]$ son degré. J'utilise [Tau21], paragraphe 4.5 en référence, mais vous pouvez également trouver ces choses là dans [Goz97].

Définition 32 (Norme). Soit $\alpha \in L$. On définit l'application :

$$\begin{aligned} m_\alpha : L &\rightarrow L \\ x &\mapsto \alpha x \end{aligned}$$

C'est une application K -linéaire. On définit alors la norme de α par :

$$N_{L/K}(\alpha) := \det(m_\alpha) \in K \quad (1)$$

Avant d'aller plus loin, remarquons ceci :

Lemme 33. L'application :

$$\begin{aligned} L &\rightarrow \mathcal{L}_K(L) \\ \alpha &\mapsto m_\alpha \end{aligned}$$

est un morphisme injectif de K -algèbres.

Démonstration. Soit $\lambda \in K$, $(\alpha, \beta) \in L^2$. On a alors :

$$\forall x \in L, m_{\lambda\alpha + \beta}(x) = (\lambda\alpha + \beta)x = \lambda m_\alpha(x) + m_\beta(x) \quad (2)$$

Donc m est K -linéaire. De plus :

$$\forall x \in L, m_{\alpha\beta}(x) = \alpha\beta x = m_\alpha \circ m_\beta(x) \quad (3)$$

m est donc bien un morphisme de K -algèbres. Par ailleurs, si $\alpha \in \text{Ker}(m)$:

$$0 = m_\alpha(1) = \alpha \quad (4)$$

Donc ce morphisme est injectif. $\square$

On peut maintenant établir quelques propriétés de la norme :

Proposition 34 (Propriétés de la norme).

1. La norme est multiplicative, c'est-à-dire :

$$\forall (\alpha, \beta) \in L^2, N_{L/K}(\alpha\beta) = N_{L/K}(\alpha)N_{L/K}(\beta) \quad (5)$$

2. Soit $\alpha \in L$. Notons $d := \deg(\pi_{L/K, \alpha})$. On a alors :

$$N_{L/K}(\alpha) = (-1)^n [\pi_{L/K, \alpha}(0)]^{n/d} \quad (6)$$

3. Si M/L est une extension telle que $\pi_{L/K, \alpha}$ est scindé sur M , on a :

$$N_{L/K}(\alpha) = \left[\prod_{\omega \in \mathcal{Z}(\pi_{L/K, \alpha}, M)} \omega^{k_\omega} \right]^{n/d} \quad (7)$$

où k_ω est la multiplicité de la racine $\omega^{(\text{xvii})}$.

Démonstration.

1. C'est une conséquence directe du lemme précédent et de la multiplicativité du déterminant. Prenons $(\alpha, \beta) \in L^2$. On a alors :

$$N_{L/K}(\alpha\beta) = \det(m_{\alpha\beta}) = \det(m_\alpha \circ m_\beta) = \det(m_\alpha) \det(m_\beta) \quad (8)$$

d'où le résultat.

2. Commençons par remarquer que le polynôme minimal de m_α (en tant qu'application K -linéaire) est $\pi_{L/K, \alpha}$. En effet :

$$\forall P \in K[X], P(m_\alpha) = 0 \iff m_{P(\alpha)} = 0 \iff P(\alpha) = 0 \quad (9)$$

Donc l'annulateur de m_α est également celui de α , et nécessairement les polynômes minimaux coïncident.

On utilise désormais le fait que le polynôme caractéristique et le polynôme minimal d'un endomorphisme ont les mêmes facteurs irréductibles^(xviii). Dans ce cas précis, cela donne quelque chose de très fort, car $\pi_{L/K, \alpha}$ est irréductible dans $K[X]$. Donc χ_{m_α} est une puissance de $\pi_{L/K, \alpha}$, et par comparaison des degrés, on a immédiatement :

$$\chi_{m_\alpha} = \pi_{L/K, \alpha}^{n/d} \quad (10)$$

Or on utilise fréquemment en algèbre linéaire le fait que le terme constant du polynôme caractéristique est, au signe près, le déterminant. Donc :

$$N_{L/K}(\alpha) = (-1)^n \chi_{m_\alpha}(0) = (-1)^n \pi_{L/K, \alpha}(0)^{n/d} \quad (11)$$

3. C'est une conséquence directe de notre travail précédent, puisque le déterminant est le produit des racines du polynôme caractéristique, qui ici sont celles du polynôme minimal.

□

Carrés dans un corps fini

Je n'ai pas de référence pour cette partie...

Soit p un nombre premier, n un entier non nul et $q = p^n$. On note $\mathbb{F}_q$ le corps à q -éléments.

Proposition 35. Soit $x \in \mathbb{F}_q$. On a :

$$N_{\mathbb{F}_q/\mathbb{F}_p}(x) = x^{\frac{q-1}{p-1}} \quad (12)$$

Démonstration. Soit $\alpha \in \mathbb{F}_q^\times$ un générateur du groupe multiplicatif de $\mathbb{F}_q$. On va commencer par montrer ce résultat pour α , et on en déduira facilement le cas général par multiplicativité de la norme. On va pour cela exploiter le morphisme de Frobenius $\text{Frob} : x \mapsto x^p$.

(xvii). Je me permets une petite remarque à ce sujet parce que je me fais régulièrement avoir : un polynôme irréductible n'est pas nécessairement, en tout généralité, à racines simples dans une extension de décomposition (mais ça n'arrive qu'en caractéristique positive et dans le cas où le morphisme de Frobenius n'est pas surjectif sur K).

(xviii). Vous en doutez ? On se retrouve en annexe !

On a $\mathbb{F}_q = \mathbb{F}_p(\alpha)$, donc $\deg(\pi_{\mathbb{F}_q/\mathbb{F}_p, \alpha}) = [\mathbb{F}_q : \mathbb{F}_p] = n$. Par ailleurs, l'application :

$$\begin{aligned} \text{Aut} &\rightarrow \mathcal{Z}(\pi_{\mathbb{F}_q/\mathbb{F}_p, \alpha}, \mathbb{F}_q) \\ \Phi &\mapsto \Phi(\alpha) \end{aligned}$$

est injective car α est un élément primitif de l'extension. Comme α est, par définition, d'ordre $q^n - 1$, les éléments $(\alpha, \text{Frob}_p(\alpha), \dots, \text{Frob}_p^{\circ n-1}(\alpha))$ forment n racines distinctes de $\pi_{\mathbb{F}_q/\mathbb{F}_p, \alpha}$, qui est donc scindé à racines simples dans $\mathbb{F}_q$ et ses racines sont données par les itérées du Frobenius appliquées à $\alpha^{(\text{xix})}$. Ainsi, on peut appliquer le dernier point de la proposition 34 :

$$N_{\mathbb{F}_q/\mathbb{F}_p}(\alpha) = \left[\prod_{k=0}^{n-1} \text{Frob}_p^{\circ k}(\alpha) \right]^{n/n} \quad (13)$$

$$= \prod_{k=0}^{n-1} \alpha^{p^k} \quad (14)$$

$$= \alpha^{\sum_{k=0}^{n-1} p^k} \quad (15)$$

$$= \alpha^{\frac{p^n - 1}{p - 1}} \quad (16)$$

$$= \alpha^{\frac{q-1}{p-1}} \quad (17)$$

Prenons maintenant $x \in \mathbb{F}_q^\times$. Puisque α engendre le groupe multiplicatif de $\mathbb{F}_q$, il existe $k \in \mathbb{N}$ tel que $x = \alpha^k$. Donc :

$$N_{\mathbb{F}_q/\mathbb{F}_p}(x) = N_{\mathbb{F}_q/\mathbb{F}_p}(\alpha^k) = N_{\mathbb{F}_q/\mathbb{F}_p}(\alpha)^k = (\alpha^k)^{\frac{q-1}{p-1}} = x^{\frac{q-1}{p-1}} \quad (18)$$

Enfin, la formule est clairement vraie pour $x = 0$, ce qui achève notre démonstration. $\square$

On peut enfin démontrer le théorème portant sur les carrés de $\mathbb{F}_q$:

Théorème 36 (Caractérisation des carrés de $\mathbb{F}_q$). *Soit $x \in \mathbb{F}_q$. Les assertions suivantes sont équivalentes :*

1. x est un carré dans $\mathbb{F}_q$
2. $N_{\mathbb{F}_q/\mathbb{F}_p}(x)$ est un carré dans $\mathbb{F}_p$.

Démonstration. Traitons tout d'abord le cas de la caractéristique 2. Celui-ci est trivial car le passage au carré correspond alors au morphisme de Frobenius. C'est donc un morphisme de corps (donc injectif), et nécessairement, c'est une bijection. Donc tous les éléments de F_q (et de F_p) sont des carrés.

On suppose désormais que p est impair. Si $x = 0$, le résultat est évident. Supposons $x \neq 0$. On va utiliser la caractérisation suivante :

Lemme 37. *Soit K un corps fini de cardinal r impair. Les carrés de $K^\times$ sont exactement les racines de $X^{\frac{r-1}{2}}$.*

(xix). Remarquez qu'on a au passage prouvé que Frob_p est un générateur cyclique du groupe $\text{Aut}(\mathbb{F}_q/\mathbb{F}_p)$ et que ce groupe est d'ordre n , un résultat important en théorie de Galois !

Il suffit en effet de considérer le morphisme de groupes :

$$\begin{aligned} K^\times &\rightarrow K^\times \\ y &\mapsto y^2 \end{aligned}$$

Celui-ci est évidemment surjectif dans le groupe des carrés, et son noyau est $\{-1, 1\}$, donc de cardinal 2 car r est impair. Ainsi, par premier théorème d'isomorphisme, le groupe des carrés inversibles est isomorphe au quotient de $K^\times$ par ce noyau, et est donc de cardinal $\frac{r-1}{2}$.

De plus, si $y = z^2$ est un carré inversible, on a :

$$y^{\frac{r-1}{2}} = z^{r-1} = 1 \quad (19)$$

par théorème de Lagrange. Tous les carrés inversibles sont donc racines de $X^{\frac{r-1}{2}} - 1$ et par cardinalité, toutes les racines de ce polynôme sont des carrés inversibles.

Ce lemme en poche, le résultat tombe de suite. On a en effet l'égalité :

$$x^{\frac{q-1}{2}} = \left(x^{\frac{q-1}{p-1}}\right)^{\frac{p-1}{2}} \quad (20)$$

et donc il est clair que x est un carré dans $\mathbb{F}_q$ si, et seulement si, $N_{\mathbb{F}_q/\mathbb{F}_p}$ est un carré dans $\mathbb{F}_p$. $\square$

Remarque : On est en droit de se demander en quoi avoir rammené le problème à détecter les carrés de $\mathbb{F}_p$ nous avance. Il se trouve qu'on a une procédure algorithmique exploitant la loi de réciprocité quadratique pour détecter les carrés de $\mathbb{F}_p$. De plus, les éléments de $\mathbb{F}_p$ sont quand même moins sauvages que ceux de $\mathbb{F}_q$, car il s'agit de simples classes de congruence d'entiers, mais je laisse ces considérations aux options C! $\blacklozenge$

Remarque : Ce résultat trouve des applications concrètes. Par exemple, étant donnée une forme quadratique non dégénérée sur un $\mathbb{F}_q$ -espace vectoriel de dimension finie, sa classe d'isométrie est entièrement déterminée par la classe modulo les carrés de $\mathbb{F}_q$ de son déterminant pris dans n'importe quelle base (voir [Rom], théo 16.19) . Ainsi, en sachant réduire les éléments de $\mathbb{F}_q$ modulo les carrés, on sait déterminer si deux formes quadratiques sur un corps fini sont isométriques! $\blacklozenge$

Remarque : Attention à la formulation du théorème : le résultat à présenter est bien que x est un carré dans $\mathbb{F}_q$ si, et seulement si $N(x)$ est un carré dans $\mathbb{F}_p$, et non pas si $x^{\frac{q-1}{p-1}}$ est un carré dans $\mathbb{F}_p$, car ce dernier résultat est facile à montrer. En effet, $(x^{\frac{q-1}{p-1}})^{p-1} = x^{q-1} = 1$ lorsque x est inversible et donc $x^{\frac{q-1}{p-1}}$ est dans $\mathbb{F}_p$. L'équivalence sur les carrés coule alors de source sans avoir besoin de parler de norme... $\blacklozenge$

Caractérisation des inversibles dans l'anneau des entiers d'un corps de nombres

Là encore, une proposition de Matthias Hosten pour laquelle je n'ai pas de référence à fournir. Personnellement, c'est ma partie préférée du développement, mais elle est clairement plus difficile que celle sur les corps finis. C'était l'un des développements de la leçon que j'ai présenté à l'oral et j'ai eu plein de questions sur les entiers algébriques, il faut donc bien se préparer en

amont pour monter cette partie (en plus, ce n'est même pas le dev qu'à choisi le jury).

Commençons par quelques définitions pour fixer les idées. On se donne K un corps de nombres, c'est-à-dire une extension finie de $\mathbb{Q}$. Comme on ne travaillera qu'avec des extensions de $\mathbb{Q}$ dans cette partie, j'ai simplifié les notations : toutes les normes et tous les polynômes minimaux considérés sont sur $\mathbb{Q}$.

Définition 38 (Entiers sur K). Un nombre complexe $z \in K$ est dit entier sur K s'il existe un polynôme unitaire à coefficients dans $\mathbb{Z}$ qui annule z . On note $\mathfrak{O}_K$ l'ensemble des entiers de K . C'est un sous-anneau de $K^{(\text{xx})}$.

Lemme 39. Soit $z \in K$. Les assertions suivantes sont équivalentes :

1. $z \in \mathfrak{O}_K$
2. $\pi_z \in \mathbb{Z}[X]$

Démonstration. Le sens réciproque est immédiat, on se contentera donc du sens direct. Pour cela, on va utiliser le lemme de Gauss affirmant que le contenu d'un polynôme est multiplicatif^(xxi). Soit P un polynôme unitaire de $\mathbb{Z}[X]$ annulant z . Par définition du polynôme minimal :

$$\exists Q \in \mathbb{Q}[X] : P = \pi_z Q \quad (21)$$

On va chasser les dénominateurs : soient r et s deux entiers non nuls tels $r\pi_z \in \mathbb{Z}[X]$ et $sQ \in \mathbb{Z}[X]$. Comme P et π_z sont unitaires, Q l'est également. Ainsi, le contenu de $r\pi_z$ (resp. de sQ) divise r (resp. s). Quitte à diviser par le contenu, on peut supposer $r\pi_z$ et sQ primitifs (c'est-à-dire de contenu 1). Mais alors, par multiplicativité du contenu :

$$C(rsP) = C(r\pi_z)C(sQ) \quad (22)$$

Comme P est primitif (puisqu'il est unitaire) :

$$rs = 1 \quad (23)$$

r et s étant entiers, ils valent ± 1 et donc $\pi_z \in \mathbb{Z}[X]$. □

On en arrive au résultat qui nous intéresse :

Théorème 40 (Caractérisation des inversibles de $\mathfrak{O}_K$). Soit $z \in \mathfrak{O}_K$. Les assertions suivantes sont équivalentes :

1. $z \in \mathfrak{O}_K^\times$
2. $N(z) = \pm 1$

Démonstration. On note d le degré de π_z et $n := [K : \mathbb{Q}]$.

(xx). Ça n'est pas trivial ! Soyez sûr.e de savoir démontrer cela avant de vous embarquer dans cette preuve. J'ai ajouté une démonstration en annexe.

(xxi). Il est possible de faire une démonstration plus élémentaire, mais cette méthode se rencontre assez fréquemment pour résoudre d'autres problèmes, c'est pourquoi j'ai trouvé intéressant de présenter cette preuve-ci.

Dans le sens direct : on exploite la multiplicativité de la norme. On a en effet

$$N(zz^{-1}) = N(1) = 1 = N(z)N(z^{-1}) \quad (24)$$

Or la norme d'un entier algébrique est un entier, car d'après la proposition 34, c'est au signe près une puissance du terme constant du polynôme minimal, qui est à coefficients dans $\mathbb{Z}$ d'après le lemme précédent.

Dans le sens indirect : on considère $z_1, \dots, z_d$ les conjugués de z (c'est-à-dire les différentes racines de π_z dans $\mathbb{C}$)^(xxii). Quitte à permuter, on suppose $z = z_1$. En exploitant la proposition 34, on a :

$$N(z_1) = \left[\prod_{i=1}^d z_i \right]^{n/d} = z_1^{n/d} \left[\prod_{i=2}^d z_i \right]^{n/d} = \pm 1 \quad (25)$$

Ainsi, l'élément $\left[\prod_{i=2}^d z_i \right]^{n/d}$, est, au signe près, l'inverse de $z_1^{n/d}$, et c'est donc un élément de K . On considère le polynôme suivant :

$$P := \prod_{j=1}^d \left(X - \prod_{\substack{i=1 \\ i \neq j}}^d T_i^{n/d} \right) \in \mathbb{Z}[X][T_1, \dots, T_d] \quad (26)$$

Ce polynôme est symétrique en les (T_i) . Par théorème de structure des polynômes symétriques, il existe un polynôme $R \in \mathbb{Z}[X][S_1, \dots, S_d]$ tel que :

$$P(X) = R(\sigma_{d,1}(T_1, \dots, T_d), \dots, \sigma_{d,d}(T_1, \dots, T_d)) \quad (27)$$

où les $(\sigma_{d,k})$ sont les polynômes symétriques élémentaires en d indéterminées. En évaluant les (T_i) en les (z_i) , il vient :

$$P(z_1, \dots, z_d) = \prod_{j=1}^d \left(X - \prod_{\substack{i=1 \\ i \neq j}}^d z_i^{n/d} \right) \in \mathbb{Z}[X] \quad (28)$$

car les (z_i) étant les racines d'un polynôme de $\mathbb{Z}[X]$, les fonctions symétriques élémentaires appliquées en les (z_i) renvoient des entiers. Ce polynôme est unitaire et annule clairement

$\prod_{i=2}^d z_i^{n/d}$: cet élément est un entier de K et $z^{n/d} \in \mathfrak{O}_K^\times$. Cela implique bien que z est dans $\mathfrak{O}_K^\times$, puisque on a :

$$1 = z^{\frac{n}{d}} z^{-\frac{n}{d}} = z \times \left(z^{\frac{n}{d}-1} z^{-\frac{n}{d}} \right) \quad (29)$$

□

(xxii). Deux petites remarques qu'il est important d'avoir en tête :

1. Il y a bien d conjugués différents ! π_z est irréductible et on travaille en caractéristique 0, donc il est premier avec son polynôme dérivé.
2. Les conjugués de z n'ont aucune raison a priori d'être dans K .

Annexe I : facteurs irréductibles du polynôme caractéristique

On a utilisé en chemin un résultat plus ou moins classique, que je me propose de démontrer avec des outils très élémentaires. Notez que dans la preuve de Matoumatheux (ainsi que celles de Gozard et de Tauvel), on évite ce résultat d'une habile pirouette qui repose sur le théorème de la base télescopique. Une affaire de goût, sans doute.

Lemme 41. *Soit E un K -espace vectoriel de dimension finie. Soit $f \in \mathcal{L}(E)$. Les polynômes minimal π_f et caractéristique χ_f de f ont les mêmes facteurs irréductibles dans $K[X]$.*

Démonstration. Le théorème de Cayley-Hamilton implique que $\pi_f | \chi_f$, donc χ_f est divisé par tous les facteurs irréductibles de π_f . Réciproquement, notons :

$$\chi_f = \prod_{i=1}^r P_i^{\alpha_i} \quad (30)$$

où les (P_i) sont des polynômes deux à deux distincts, irréductibles dans $K[X]$. Soit L/K une extension de décomposition de χ_f . Pour $i \in \llbracket 1, r \rrbracket$, on considère $\lambda \in L$ une racine de P_i . λ est alors valeur propre de f car racine de $\det(XI_d - f)$. Ainsi, λ est racine de π . Ceci implique que le polynôme minimal de λ sur K divise π . Mais ce polynôme minimal est nécessairement P_i car celui-ci est irréductible sur K et annule λ ! $\square$

Annexe II - $\mathfrak{D}_k$ est un anneau

Il n'est pas tout à fait évident de montrer que l'ensemble des entiers sur un corps de nombres est un anneau. Dans la littérature, j'ai trouvé essentiellement deux preuves : l'une utilisant les résultants (chose que j'ai voulu absolument éviter car c'est lourd, pénible et pas au programme si vous ne faites option C ; elle a toutefois l'avantage de construire explicitement des polynômes annulateurs), l'autre utilisant des notions de théorie des modules (largement hors-programme donc, mais cela ressemble plus à ce que l'on fait pour prouver qu'une somme et un produit d'éléments algébriques est encore algébrique ; voir [DLQ14]). Un ami m'a suggéré une preuve beaucoup plus élémentaire, qui ne repose que sur les polynômes symétriques. Cependant, je n'ai pas trouvé de référence pour cette preuve. C'est celle que je vais présenter ici, dans un cadre très général, car c'est la plus simple et sûrement la moins connue.

Proposition 42 ($\mathfrak{D}_K$ est un anneau). *Soit A un anneau intègre de corps des fractions K et L/K une extension. On note $\mathfrak{D}$ l'ensemble des éléments de L qui sont annulés par un polynôme unitaire à coefficients dans A . C'est un anneau.*

Démonstration. Notons qu'il est clair que $\mathfrak{D}$ contient 0 et 1.

Soient x et y deux entiers sur A . Soient P et Q des polynômes unitaires de $A[X]$ annulant respectivement x et y . Quitte à plonger L dans des extensions de décomposition successives, on peut supposer que P et Q sont scindés dans L , de racines $(x_1, \dots, x_p)$ et $(y_1, \dots, y_q)$. On note $\mathcal{Z} := \{x_1, \dots, x_p\} \cup \{y_1, \dots, y_q\}$.

Stabilité par somme : On considère le polynôme :

$$S := \prod_{(a,b) \in \mathcal{Z}^2} (X - (a + b)) \quad (31)$$

Il s'agit d'un polynôme dont les coefficients sont symétriques en les éléments de $\mathcal{Z}$. Or $\mathcal{Z}$ est l'ensemble des racines du polynôme $PQ \in A[X]$. Une application du théorème de structure des polynômes symétriques prouve que S est à coefficients dans A . Comme il annule clairement $x + y$, on a $x + y \in \mathfrak{D}$.

Stabilité par produit : Le raisonnement est similaire. On pose :

$$S := \prod_{(a,b) \in \mathcal{Z}^2} (X - ab) \quad (32)$$

qui est encore un polynôme unitaire à coefficients dans A , et S annule xy .

□

Remarque : Cette preuve, outre qu'elle ne repose sur pas grand'chose une fois qu'on a bien compris la démarche, permet d'obtenir beaucoup de résultats classiques qu'il n'est pas toujours évident de démontrer, par exemple $\mathbb{Q}$ est un corps, l'ensemble des entiers algébriques de $\mathbb{C}$ est un anneau, etc. Avouez que c'est plus chouette que des vilains résultants non ? En plus ça se recase à merveille dans la 144 ! ♦

1.9 Par cinq points passe une conique ★★

Recasages possibles : 149, 162, 171, 181^(xxiii), 191

On va dans ce développement s'intéresser à un problème de géométrie : étant prescrits un certains nombres de points du plan, est-il possible de les interpoler par une conique ? Si oui, cette conique est-elle unique ? Si la nature géométrique de ce problème est facile à saisir, sa démonstration va requérir de déployer des outils purement algébriques. Il serait alors facile de s'enfoncer dans des calculs aveugles qui masqueraient la réalité géométrique du problème. Tout le jeu de cette preuve va être de jongler entre des points de vue géométriques et algébriques pour exploiter au mieux chacun des deux et s'épargner le plus de peine possible pour répondre au problème.

Il s'agit d'un développement très classique, pourtant je suis convaincu qu'il est amplement suffisant pour la leçon 171 étant donné que de nombreux.ses candidat.e.s font l'impasse sur les coniques. Il est par ailleurs très agréable à mener, et se recase à merveille dans plusieurs leçons difficiles à remplir du programme d'algèbre. Il a sauvé un certain nombre de mes plans, aussi ai-je une affection toute particulière pour celui-ci ! En plus, c'est le moment de travailler vos barycentres ;)

On considère A, B, C, D, E cinq points distincts du plan affine $\mathcal{P}$. Si quatre d'entre eux sont alignés, alors l'union de la droite passant par ces quatre points et de n'importe quelle droite passant par le cinquième est une conique passant par nos cinq points. On s'intéresse donc au cas où ces points sont quatre à quatre non alignés. Quitte à permuter, on suppose que A, B et C ne sont pas alignés. Ces trois points forment alors une base affine du plan ; toutes les coordonnées qu'on considérera dans la suite seront prises relativement à cette base. L'objectif est de démontrer :

Théorème 43 (Conique passant par cinq points ([Eid09], chap II, 2.2)). *Il existe une unique conique passant par cinq points quatre à quatre non alignés. De plus, cette conique est non dégénérée^(xxiv) si, et seulement si, ces points sont trois à trois non alignés.*

Pour cela, on va faire un usage abondant des coordonnées barycentriques, grâce auxquelles le problème va prendre une forme qu'on sait traiter avec des outils élémentaires d'algèbre linéaire. Dans une première partie, on commencera par montrer la forme de l'équation générale d'une conique en coordonnées barycentriques, puis on montrera le théorème en travaillant sur cette forme.

Equation barycentrique d'une conique

On va montrer :

(xxiii). L'intitulé de la leçon 181 n'inclut pas explicitement les barycentres, toutefois j'ai discuté avec l'un des professeurs qui encadraient la préparation à l'agreg et au vu du rapport, dont le premier point mentionne encore l'importance des barycentres dans la leçon, il pense que ce développement est toujours tout à fait d'actualité dans cette leçon.

(xxiv). Les conventions varient sur la définition de conique *non dégénérée*. Si $\varphi \in \mathbb{R}[U, V]$ est un polynôme de degré 2, on dit que l'ensemble de ses points d'annulation est une conique non dégénérée si φ ne se factorise pas en un produit de polynômes de degré 1. En termes géométriques, une conique dégénérée est la réunion de deux droites, éventuellement confondues.

Proposition 44 ([Eid09], chap II, 2.1). *Un ensemble de points de $\mathcal{P}$ est une conique passant par les points A , B et C si, et seulement si, il est décrit par une équation en coordonnées barycentriques de la forme :*

$$pXY + qXZ + rYZ = 0 \quad (1)$$

Partons de l'équation générale d'une conique **en coordonnées cartésiennes**. Notons :

$$\varphi(U, V) = aU^2 + bUV + cV^2 + dU + eV + f \quad (2)$$

On va alors utiliser la correspondance usuelle entre coordonnées cartésiennes et barycentriques :

$$(U, V) \mapsto (1 - U - V, U, V) \\ \left(\frac{Y}{X + Y + Z}, \frac{Z}{X + Y + Z} \right) \leftarrow (X, Y, Z)$$

Ainsi, un point $M \in \mathcal{P}$ de coordonnées barycentriques (X, Y, Z) appartient à la conique définie par φ si, et seulement si :

$$aY^2 + bYZ + cZ^2 + (X + Y + Z)(dY + eZ) + f(X + Y + Z)^2 = 0 \quad (\star)$$

Par ailleurs, cette conique contient A , B et C si, et seulement si, les triplets $(1, 0, 0)$, $(0, 1, 0)$ et $(0, 0, 1)$ sont solutions de cette équation. Cela donne immédiatement les conditions :

$$\begin{cases} f = 0 \\ a = -d \\ c = -e \end{cases} \quad (3)$$

En réécrivant l'équation $(\star)$ avec ces conditions, et après simplification, on trouve :

$$-aXY - cCZ + (b - a - c)YZ = 0 \quad (4)$$

et cette équation est de la forme voulue ! Réciproquement, on voit qu'une telle équation est une équation de conique satisfaite par les coordonnées de A , B et C .

Conique passant par cinq points

On cherche à montrer qu'il existe une et une seule conique passant par nos cinq points. En d'autres termes, on cherche une équation de la forme :

$$pXY + qXZ + rYZ = 0 \quad (\star\star)$$

satisfaite par les coordonnées barycentriques de D et E . Notons ces coordonnées (x_D, y_D, z_D) et (x_E, y_E, z_E) . Alors, la conique d'équation $(\star\star)$ passe par les cinq points si, et seulement si :

$$\begin{cases} px_D y_D + qx_D z_D + ry_D z_D = 0 \\ px_E y_E + qx_E z_E + ry_E z_E = 0 \end{cases} \quad (\star\star\star)$$

On interprète ce système comme un système linéaire dont les inconnues sont (p, q, r) . Puisqu'on cherche à montrer qu'il existe une seule conique circonscrite aux cinq points, il suffit de

montrer que ce système est de rang 2. Dans ce cas, l'ensemble des solutions est une droite vectorielle, et comme on ne change pas la conique géométrique en multipliant tous les coefficients par un même scalaire, tous les triplets non nuls (p, q, r) de cette droite donneront naissance à la même conique. Montrons donc que ce système est de rang 2. Pour cela, intéressons-nous à ses mineurs de taille 2 :

$$\begin{cases} \Delta_1 = x_D x_E (y_D z_E - y_E z_D) \\ \Delta_2 = y_D y_E (x_D z_E - x_E z_D) \\ \Delta_3 = z_D z_E (x_D y_E - x_E y_D) \end{cases}$$

Il suffit donc de montrer que ces trois quantités ne sont pas toutes nulles.

Maintenant que le contenu algébrique du problème a bien été saisi, tout le jeu de cette démonstration va être de savoir passer d'une observation algébrique à sa traduction géométrique et réciproquement. Pour cela, on considère le dessin suivant, qui nous servira souvent d'appui :

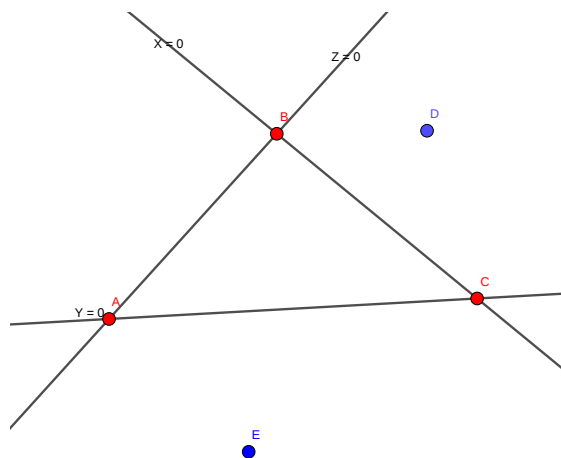


FIGURE 1.1 – Triangle fondamental

Remarquons que les trois droites noires représentent des lieux géométriques d'équation très simple. Nous reviendront souvent sur ce point par la suite. On va maintenant faire une petite étude de cas.

On suppose que $D, E \notin (BC)$. D'un point de vue algébrique, cette hypothèse se traduit par $x_D \neq 0, x_E \neq 0$, comme on le voit sur la figure 1.1. Si $\Delta_1 \neq 0$, il n'y a rien à démontrer. Sinon, observons :

$$\Delta_1 = 0 \iff \begin{vmatrix} 1 & 0 & 0 \\ x_D & y_D & z_D \\ x_E & y_E & z_E \end{vmatrix} = 0 \quad (5)$$

$$\iff D, E \text{ et } A \text{ sont alignés}^{(xxv)} \quad (6)$$

Ainsi, comme les cinq points sont quatre à quatre non alignés, si $\Delta_1 = 0$, D et E ne peuvent appartenir aux droites (AB) et (AC) , c'est-à-dire y_D, y_E, z_D et z_E sont non nuls.

(xxv). C'est un résultat général sur les coordonnées barycentriques, qu'il est indispensable de savoir démontrer pour faire ce développement. J'en ai mis une preuve en annexe.

On peut alors mener le même raisonnement sur Δ_2 et obtenir que celui-ci est nul si, et seulement si, B , D et E sont alignés, et A appartiendrait également à cet alignement, ce qui est absurde.

On suppose que $D \in (BC)$. Dans ce second cas, $x_D = 0$. Ainsi :

$$\begin{cases} \Delta_2 = 0 \\ \Delta_3 = 0 \end{cases} \iff \begin{cases} y_D y_E x_E z_D = 0 \\ z_D z_E x_E y_D = 0 \end{cases} \quad (7)$$

A présent, simplifions ces écritures à partir de considérations géométriques.

- Nos points étant supposés distincts, $x_D = 0$ entraîne $y_D \neq 0$ et $z_D \neq 0$, car sinon D serait confondu avec B ou C (voir figure 1.1).
- $x_E \neq 0$ car sinon $E \in (BC)$ et quatre de nos points seraient alignés.

Ainsi, le système se simplifie en :

$$\begin{cases} \Delta_2 = 0 \\ \Delta_3 = 0 \end{cases} \iff \begin{cases} y_E = 0 \\ z_E = 0 \end{cases} \iff E = A \quad (8)$$

Comme cette dernière assertion est fautive par hypothèse, nécessairement, $\Delta_2 \neq 0$ ou $\Delta_3 \neq 0$.

On a donc prouvé qu'au moins un des trois mineurs est non nul, et donc le système $(\star\star\star)$ est de rang 2, ce qui se traduit géométriquement par l'existence d'une unique conique interpolant nos cinq points !

Montrons finalement que cette conique est dégénérée si, et seulement si, trois points sont alignés.

Si trois points sont alignés, alors la réunion de la droite passant par ces trois points et de la droite passant par les deux points restants est une conique dégénérée interpolant les cinq points.

Si la conique est dégénérée, c'est une réunion de deux droites. Dans ce cas, le lemme des tiroirs indique qu'une de ces droites contient au moins trois de nos points, et donc trois d'entre eux sont alignés.

Annexe - Alignement en coordonnées barycentriques

Dans cette annexe, je propose une petite liste de résultats utilisant les coordonnées barycentriques qu'il est nécessaire de bien maîtriser pour aborder ce développement. Tout ceci est bien expliqué dans [Eid09].

Dans toute l'annexe, on considère une base affine du plan (A, B, C) dans lequel on exprime par X, Y, Z les coordonnées barycentriques.

Proposition 45 (Equation barycentrique de droite ([Eid09], chap I, 4.2)). *Un lieu géométrique est une droite affine si, et seulement si, il admet une équation en coordonnées barycentriques de la forme :*

$$aX + bY + cZ = 0 \quad (9)$$

avec a, b et c des réels distincts.

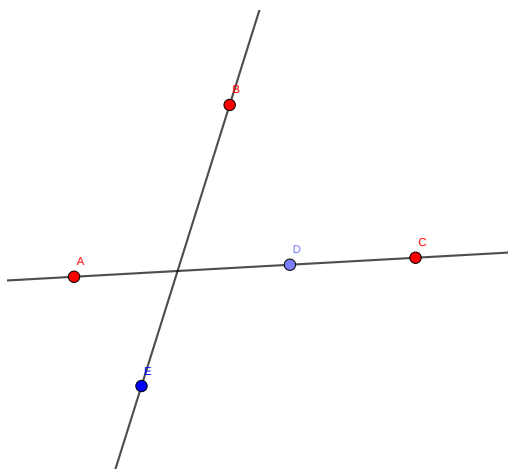


FIGURE 1.2 – Cas dégénéré

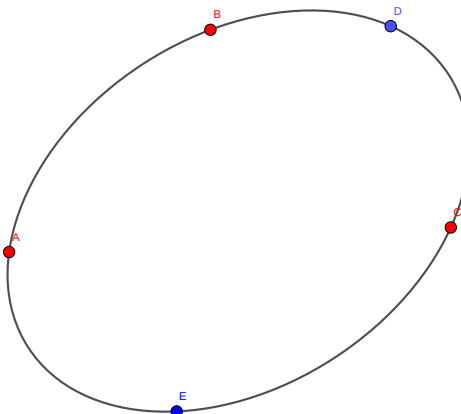


FIGURE 1.3 – Cas non dégénéré

Démonstration. On part de l'équation cartésienne d'une droite affine quelconque :

$$\alpha U + \beta V + \gamma = 0 \quad (10)$$

où $(\alpha, \beta) \neq (0, 0)$. En substituant les coordonnées barycentriques aux coordonnées cartésiennes, il vient :

$$\alpha Y + \beta Z + (X + Y + Z)\gamma = \gamma Z + (\alpha + \gamma)Y + (\beta + \gamma)Z = 0 \quad (11)$$

et comme α et β ne sont pas tous deux nuls, les trois coefficients qui apparaissent sont distincts.

Réciproquement, si on a une équation :

$$aX + bY + cZ = 0 \quad (12)$$

alors en passant aux coordonnées cartésiennes, il vient :

$$a(1 - U - V) + bU + cV = (b - a)U + (c - a)V + a = 0 \quad (13)$$

et $b - a$ et $c - a$ ne sont pas tous deux nuls sinon les trois coefficients a, b et c seraient égaux. Donc il s'agit là d'une équation cartésienne de droite affine. $\square$

Proposition 46 (Condition d'alignement ([Eid09], chap I, 5.1)). *Trois points (M_1, M_2, M_3) de coordonnées barycentriques $(x_1, y_1, z_1), (x_2, y_2, z_2)$ et (x_3, y_3, z_3) sont alignés si, et seulement si :*

$$\begin{vmatrix} x_1 & y_1 & z_1 \\ x_2 & y_2 & z_2 \\ x_3 & y_3 & z_3 \end{vmatrix} = 0 \quad (14)$$

Démonstration. Pour voir cela, on remplace la dernière ligne par les coordonnées génériques d'un point :

$$\begin{vmatrix} x_1 & y_1 & z_1 \\ x_2 & y_2 & z_2 \\ X & Y & Z \end{vmatrix} = \Delta_1 X + \Delta_2 Y + \Delta_3 Z \quad (15)$$

où les (Δ_i) sont, au signe près, les mineurs de taille 2 de la matrice. Supposons sans perte de généralité $M_1 \neq M_2$ (le résultat est trivial dans le cas contraire). Alors les (Δ_i) ne sont pas tous nuls, car sinon la matrice :

$$\begin{pmatrix} x_1 & y_1 & z_1 \\ x_2 & y_2 & z_2 \end{pmatrix} \quad (16)$$

serait de rang 1. En d'autres termes, ses deux lignes seraient proportionnelles, ce qui reviendrait à dire que $M_1 = M_2$. De plus, les (Δ_i) ne sont pas tous égaux. En effet, (x_1, y_1, z_1) est évidemment solution de l'équation :

$$\begin{vmatrix} x_1 & y_1 & z_1 \\ x_2 & y_2 & z_2 \\ X & Y & Z \end{vmatrix} = \Delta_1 X + \Delta_2 Y + \Delta_3 Z = 0 \quad (E)$$

et donc on aurait $\Delta_1(x_1 + y_1 + z_1) = 0$ avec $\Delta_1 \neq 0$. Mais ce cas est exclu car les coordonnées barycentriques d'un point ne peuvent sommer à 0. Ainsi, (E) est une équation de droite affine passant par M_1 et M_2 . Il en résulte immédiatement :

$$M_1, M_2, M_3 \text{ sont alignés} \iff M_3 \in (M_1 M_2) \iff \begin{vmatrix} x_1 & y_1 & z_1 \\ x_2 & y_2 & z_2 \\ x_3 & y_3 & z_3 \end{vmatrix} = 0 \quad (17)$$

□

1.10 Probabilités sur $GL_2(\mathbb{F}_q)$ ★★

Recasages possibles : 101, 103, 104, 106, 190

L'utilisation des probabilités uniformes permet de rendre compte de la proportion de couples commutant dans $GL(E)^2$, et donc d'avoir une idée quantitative d'à quel point ce groupe «n'est pas abélien». Ce genre de résultat est décrit plus généralement par un théorème dû à Dixon, qui affirme qu'en remplaçant $GL(E)$ par un groupe fini non-abélien quelconque, on peut toujours majorer cette probabilité par $5/8$.

Tout le développement est extrait de [Cal23], exercice 108, section XI-5, où seul le cas $q = 5$ est traité, mais c'est exactement pareil pour n'importe quel corps fini.

Dans toute la suite, on fixe q une puissance non-nulle d'un nombre premier et on note $\mathbb{F}_q$ le corps fini à q éléments. Posons $E := \mathbb{F}_q^2$.

Comme $GL(E)$ est un ensemble fini, on peut naturellement considérer des variables aléatoires suivant une loi uniforme sur $GL(E)$. Nous allons démontrer :

Théorème 47. Soient X et Y deux variables aléatoires indépendantes de loi $\mathcal{U}(GL(E))$. Alors :

$$\mathbb{P}(XY = YX) = \frac{1}{q^2 - q} \quad (1)$$

Un lemme de théorie des groupes

Notre démonstration va être dirigée par le lemme général suivant :

Lemme 48. On fait agir $GL(E)$ sur lui-même par conjugaison et on note k le nombre d'orbites obtenues. Alors :

$$\mathbb{P}(XY = YX) = \frac{k}{|GL(E)|} \quad (2)$$

Pour tout $g \in GL(E)$, on note $\text{Fix}(g) := \{x \in G \mid gxg^{-1} = x\}$ le fixateur de g sous l'action par conjugaison. On applique alors la formule des probabilités totales :

$$\mathbb{P}(XY = YX) = \mathbb{P}(Y = XYx^{-1}) \quad (3)$$

$$= \sum_{x \in GL(E)} \mathbb{P}(Y = xYx^{-1} \mid X = x) \mathbb{P}(X = x) \quad (4)$$

$$= \sum_{x \in GL(E)} \mathbb{P}(Y \in \text{Fix}(x)) \mathbb{P}(X = x) \quad (5)$$

$$= \frac{1}{|G|^2} \sum_{x \in G} |\text{Fix}(x)| \quad (6)$$

par indépendance de X et de Y , et car ces variables suivent la loi uniforme sur $GL(E)$. Il suffit donc d'appliquer la formule de Burnside pour obtenir le résultat souhaité.

Dénombrement des orbites

Grâce au lemme 48, il suffit pour calculer notre probabilité de calculer le nombre d'orbites de $GL(E)$ sous l'action de conjugaison. Pour cela, commençons par faire une observation.

On note :

$$\mathcal{D} = \{u \in GL(E) \mid u \text{ est diagonalisable}\}, \quad (7)$$

$$\mathcal{T} := \{u \in GL(E) \mid u \text{ est trigonalisable mais pas diagonalisable}\}, \quad (8)$$

$$\mathcal{N} := \{u \in GL(E) \mid u \text{ n'est pas trigonalisable}\}. \quad (9)$$

Alors, il est clair qu'on a une partition :

$$GL(E) = \mathcal{D} \sqcup \mathcal{T} \sqcup \mathcal{N} \quad (10)$$

et par ailleurs cette partition est stabilisée par l'action de conjugaison. On va donc compter séparément les orbites contenues dans $\mathcal{D}$, $\mathcal{T}$ et $\mathcal{N}$ respectivement.

Dans toute la suite, on se donne $u \in GL(E)$.

Etude des orbites de $\mathcal{D}$

On remarque :

$$u \in \mathcal{D} \iff \exists(\lambda, \mu) \in (\mathbb{F}_q^\times)^2 \text{ (xxvi)}, \exists \mathcal{B} \text{ une base de } E : \text{Mat}_{\mathcal{B}}(u) = \begin{pmatrix} \lambda & 0 \\ 0 & \mu \end{pmatrix} \quad (11)$$

Par ailleurs, un tel (λ, μ) est nécessairement décrit, à permutation près, par le spectre de u . Aussi, l'orbite de u est entièrement caractérisée par la classe de (λ, μ) modulo l'action naturelle de $\mathfrak{S}_2$. En d'autres termes, on a une bijection :

$$\begin{aligned} (\mathbb{F}_q^\times)^2 / \mathfrak{S}_2 &\rightarrow GL(E) \cdot \mathcal{D} \\ (\lambda, \mu) &\mapsto GL(E) \cdot \begin{pmatrix} \lambda & 0 \\ 0 & \mu \end{pmatrix} \end{aligned}$$

et donc le nombre d'orbites dans $\mathcal{D}$ est égal à :

$$|GL(E) \cdot \mathcal{D}| = \left| (\mathbb{F}_q^\times)^2 / \mathfrak{S}_2 \right| \quad (12)$$

$$= |\mathbb{F}_q^\times| + |\mathcal{P}_2(\mathbb{F}_q^\times)| \quad (13)$$

$$= q - 1 + \binom{q-1}{2} \quad (14)$$

$$= \frac{q(q-1)}{2} \quad (15)$$

où $\mathcal{P}_2(\mathbb{F}_q^\times)$ désigne l'ensemble des parties à deux éléments de $\mathbb{F}_q^\times$.

(xxvi). On prend les valeurs propres dans $\mathbb{F}_q^\times$ car notre endomorphisme est inversible.

Etude des orbites de $\mathcal{T}$

Supposons maintenant $u \in \mathcal{T}$. Comme par hypothèse u est trigonalisable, son spectre est non-vide, mais celui-ci est nécessairement réduit à un élément. En effet, dans le cas contraire, son polynôme caractéristique serait scindé à racines simples et donc u serait diagonalisable. Notons donc $\lambda \in \mathbb{F}_q^\times$ l'unique valeur propre de u . Par trigonalisation, il existe une base $\mathcal{B}$ de E telle que :

$$Mat_{\mathcal{B}} = \begin{pmatrix} \lambda & 1 \\ 0 & \lambda \end{pmatrix} \quad (\text{xxvii}) \quad (16)$$

Posons donc l'application :

$$\begin{aligned} \mathbb{F}_q^\times &\rightarrow GL(E) \cdot \mathcal{T} \\ \lambda &\mapsto GL(E) \cdot \begin{pmatrix} \lambda & 1 \\ 0 & \lambda \end{pmatrix} \end{aligned}$$

Cette application est surjective d'après le raisonnement précédemment mené. Par ailleurs, elle est nécessairement injective, puisque chaque antécédant d'une orbite est dans le spectre de tout endomorphisme de l'orbite, lequel est réduit à un élément. Donc :

$$|GL(E) \cdot \mathcal{T}| = |\mathbb{F}_q^\times| \quad (17)$$

$$= q - 1 \quad (18)$$

Etude des orbites de $\mathcal{N}$

On suppose enfin que $u \in \mathcal{N}$. Puisque u n'est pas trigonalisable, son polynôme caractéristique n'est pas scindé, et donc, celui-ci étant de degré 2, il n'a pas de racines dans $\mathbb{F}_q$. Montrons que dans ce cas, le polynôme caractéristique caractérise entièrement les orbites. Il est clair d'une part que c'est un invariant sous l'action par conjugaison. Réciproquement, on note $\chi_u = X^2 + aX + b$. Soit $e_1 \in E \setminus \{0\}$ et $e_2 := u(e_1)$. Comme u n'a pas de valeur propre, $\mathcal{B} := (e_1, e_2)$ est libre. De plus, par le théorème de Cayley-Hamilton :

$$u(e_2) = u^2(e_1) = -au(e_1) - be_1 = -be_1 - ae_2 \quad (19)$$

donc :

$$Mat_{\mathcal{B}}(u) = \begin{pmatrix} 0 & -b \\ 1 & -a \end{pmatrix} \quad (20)$$

Cette forme matricielle ne dépend que du polynôme caractéristique de u : elle détermine complètement l'orbite de u ! Ainsi :

$$|GL(E) \cdot \mathcal{N}| = |\{P \in \mathbb{F}_q[X] \mid \deg(P) = 2, P \text{ est unitaire et sans racine dans } \mathbb{F}_q\}| \quad (21)$$

$$= |\{X^2 + aX + b, (a, b) \in \mathbb{F}_q^2\} \setminus \{(X - \lambda)(X - \mu), (\lambda, \mu) \in (\mathbb{F}_q)^2 / \mathfrak{S}_2\}| \quad (22)$$

$$= q^2 - \frac{q(q+1)}{2} \quad (23)$$

$$= \frac{q^2 - q}{2} \quad (24)$$

(xxvii). Il suffit de prendre une base trigonalisante et de renormaliser le premier vecteur pour imposer le 1 en haut à droite.

Calcul de la probabilité

Ainsi, en combinant tous nos résultats, le nombre d'orbites de $GL(E)$ sous son action par conjugaison est donné par :

$$|GL(E) \cdot \mathcal{D}| + |GL(E) \cdot \mathcal{T}| + |GL(E) \cdot \mathcal{N}| = \frac{q(q-1)}{2} + q - 1 + \frac{q(q+1)}{2} = q^2 - 1 \quad (25)$$

Enfin, une application directe du lemme 48 donne le résultat :

$$\mathbb{P}(XY = YX) = \frac{q^2 - 1}{|GL_2(\mathbb{F}_q)|} = \frac{q^2 - 1}{(q^2 - q)(q^2 - 1)} = \frac{1}{q^2 - q} \quad (26)$$

Remarque : Pour mettre en lien avec le résultat du théorème de Dixon, notons que la fonction $x \mapsto x^2 - x$ est croissante pour $x \geq 1/2$. Ainsi, la probabilité que deux éléments commutent dans $GL_2(\mathbb{F}_q)$ est maximale pour $q = 2$ et vaut dans ce cas $1/2$, donc le groupe linéaire d'ordre 2 sur un corps fini n'atteint jamais l'optimum (notons que la valeur $5/8$ peut être atteinte, c'est le cas par exemple dans D_4 ou Q_8). Par ailleurs, $\mathbb{P}(XY = YX)$ tend vers 0 lorsque q tend vers $+\infty$.

Une question naturelle à se poser et de savoir si la preuve se généralise à $GL_n(\mathbb{F}_q)$. La réponse est non, car la description des matrices non-trigonalisables devient plus compliquée lorsque la dimension de E augmente. En particulier, le polynôme caractéristique ne caractérise plus l'orbite par conjugaison dès que $n \geq 3$ et n'est pas forcément sans racine. On pourrait éventuellement envisager la question du dénombrement en utilisant le théorème de réduction de Frobenius qui permettrait d'exhiber des invariants totaux, mais le problème deviendrait alors beaucoup plus compliqué et je n'ai pas creusé cette piste, je ne garantis donc pas qu'elle puisse aboutir. ♦

1.11 Réduction de Frobenius ★★

Recasages possibles : 148, 150, 151, 159

Dans ce développement, on va démontrer la classique *théorème de réduction de Frobenius*. Il s'agit d'un résultat qui a d'importantes conséquences théoriques et qui répond au problème de classier toutes les orbites de matrices sous l'action de similitude du groupe linéaire, sans aucune hypothèse sur le corps de base. On commencera par introduire la forme de Frobenius, les invariants de similitude dont on démontrera l'existence et l'unicité, et on terminera par une application intéressante qu'on peut faire de ce résultat.

Remarquez que le développement est **extrêmement** long. Je pense que vous pouvez traiter uniquement l'existence à l'oral ou alors l'unicité et l'application selon la leçon présentée et/ou vos affinités, mais que tout faire est déraisonnable.

Soient K un corps quelconque ^(xxviii) et n un entier naturel non-nul. Etant donné un polynôme $P \in K[X]$ de degré n unitaire, qu'on note sous forme développée :

$$P = X^n + \sum_{i=0}^{n-1} a_i X^i \quad (1)$$

on note sa matrice compagnon :

$$C_P := \begin{pmatrix} 0 & 0 & \dots & 0 & -a_0 \\ 1 & 0 & \dots & 0 & -a_1 \\ 0 & 1 & \dots & 0 & -a_2 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & \dots & 1 & -a_{n-1} \end{pmatrix} \quad (2)$$

Etant donné u un endomorphisme d'un K -espace vectoriel de dimension finie, on notera π_u son polynôme minimal. Si de plus x est un vecteur de l'espace sous-jacent, on note $\pi_{u,x}$ le polynôme minimal de u en x , c'est-à-dire l'unique générateur unitaire de l'idéal :

$$\{P \in K[X] \mid P(u)(x) = 0\} \quad (3)$$

Théorème 49 (Frobenius ([Goua], théorème B.2.1)). Soit E un K -espace vectoriel de dimension n et soit $u \in \mathcal{L}(E)$. Il existe une suite $(E_1, \dots, E_r)$ de sous-espaces vectoriels de E tels que :

1. $E = \bigoplus_{i=1}^r E_i$
2. u stabilise chacun des E_i et $u|_{E_i}$ induit un endomorphisme cyclique.
3. La suite $(P_1, \dots, P_r)$ des polynômes minimaux des endomorphismes induits sur les (E_i) est telle que $P_r \mid \dots \mid P_1$

De plus, la suite de polynômes (P_i) ainsi définis est unique : on l'appelle suite des invariants de similitudes de u .

(xxviii). Dans certains livres, comme le Rombaldi, le théorème de Frobenius suppose que le corps K est infini. Cette hypothèse intervient pour la preuve d'un lemme qui est fondamental pour le développement : pour $u \in \mathcal{L}(E)$, il existe un vecteur $x \in E$ tel que le polynôme minimal en x de u soit égal au polynôme minimal. La preuve est quasi immédiate si K est supposé infini, mais ce résultat reste vrai avec K quelconque. J'ai mis en annexe une preuve qui fonctionne dans le cas général.

Le théorème de Frobenius répond effectivement au problème de classification des orbites de similitude de $\mathbb{M}(K)$:

Corollaire 50. *Deux matrices de $\mathbb{M}(K)$ sont semblables si, et seulement si, elles ont les mêmes invariants de similitude.*

Existence de la forme de Frobenius

On raisonne par récurrence sur n .

Pour $n = 1$, le résultat est évident : on prend $r = 1$, $E_1 = E$ et on a $P_1 = X - u(1)$ en assimilant E à K .

Supposons le résultat acquis si la dimension de E est inférieure stricte à n . Dans ce cas, soit $x \in E \setminus \{0\}$ tel que $\pi_{u,x} = \pi_u$. On pose E_1 l'espace cyclique engendré par u et x , c'est-à-dire :

$$E_1 := \text{Vect}(u^k(x) \mid k \in \mathbb{N}) \quad (4)$$

Tout l'enjeu du développement est de trouver un supplémentaire de E_1 stable par u pour pouvoir y appliquer l'hypothèse de récurrence. Pour cela, on raisonne par dualité. Soit $(e_1, \dots, e_p)$ la base de E_1 formée des $e_i := u^{i-1}(x)$ avec p le degré du polynôme minimal en x de u . On la complète en une base de E . Notons (e_i^*) la base duale ainsi obtenue. On considère alors le sous-espace de E^* :

$$G := \text{Vect}((u^T)^k(e_p^*) \mid 0 \leq k \leq p-1) \quad (5)$$

G est stable par u^T . En effet, pour $0 \leq k < p-1$, on a :

$$u^T((u^T)^k(e_p^*)) = (e^T)^{k-1}(e_p^*) \in G \quad (6)$$

et pour $k = p-1$, on utilise le fait que $\pi_u(u) = 0$ et p est le degré de π_u , donc u^p est une combinaison linéaire des $(u^k)_{0 \leq k \leq p-1}$. Il en découle immédiatement que $u^T((u^T)^{p-1}(e_p^*)) \in G$.

Par ailleurs, G est de dimension p . En effet, soient $(\lambda_1, \dots, \lambda_p)$ des scalaires tels que :

$$\sum_{k=1}^p \lambda_k (u^T)^{k-1}(e_p^*) = 0 \quad (7)$$

Par l'absurde, on suppose que les (λ_i) ne sont pas tous nuls et on note r le plus grand indice tel que $\lambda_r \neq 0$. Alors :

$$0 = \left[\sum_{k=1}^r \lambda_k (u^T)^{k-1}(e_p^*) \right] (u^{p-r}(x)) \quad (8)$$

$$= e_p^* \underbrace{\lambda_r u^{p-1}(x)}_{=: e_p} + e_p^* \underbrace{\left(\sum_{k=1}^{r-1} \lambda_k (u^{k+p-r-1}(x)) \right)}_{=0 \text{ car } k+p-r-1 < p-1} \quad (9)$$

$$= \lambda_r \quad (10)$$

ce qui fournit une contradiction.

L'annulateur de G , qu'on note G° , est donc un sous-espace de E de dimension $n - p$, stable par u . Montrons que c'est un supplémentaire de E_1 . Par relation sur les dimensions, il suffit de montrer que leur intersection est réduite à $\{0\}$. Soit $z \in E_1 \cap G^\circ$. On décompose z sous la forme :

$$z = \sum_{k=1}^p \lambda_k u^{k-1}(x) \quad (11)$$

et on suppose par l'absurde que z est non nul. En prenant encore une fois r le plus grand indice tel que $\lambda_r \neq 0$ et en appliquant $(u^T)^{p-r-1}(e_p^*) \in G$ à z , on obtient :

$$0 = \lambda_r \text{ (xxix)} \quad (12)$$

Ainsi, on peut appliquer l'hypothèse de récurrence à l'endomorphisme induit sur G° : il existe une suite $(E_2, \dots, E_r)$ de sous-espaces vectoriels de G° telle que $G^\circ = \bigoplus_{i=2}^r E_i$ et telle que les polynôme minimaux des endomorphismes induits (qu'on appelle (P_i)) satisfont les bonnes relations de divisibilité.

Attention, la preuve n'est pas encore terminée ! En effet, il n'y a a priori aucune raison pour que $P_2|P_1$. Pour obtenir cette dernière relation, il faut remarquer que $P_1 = \pi_{u,x}$, qui par construction est le polynôme minimal de u . Donc $P_1(u|_{E_2}) = 0$ et puisque P_2 est le polynôme minimal de $u|_{E_2}$, on a bien $P_2|P_1$.

Unicité des facteurs invariants

Le raisonnement est classique : on se donne $(F_1, \dots, F_r)$ et $(G_1, \dots, G_s)$ deux suites de sous-espaces vectoriels comme dans le théorème, et on note $(P_1, \dots, P_r)$ et $(Q_1, \dots, Q_s)$ les polynômes minimaux des endomorphismes induits. On note de plus, pour alléger la suite, $(v_1, \dots, v_r)$ et $(w_1, \dots, w_s)$ les endomorphismes induits sur ces espaces. Par l'absurde, on suppose que les suites (P_i) et (Q_j) sont distinctes et on note i_0 l'indice minimal tel que $P_{i_0} \neq Q_{i_0}$ (xxx). Comme les (F_i) et les (G_j) sont stables par u , ils sont stables par tout polynôme en u . En particulier :

$$P_{i_0}(u)(E) = \bigoplus_{i=1}^r P_{i_0}(u)(E_i) = \bigoplus_{j=1}^s P_{i_0}(u)(F_j) \quad (13)$$

et comme, pour $i, j > i_0$, les polynômes minimaux des endomorphismes induits par u sur E_i et F_j divisent P_{i_0} , on a :

$$\forall i, j > i_0, \{0\} = P_{i_0}(E_i) = P_{i_0}(F_j) \quad (14)$$

et cette égalité reste vraie pour $i = i_0$. Ainsi :

$$P_{i_0}(u)(E) = \bigoplus_{i=1}^{i_0-1} P_{i_0}(u)(E_i) = P_{i_0}(G_{i_0}) \oplus \bigoplus_{j=1}^{i_0-1} P_{i_0}(u)(F_j) \quad (\star)$$

Utilisons maintenant le fait que les (v_i) et les (w_j) sont cycliques : ils sont donc représentés dans une certaine base par la matrice compagnon de leurs polynômes minimaux respectifs. En

(xxix). Moi aussi ça m'embête de faire deux fois exactement la même preuve mais je n'ai malheureusement pas trouvé de moyen simple pour expliquer pourquoi cette deuxième preuve découle immédiatement de notre travail précédent.

(xxx). qui existe même si $r \neq s$ car $\sum_{i=1}^r \deg(P_i) = \sum_{j=1}^s \deg(Q_j) = n$ donc si $r < s$ et les (P_i) sont les r -premiers (Q_j) , par égalité sur les degrés, les (Q_j) restants sont des constantes non nulles, ce qui est absurde car ce sont des polynômes minimaux.

particulier, pour $i < i_0$, v_i et w_i sont représentés par une même matrice ; de même donc pour $P_{i_0}(v_i)$ et $P_{i_0}(w_i)$. Ceci implique que ces endomorphismes ont même rang, et donc on a :

$$\forall i < i_0, \dim(P_{i_0}(u)(F_i)) = \dim(P_{i_0}(u)(G_i)) \quad (15)$$

On passe alors à la dimension dans $(\star)$ et on montre :

$$\dim(P_{i_0}(u)(G_{i_0})) = \dots = \dim(P_{i_0}(u)(G_s)) = 0 \quad (16)$$

et donc P_{i_0} annule w_{i_0} , ce qui implique qu'il est divisé par son polynôme minimal Q_{i_0} .

Le raisonnement étant symétrique en les (P_i) et les (Q_j) , on obtient de même que $P_{i_0} | Q_{i_0}$ et donc ces polynômes sont égaux, ce qui prouve l'unicité des invariants de similitude !

Réponse au problème de classification

On va maintenant démontrer le corollaire, bien que je pense que cette partie puisse être passée à l'oral si vous regardez bien le jury dans les yeux, parce qu'elle se résume franchement à un jeu d'écriture.

Soient A et B dans $\mathbb{M}(K)$. Le théorème de Frobenius affirme qu'il existe des bases $\mathcal{B}_1$ et $\mathcal{B}_2$ de K^n dans lesquelles les matrices des endomorphismes canoniquement associés à A et B sont :

$$\begin{pmatrix} C_{P_1} & 0 & 0 & \dots & 0 \\ 0 & C_{P_2} & 0 & \dots & 0 \\ \vdots & \ddots & \ddots & \ddots & \vdots \\ 0 & 0 & 0 & \dots & C_{P_r} \end{pmatrix} \quad \text{et} \quad \begin{pmatrix} C_{Q_1} & 0 & 0 & \dots & 0 \\ 0 & C_{Q_2} & 0 & \dots & 0 \\ \vdots & \ddots & \ddots & \ddots & \vdots \\ 0 & 0 & 0 & \dots & C_{Q_s} \end{pmatrix} \quad (17)$$

où les (P_i) et (Q_j) sont les invariants de similitude respectifs de A et de B . Il est donc clair que si A et B sont semblables si, et seulement si, elles ont les mêmes invariants.

Une application

Une conséquence théorique intéressante de ce résultat est la suivante :

Proposition 51. *Soit L/K une extension de corps et soit A et B deux matrices carrées de taille n à coefficients dans K . Les assertions suivantes sont équivalentes :*

1. *A et B sont K -semblables.*
2. *A et B sont L -semblables.*

Démonstration. Supposons qu'il existe P une matrice carrée inversible à coefficients dans L telle que $P^{-1}AP = B$. Ceci implique que A et B ont les mêmes invariants de similitude en tant que matrices de $\mathbb{M}(L)$. Mais les invariants de similitude à coefficients dans K satisfont les hypothèses du théorème de Frobenius lorsqu'on plonge A et B dans $\mathbb{M}(L)$: par unicité, les invariants de similitude de A et B vues comme éléments de $\mathbb{M}(K)$ sont identiques, et donc A et B sont semblables dans $\mathbb{M}(K)$. $\square$

Annexe

Comme je l'indiquais plus haut, certains livres supposent le corps K infini pour faciliter la preuve d'un lemme qu'on a utilisé plus haut. Afin que la preuve soit complète, j'ajoute une démonstration de ce lemme dans le cas général, car je pense qu'il est important de la connaître si on veut traiter le théorème de Frobenius dans toute sa généralité.

Lemme 52 ([Goua], chapitre IV, paragraphe 2, exercice 3). Soit $u \in \mathcal{L}(E)$. Il existe $x \in E$ tel que $\pi_{u,x} = \pi_u$.

Démonstration. Etudions d'abord le cas où π_u est une puissance d'un polynôme irréductible dans $K[X]$, qu'on note P . On a alors $\pi_u = P^\alpha$. Comme on a toujours, pour tout $x \in E$, $\pi_{u,x} | \pi_u$, il existe un entier $\beta(x)$ tel que $\pi_{u,x} = P^{\beta(x)}$. Il est alors assez immédiat de remarquer :

$$\alpha = \max_{x \in E} \beta(x) \quad (18)$$

où le maximum est justifié car β est à valeurs dans une partie finie de $\mathbb{N}$. Donc il existe x tel que $\beta(x) = \alpha$, c'est-à-dire tel que $\pi_{u,x} = \pi_u$.

Dans le cas général, décomposons π_u en produit d'irréductibles :

$$\pi_u = \prod_{i=1}^r P_i^{\alpha_i} \quad (19)$$

On utilise alors le lemme des noyaux et le fait que $\pi_u(u) = 0$:

$$E = \bigoplus_{i=1}^r \text{Ker}(P_i^{\alpha_i}(u)) \quad (20)$$

Pour $i \in \llbracket 1, r \rrbracket$, notons u_i l'endomorphisme induit sur $\text{Ker}(P_i^{\alpha_i}(u))$. On montre facilement que le polynôme minimal de u_i est $P_i^{\alpha_i}$, donc on se ramène au cas précédent. Il existe donc $x_i \in \text{Ker}(P_i^{\alpha_i}(u))$ tel que $\pi_{u_i, x_i} = \pi_{u_i}$.

On pose alors $x = x_1 + \dots + x_r$. On a alors, pour $Q \in K[X]$:

$$Q(u)(x) = 0 \iff \sum_{i=1}^r Q(u_i)(x_i) = 0 \quad (21)$$

$$\iff \forall i \in \llbracket 1, r \rrbracket, Q(u_i)(x_i) = 0 \text{ (xxxi)} \quad (22)$$

$$\iff \forall i \in \llbracket 1, r \rrbracket, P_i^{\alpha_i} | Q \quad (23)$$

$$\iff \pi_u | Q \quad (24)$$

ce qui prouve que $\pi_{u,x} = \pi_u$. □

Et en pratique ?

Le forme normale de Frobenius, comme les autres formes normales, est une forme réduite de matrices dont on espère pouvoir tirer des informations. Hélas, comme vous l'aurez sans doute

(xxxi). Car $Q(u)$ stabilise la somme directe en espaces caractéristiques.

remarqué, elle ne coïncide pas avec les autres formes agréables (par exemple, si A est diagonalisable, sa réduite de Frobenius n'est en général pas diagonale). Toutefois, elle a un avantage immense par rapport à celles-ci : elle est calculable !

En fait, ce développement est intimement lié à un autre résultat : le théorème de réduction de Smith dont voici l'énoncé :

Théorème 53 (Smith^(xxxii)). *Soit $M \in \mathbb{M}_{n,p}(\mathbb{A})$. Il existe un unique entier r et une suite $(d_1, \dots, d_r)$ d'éléments de $\mathbb{A}$, uniques modulo l'association, tels que $d_r | \dots | d_1$ et M soit équivalente à la matrice :*

$$\left(\begin{array}{ccc|c} d_1 & & & 0 \\ & \ddots & & \\ & & d_r & \\ \hline & 0 & & 0 \end{array} \right) \quad (25)$$

Par ailleurs, la preuve du théorème de Smith fournit immédiatement un algorithme pour calculer cette forme réduite, qui repose dans le cas d'un anneau euclidien sur l'algorithme d'Euclide de calcul du pgcd. Si on applique cet algorithme à la matrice $XI_n - C_P \in \mathbb{M}(K[X])$, on trouve une réduite de Smith dont tous les coefficients diagonaux sont égaux à 1, à l'exception d'un qui est égal à P . Etendant ce résultat à toutes les matrices, qui sont compagnons par blocs d'après le théorème de Frobenius, on trouve que les éléments différents de 1 dans la réduite de Smith de $XI_n - A$ sont exactement les facteurs invariants de la forme de Frobenius.

Cette « coïncidence » est loin d'être anodine et repose en réalité sur une théorie profonde qui est celle des modules de type fini sur les anneaux principaux, dont le théorème de Smith est l'un des fondements. Peut-être aurez-vous remarqué l'étonnante similitude entre le théorème de Frobenius et le théorème de structure des groupes abéliens de types finis (« il existe une unique suite $d_R | \dots | d_1$ etc. ») ? Ces deux théorèmes sont en fait deux instances d'un théorème beaucoup plus général qui découlent de la forme normale de Smith ! Si ce point de vue vous intéresse, Matthias Hostein a écrit un poly fort instructif sur le sujet : [Hos24]. On pourra également trouver ces considérations dans des livres plus ou moins classiques de l'agrégation, comme [BMP05] ou encore [DLQ14].

(xxxii). Pssst... j'ai aussi tapé un poly sur celui-ci !

1.12 Réduction de Smith - anneaux principaux ★★★★★

Quoi qu'on en dise, l'algorithme du pivot de Gauss est un outil d'algèbre fondamental, tant du point de vue de la pratique que pour l'immense variété de ses conséquences théoriques. Hélas, celui-ci repose fortement sur la possibilité de diviser par des scalaires, c'est-à-dire sur le fait qu'on travaille sur un corps. Pourtant, de nombreux champs d'algèbre exploitent de façon cruciale les modules, qui sont l'équivalent des espaces vectoriels où l'on a substitué un anneau au corps de base, et dans ce contexte, le pivot de Gauss échoue à produire quelque résultat que ce soit...

On va, dans ce développement, étudier un algorithme d'élimination des coefficients pour les systèmes linéaires à coefficients dans des anneaux principaux, et appliquer cet algorithme à l'élaboration d'une forme normale réduite pour les matrices à coefficients dans notre anneau. Cette forme, dite de Smith, est l'analogue dans les anneaux principaux du représentant canonique des orbites d'équivalence des matrices à coefficients dans un corps.

D'une façon surprenante peut-être (en tout cas qui moi me surprend) la forme réduite de Smith a des conséquences théoriques absolument dantesques, parmi lesquelles on peut citer le théorème de structure des modules de type fini sur un anneau principal, dont le théorème de structure des groupes abéliens de type fini et le théorème de Frobenius sont deux corollaires célèbres.

La présentation qu'on va en faire ici est très tournée algorithmique (on va même réaliser des preuves formelles de terminaison d'algorithmes!) et se place dans le cas principal, que je trouve plus intéressant que le cas euclidien, même si celui-ci est certainement plus intuitif.

Remarque : On pourrait présenter ce développement dans le cadre des anneaux euclidiens et se simplifier la tâche, d'autant que les « vraies » applications se font toujours dans le cadre des anneaux euclidiens. J'ai préféré ce cadre abstrait pour deux raisons : la première est qu'il met plus en valeur les propriétés profondes de l'anneaux qui sont à l'œuvre (et personnellement j'aime beaucoup la théorie des anneaux). La deuxième est qu'on atteint ainsi un cadre minimal pour faire fonctionner l'algorithme, ce sur quoi je reviendrai plus en détails dans l'annexe. Soyez simplement prévenu.e.s que cette version du développement nécessite de rentrer dans des considérations sur les anneaux noetheriens, qui sont hors-programmes. ♦

Dans toute la suite, on fixe $\mathbb{A}$ un anneau principal. Etant donnée une matrice $M \in \mathbb{M}(\mathbb{A})$, $i \in \llbracket 1, n \rrbracket$ et $j \in \llbracket 1, p \rrbracket$, on notera $M(i, j)$ le coefficient de M situé en i -ème ligne et j -ième colonne^(xxxiii). Etant donné $S \subset \mathbb{A}$, on notera $\langle S \rangle$ l'idéal de $\mathbb{A}$ engendré par S . Notre objectif est de démontrer :

Théorème 54 (Smith). *Soit $M \in \mathbb{M}_{n,p}(\mathbb{A})$. Il existe un unique entier r et une suite $(a_1, \dots, a_r)$ d'éléments de $\mathbb{A}$, uniques modulo l'association, tels que $a_r | \dots | a_1$ et M soit équivalente à la matrice :*

$$\left(\begin{array}{ccc|c} d_1 & & & 0 \\ & \ddots & & \\ & & d_r & \\ \hline & 0 & & 0 \end{array} \right) \quad (1)$$

(xxxiii). Quitte à parler d'algorithmes, autant le faire dans la tradition des informaticien.ne.s!

Préliminaires : matrices de Bézout et technique d'élimination ([DLQ14], chapitre IV, théorème 1.1)

Cette première partie vise à introduire les outils et techniques de base qui seront ceux qu'on utilisera pour arriver à la forme de Smith.

Considérons a et b deux éléments de $\mathbb{A}$ avec $a \neq 0$. L'objectif est de trouver une matrice inversible B telle que :

$$B \begin{pmatrix} a \\ b \end{pmatrix} = \begin{pmatrix} g \\ 0 \end{pmatrix} \quad (2)$$

où g est un autre élément de $\mathbb{A}$ ayant de bonnes propriétés. On distingue trois cas :

1. *Le cas trivial* où $b = 0$. Dans ce cas, il n'y a rien à faire.
2. *Le cas simple* où $a|b$. Dans ce cas, on écrit $b = da$ et on résout le problème par transvection, en posant :

$$B = \begin{pmatrix} 1 & 0 \\ d & 1 \end{pmatrix} \quad (3)$$

Dans ce cas, $g = a$.

3. *Le cas décisif* où a ne divise pas b . Tout repose alors sur les propriétés des anneaux principaux. D'une part, a et b admettent un pgcd noté g . D'une autre, un anneau principal satisfait le théorème de Bézout : on dispose de u et v tels que $au + bv = g$. Notons $a = a'g$, $b = b'g$. Si l'on pose :

$$B = \begin{pmatrix} u & v \\ -a' & b' \end{pmatrix} \quad (4)$$

on obtient bien :

$$B \begin{pmatrix} a \\ b \end{pmatrix} = \begin{pmatrix} g \\ 0 \end{pmatrix} \quad (5)$$

En effet :

$$0 = ab - ab = g(ab' - a'b) \quad (6)$$

et comme $\mathbb{A}$ est intègre, on trouve effectivement $ab' - a'b = 0$. La matrice B est par ailleurs inversible, puisque :

$$g = g(a'u + b'v) \quad (7)$$

et comme $\mathbb{A}$ est intègre, cela implique que $\det(B) = a'u + b'v = 1$.

Remarque : En réalité, le cas décisif contient les deux cas précédents à lui tout seul, mais bien faire la différence entre ces trois cas aidera grandement à comprendre l'algorithme de Smith. ♦

Cette étude contient à elle seule (presque) tous les ingrédients nécessaires pour la mise en place de l'algorithme. Si vous l'avez comprise, vous avez (presque) tout compris au développement !

Formalisons un peu tout ça :

Définition 55 (Matrice de Bézout). On appelle matrice de Bézout toute matrice carrée de $\mathbb{M}_n(\mathbb{A})$ de la forme :

$$\begin{pmatrix} 1 & & & & \\ & \ddots & & & \\ & & u & v & \\ & & -a' & b' & \\ & & & \ddots & 1 \end{pmatrix} \begin{matrix} \\ \\ \leftarrow i \\ \leftarrow j \\ \end{matrix}$$

$\begin{matrix} \uparrow & \uparrow \\ i & j \end{matrix}$

(où les pointillés figurent des 1) avec :

$$ub' + va' = 1 \quad (8)$$

Remarque : La multiplication à droite par une matrice de Bézout correspond à faire simultanément les opérations :

$$\begin{cases} L_i \leftarrow uL_i + vL_j \\ L_j \leftarrow -a'L_i + b'L_j \end{cases} \quad (9)$$

ce qui correspond à notre technique d'élimination dans le cas décisif. $\blacklozenge$

Tout est maintenant en place pour mettre en place l'algorithme.

Existence de la forme de Smith ([DLQ14], chapitre IV, 2.2 et 2.3)

Soit $M \in \mathbb{M}(\mathbb{A})$. Si M est nulle, elle est déjà sous forme de Smith. On suppose donc que M admet un coefficient non nul. Quitte à permuter des lignes et / ou des colonnes (ce qui correspond à multiplier à gauche et / ou à droite par des matrices de permutation, inversibles), on peut supposer $M(1,1) \neq 0$. On va dans un premier temps ignorer la condition de divisibilité sur les termes diagonaux et montrer que M est équivalente à une matrice diagonale, dont on s'occupera ensuite d'ordonner les termes. Pour cela, on procède de façon récursive sur la taille de la matrice. On commence par proposer un algorithme, dont on prouvera la terminaison et la correction, qui prenant la matrice M en entrée, renvoie une matrice équivalente à M de la forme :

$$\left(\begin{array}{c|c} a_1 & 0 \\ \hline 0 & M' \end{array} \right) \quad (\star)$$

avec M' une matrice de $\mathbb{M}_{n-1,p-1}(\mathbb{A})$ à laquelle on pourra de nouveau appliquer l'algorithme. Donnons dans un premier temps la procédure, qu'on expliquera ensuite :

Algorithme :

Entrée : $M \in \mathbb{M}(\mathbb{A})$.

Sortie : une matrice équivalente à M sous forme $(\star)$.

1. Etape préliminaire :

- (a) Si M est nulle, retourner M
- (b) Sinon, multiplier M par des matrices de permutation de sorte à ce que $M(1,1)$ soit non nul.

2. Pour i allant de 2 à n :
 - (a) Si $M(i, 1) = 0$, ne rien faire. (*Cas trivial*)
 - (b) Si $M(1, 1) | M(i, 1)$, éliminer $M(i, 1)$ par transvection. (*Cas simple*)
 - (c) Sinon, éliminer $M(i, 1)$ par transformation de Bézout. (*Cas décisif*)
3. Pour j allant de 2 à p :
 - (a) Si $M(1, j) = 0$, ne rien faire. (*Cas trivial*)
 - (b) Si $M(1, 1) | M(1, j)$, éliminer $M(1, j)$ par transvection. (*Cas simple*)
 - (c) Sinon, éliminer $M(1, j)$ par transformation de Bézout. **Stopper l'exécution de l'étape 3 et recommencer l'étape 2.** (*Cas décisif*)

Retourner M .

Des explications s'imposent !

- A chaque tour de boucle de l'étape 2, un coefficient de la première colonne est éliminé. Chacune de ces étapes ne perturbe aucune autre ligne que la première et la colonne cible, aussi, en sortie du i -ème tour de boucle, les coefficients 2 à i de la première colonne sont nuls.
- Idem pour l'étape 3 qui élimine successivement les coefficients de la première ligne.
- Le cas 3c pose un problème majeur : il perturbe la première colonne, en y réintroduisant éventuellement des termes non nuls. De même, l'étape 2c perturbe la première ligne. La terminaison de l'algorithme n'est donc pas assurée.
- Le coefficient $M(1, 1)$ ne change que lors des étapes décisives.

L'algorithme mieux compris, attelons-nous à montrer qu'il fonctionne. On notera pour plus de clarté $M^{(n)}$ la matrice obtenue à l'issue de la n -ième étape (avec $M^{(0)} = M$).

Correction de l'algorithme

La correction est assez facile à voir. Premièrement, $M^{(n)}$ est équivalente à M pour tout $n \in \mathbb{N}$ par récurrence immédiate. En effet chacun des cas correspond soit à multiplier (à gauche lors de l'étape 2, à droite lors de l'étape 3) par l'identité dans le cas trivial, par une matrice de transvection dans le cas simple ou par une matrice de Bézout dans le cas décisif, et ces matrices sont toutes inversibles.

De plus, on a la propriété suivante :

Lemme 56. *A l'issue du j -ième tour de boucle de l'étape 3, les coefficients de la première ligne situés en position 2 à j sont nuls. Si de plus ce j -ième tour de boucle ne s'est pas traduit par une étape décisive, tous les coefficients de la première ligne, à l'exception du premier, sont nuls.*

Démonstration. C'est une récurrence assez immédiate. Au commencement de l'étape 3, la première colonne est nulle, à l'exception du premier terme (cf remarques ci-dessus). A chaque tour de boucle, un nouveau coefficient est éliminé, et le reste de la première ligne n'est pas perturbé (à l'exception encore du premier coefficient). Comme seule l'étape décisive peut perturber la première colonne, on obtient le lemme. $\square$

Comme il est évident que l'algorithme ne peut pas terminer par une étape décisive, l'algorithme est nécessairement correct. Reste à montrer qu'il s'arrête effectivement...

Terminaison de l'algorithme

Si vous aimez les anneaux, c'est là que ça devient intéressant ! Pour prouver la terminaison d'un algorithme, il nous faut exhiber un "variant de boucle", c'est-à-dire une quantité qui décroît à chaque tour de boucle. Le plus souvent, on choisit un entier pour variant, et comme $\mathbb{N}$ est bien ordonné, la suite ne peut être strictement décroissante, ce qui est le cœur de l'argument de terminaison. C'est ce qui se passe dans le cas euclidien : le variant de boucle est le stathme du coefficient $M^{(n)}(1, 1)$, qui ne décroît strictement que lors des étapes décisives.

Ici, nous n'avons pas de stathme. Toutefois, nous n'avons pas réellement besoin que le variant de boucle soit à valeurs dans $\mathbb{N}$: il suffit qu'il soit à valeurs dans un poset artinien, c'est-à-dire un ensemble partiellement ordonné qui n'admet aucune suite strictement décroissante. De façon duale, on peut plutôt chercher pour variant une suite croissante à valeurs dans un poset noethérien... vous voyez on je veux en venir n'est-ce pas ? C'est là la beauté de l'argument dans le cas principal : nous avons à notre disposition un splendide poset noethérien en la qualité de l'ensemble des idéaux de $\mathbb{A}$ muni de l'inclusion, puisqu'un anneau principal est toujours noethérien.^(xxxiv) Voici donc notre variant : on pose :

$$V_n := \langle M^{(n)}(1, 1) \rangle \quad (10)$$

Lemme 57. *La suite (V_n) est une suite croissante d'idéaux de $\mathbb{A}$. De plus, on a :*

$$V_n \subsetneq V_{n+1} \iff \text{la } n\text{-ième étape est un cas décisif} \quad (11)$$

Démonstration. Plaçons-nous au début de la $n + 1$ -ème étape. On a alors $\langle M^{(n)}(1, 1) \rangle = V_n$. Si cette étape se solde par un cas trivial ou simple, le coefficient en haut à gauche n'est pas perturbé, c'est-à-dire que $M^n(1, 1) = M^{n+1}(1, 1)$ et donc $V_n = V_{n+1}$. S'il s'agit par contre d'un cas décisif, $M^{(n)}(1, 1)$ est remplacé par le pgcd de lui-même et d'un élément non-nul de $\mathbb{A}$, qui est donc un diviseur non nul strict de $M^{(n)}(1, 1)$. Cette divisibilité stricte se traduit exactement par $V_n \subsetneq V_{n+1}$. $\square$

Ainsi, par noetherianité de $\mathbb{A}$, la suite (V_n) est stationnaire à partir d'un certain rang, c'est-à-dire que que l'algorithme n'effectue qu'un nombre fini d'étapes décisives. En d'autres termes, l'algorithme termine.

Travail de la forme diagonale

Une application récursive de l'algorithme fournit donc une matrice équivalente à M dont tous les coefficients hors de la diagonale sont nuls. Notons cette matrice $\bar{M}$ et r le plus grand indice tel que $\bar{M}(r, r) \neq 0$. Il résulte de l'algorithme que pour $i \leq r$, $\bar{M}(i, i) \neq 0$. On note ces termes $a_1, \dots, a_r$. Il nous reste à montrer que l'on peut ordonner ces coefficients pour la divisibilité. Pour cela, il suffit de le constater avec une matrice de taille 2×2 . Considérons :

$$A = \begin{pmatrix} a & 0 \\ 0 & b \end{pmatrix} \quad (12)$$

avec a et b non nuls. On veut montrer que A est équivalente à une matrice de la forme :

$$\begin{pmatrix} c & 0 \\ 0 & g \end{pmatrix} \quad (13)$$

(xxxiv). Si tout ceci ne vous parle pas beaucoup, je vous invite à voir l'annexe.

où g et c sont non nuls et $g|c$. On propose une suite d'opérations élémentaires :

$$\begin{pmatrix} a & 0 \\ b & b \end{pmatrix} = \begin{pmatrix} a & 0 \\ 0 & b \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix} \quad (14)$$

Notons $g = a \wedge b$, u et v des coefficients de Bézout tels que $au + bv = g$ et a' et b' tels que $a'g = a$, $b'g = b$. On effectue une transformation de Bézout :

$$\begin{pmatrix} g & vb \\ 0 & a'b \end{pmatrix} = \begin{pmatrix} u & v \\ -b' & a' \end{pmatrix} \begin{pmatrix} a & 0 \\ b & b \end{pmatrix} \quad (15)$$

On a alors que $g|vb$ et $g|a'b$. Par transvection, on élimine le coefficient en haut à droite, puis une multiplication par une matrice de permutation donne la forme voulue :

$$\begin{pmatrix} a'b & 0 \\ 0 & g \end{pmatrix} = \begin{pmatrix} g & vb \\ 0 & a'b \end{pmatrix} \begin{pmatrix} 1 & -vb' \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad (16)$$

Ainsi, on peut utiliser cette technique pour imposer que le coefficient a_r divise tous les autres, puis on impose que le coefficient a_{r-1} divise tous ceux de rang supérieur. Malheureusement, à cette étape, le coefficient a_r pourrait ne plus diviser a_{r-1} , donc il faut recommencer avec a_r , puis avec a_{r-1} , etc. C'est encore un argument de noetherianité qui prouve que la procédure termine, et à la fin, on obtient une jolie forme de Smith de la matrice M !

Unicité des facteurs invariants ([BMP05], proposition 6.78)

Remarquons tout d'abord que l'entier r est unique. En effet, comme $\mathbb{A}$ est intègre, on peut voir M comme une matrice à coefficients dans le corps des fractions de $\mathbb{A}$, équivalente à sa forme de Smith. Or il est alors clair que r est le rang de M , qui est invariant par équivalence matricielle.

Pour montrer l'unicité modulo l'association des facteurs invariants, on va exploiter les *idéaux déterminantiels* :

Définition 58 (Idéaux déterminantiels). Lorsque M est une matrice de $\mathbb{M}(\mathbb{A})$, on appelle k -ième idéal déterminantiel de M , noté $\Delta_k(M)$ l'idéal de $\mathbb{A}$ engendré par les mineurs de taille k de M .

Proposition 59. Si M et N sont deux matrices équivalentes de $\mathbb{M}(\mathbb{A})$, alors on a :

$$\forall k \in \llbracket 1 \min\{n, p\} \rrbracket, \Delta_k(M) = \Delta_k(N) \quad (17)$$

Démonstration. Considérons tout d'abord le cas où il existe $P \in GL_n(\mathbb{A})$ tel que $M = PN$. Alors les colonnes de M , vues comme éléments de $\mathbb{A}^n$, sont combinaisons linéaires à coefficients dans $\mathbb{A}$ des colonnes de N , et par suite, les colonnes de toute matrice extraite de M sont combinaisons linéaires des colonnes de la matrice extraite de N selon les mêmes indices. Or le déterminant est multilinéaire. On en déduit que tout mineur de taille k de M est combinaison linéaire des mineurs de taille k de N , et donc $\Delta_k(M) \subset \Delta_k(N)$. Or l'inversibilité de P fait que la situation est symétrique en M et N : $N = P^{-1}M$, donc on a l'égalité des idéaux déterminantiels.

Le même raisonnement s'applique dans le cas où $M = NP$ en raisonnant sur les lignes plutôt que sur les colonnes, d'où la proposition. $\square$

Dans notre cas, la condition de divisibilité sur les facteurs invariants trivialise l'expression des idéaux déterminantiels : si $a_r | \dots | a_1$ sont les termes diagonaux d'une forme de Smith de M , alors pour $k \leq \min\{n, p\}$, on a :

$$\Delta_k(M) = \left\langle \prod_{i=0}^{\min\{k, r-1\}} a_{r-i} \right\rangle \quad (18)$$

Ainsi, si M est équivalente à une autre matrice sous forme de Smith, dont on note $b_r | \dots | b_1$ les coefficients diagonaux, on a :

$$\forall k \in \llbracket 0, r-1 \rrbracket, \left\langle \prod_{i=1}^k a_{r-i} \right\rangle = \left\langle \prod_{i=1}^k b_{r-i} \right\rangle \quad (19)$$

En particulier, pour $k = r-1$, $\langle a_r \rangle = \langle b_r \rangle$, c'est-à-dire que a_r et b_r sont associés. De même pour $a_r a_{r-1}$ et $b_r b_{r-1}$, puis de proche en proche, on en déduit que les (a_i) et les (b_i) sont associés, d'où l'unicité modulo l'association des facteurs invariants.

Annexe - anneaux noetheriens, domaines de Bézout, anneaux principaux...

J'ai trouvé intéressant dans cette annexe de discuter un peu des hypothèses qui sont mises sur l'anneau de base $\mathbb{A}$ vis-à-vis du théorème de réduction de Smith, d'une part pour observer qu'on a atteint les hypothèses optimales pour que notre algorithme fonctionne, d'une autre pour tisser des liens entre différents types d'anneaux qu'on peut être amené à manipuler.

Notre algorithme repose sur deux résultats fondamentaux qui dépendent de la nature intrinsèque de l'anneau : la noetherianité et le théorème de Bézout. On va voir qu'un anneau dans lequel on dispose de ces outils est nécessairement principal.

Théorème 60 (Bézout). *Soit $\mathbb{A}$ un anneau intègre. Les assertions suivantes sont équivalentes :*

1. *Pour tout couple $(a, b) \in \mathbb{A}^2$, il existe $g \in \mathbb{A}$ un pgcd de a et b et un couple $(u, v) \in \mathbb{A}^2$ tel que :*

$$au + bv = g \quad (20)$$

2. *Tout idéal de type fini de $\mathbb{A}$ est principal.*

Si ces conditions sont remplies, on dit que $\mathbb{A}$ est un domaine de Bézout.

Démonstration.

Dans le sens direct , considérons $I = \langle a_1, \dots, a_n \rangle$ un idéal de type fini de $\mathbb{A}$. Le théorème de Bézout donne, par récurrence immédiate, l'existence de g un pgcd de la famille $(a_1, \dots, a_n)$ et de coefficients de Bézout $u_1, \dots, u_n \in \mathbb{A}^n$ tels que :

$$\sum_{i=1}^n u_i a_i = g \quad (21)$$

Donc $I \subset \langle g \rangle$. L'inclusion réciproque est immédiate car g divise chacun des a_i . Donc I est un idéal principal.

Dans le sens réciproque , on se donne a et b des éléments de $\mathbb{A}$. Par hypothèse, il existe $g \in \mathbb{A}$ tel que :

$$\langle a, b \rangle = \langle g \rangle \quad (22)$$

Ceci implique que g est un pgcd de a et de b pour lequel on dispose de coefficients de Bézout.

□

Ainsi, il est clair que notre algorithme ne peut pas fonctionner hors d'un domaine de Bézout. Voyons ce qui se passe si on ajoute la noetherianité :

Proposition 61. *Soit $\mathbb{A}$ un anneau intègre. Les assertions suivantes sont équivalentes :*

1. $\mathbb{A}$ est principal.
2. $\mathbb{A}$ est un domaine de Bézout noetherien.

Démonstration.

Dans le sens direct , il est évident que si $\mathbb{A}$ est principal, c'est un domaine de Bézout. Considérons une suite croissante (I_n) d'idéaux de $\mathbb{A}$. Pour $n \in \mathbb{N}$, on se donne a_n un générateur de I_n et on pose :

$$\langle I \rangle := \langle a_n \mid n \in \mathbb{N} \rangle \quad (23)$$

Alors il est clair que chacun des (I_n) est contenu dans I . Si de plus g est un générateur de I , il existe une famille finie d'indices $i_1, \dots, i_k$ et des coefficients $(u_1, \dots, u_k) \in \mathbb{A}^k$ tels que :

$$g = \sum_{j=1}^k u_j a_{i_j} \quad (24)$$

Donc :

$$I \subset \sum_{j=1}^k I_{i_j} = I_{\max_{1 \leq j \leq k} i_j} \quad (25)$$

Ce qui implique que la chaîne (I_n) est ultimement stationnaire.

Dans le sens indirect , on considère I un idéal quelconque de $\mathbb{A}$ et $\mathcal{I}$ l'ensemble des idéaux de type fini de $\mathbb{A}$ contenus dans I qu'on munit de l'inclusion. Alors, par noetherianité de $\mathbb{A}$, et comme $\mathcal{I}$ contient l'idéal nul, $(\mathcal{I}, \subset)$ est un ensemble inductif non vide. Par le lemme de Zorn, il contient un élément maximal M . Il est évident que $M = I$, car sinon, on pourrait prendre un élément $a \in I \setminus M$ et considérer l'idéal de type fini contenu dans $I \setminus M + \langle a \rangle$, ce qui contredirait la maximalité de M . Donc I est de type fini, et comme $\mathbb{A}$ est un domaine de Bézout, I est principal.

□

1.13 Théorème de Lie-Kolchin ★★

Recasages possibles : 103, 106, 151, 156, 204 ^(xxxv)

On va démontrer ici un théorème de réduction simultanée pour les endomorphismes d'un $\mathbb{C}$ espace vectoriel de dimension finie, qui, s'il peut paraître anecdotique au niveau de l'agreg, est parait-il fort utile en théorie de Galois différentielle (quoi que ça veuille dire). Et qui a le bon goût de se recaser dans un joli palmarès de leçons maudites.

Il est assez connu qu'une partie commutative de $M_n(K)$ composée de matrices trigonalisables est co-trigonalisable. On aimerait renforcer ce résultat pour des parties non-commutatives. En se dirigeant du côté de la théorie des groupes, on trouve effectivement une notion qui fait l'intermédiaire entre les groupes abéliens et les groupes généraux : c'est celle de groupe résoluble. C'est dans cette direction là que nous allons nous diriger. Malheureusement, la seule résolubilité ne sera pas suffisante, aussi nous allons devoir ajouter à cela un peu de topologie !

Pour un groupe G quelconque et $g, h \in G^2$, on définit le commutateur de g et h par :

$$[g, h] := ghg^{-1}h^{-1} \quad (1)$$

On écrira $D(G)$ pour le groupe dérivé de G , c'est-à-dire le sous-groupe de G engendré par les commutateurs. Par ailleurs, pour $r \in \mathbb{N}$, on définit :

$$D^r(G) := \begin{cases} G & \text{si } r = 0 \\ D(D^{r-1}(G)) & \text{sinon} \end{cases} \quad (2)$$

L'objectif de ce développement est de démontrer :

Théorème 62 (Lie-Kolchin ([Cha05])). *Soit E un $\mathbb{C}$ -espace vectoriel de dimension $d \in \mathbb{N}$. On considère G un sous-groupe de $GL(E)$ qu'on suppose connexe et résoluble. Il existe $\mathcal{B}$ une base de E dans laquelle la matrice de tout élément de G est triangulaire supérieure.*

La démonstration repose sur la recherche de sous-espaces de E stables par tout élément de G pour pouvoir faire une récurrence sur la dimension. On utilisera sans démonstration le lemme de co-trigonalisation de Frobenius. C'est un résultat classique qu'il faut absolument savoir prouver pour s'attaquer à Lie-Kolchin, mais qui en terme de temps, pourrait être un développement en soi, et je pense qu'il vaut mieux prendre du temps pour prouver le lemme de topologie qui interviendra plus tard et qui justifie bien mieux le recasage dans la leçon de connexité.

Lemme 63 (Frobenius ([Goua], IV, théorème 3)). *Soit K un corps, et soit $S \subset M_n(K)$ un ensemble quelconque de matrices trigonalisables qui commutent deux à deux. Alors S est co-trigonalisable.*

L'initialisation de la récurrence est immédiate : si $d = 1$, alors dans n'importe quelle base, la matrice de tout élément de $GL(E)$ est diagonale. Supposons donc $d \geq 2$ et que pour tout

(xxxv). Je me méfierais quand même pour cette dernière leçon, j'aurais un peu peur qu'on m'accuse de la détourner pour faire de l'algèbre.

espace vectoriel F de dimension inférieure stricte à d et que tout sous-groupe connexe résoluble de $GL(F)$ est co-trigonalisable.

Dans un premier temps, on suppose qu'il existe $\{0\} \subsetneq V \subsetneq E$ un sous-espace non-trivial de E stable par tous les éléments de G . Soit W un supplémentaire de V dans E , et soit $\mathcal{B}$ une base de E bien adaptée à la décomposition $E = V \oplus W$. Alors, pour $g \in G$ quelconque, on a la forme matricielle suivante :

$$Mat_{\mathcal{B}}(g) = \left(\begin{array}{c|c} M_1 & * \\ \hline 0 & M_2 \end{array} \right) \quad (3)$$

Notons g_1 et g_2 les endomorphismes de V et W dont les matrices dans les bases $\mathcal{B}$ intersecté respectivement avec V et W sont $M_1(g)$ et $M_2(g)$. On écrit $p_V : E \twoheadrightarrow V$ et $\iota_V : V \hookrightarrow E$ la projection et l'inclusion canonique de V (idem pour W avec p_W et ι_W). On observe alors que les applications :

$$g \mapsto g_1 = p_V \circ g \circ \iota_V \in GL(V) \quad g \mapsto g_2 = p_W \circ g \circ \iota_W \in GL(W) \quad (4)$$

sont des morphismes des groupes continus, l'aspect morphisme découlant directement de la multiplication matricielle en exploitant la forme triangulaire par blocs. Ainsi, les ensembles :

$$\{g_1, g \in G\} \quad \{g_2, g \in G\} \quad (5)$$

s'identifient à des sous-groupes connexes de $GL(V)$ et $GL(W)$. Ces groupes sont résolubles car $g \mapsto g_1$ et $g \mapsto g_2$ sont surjectives, aussi tout commutateur de l'image est l'image d'un commutateur de G , qui est lui-même résoluble. Donc par hypothèse de récurrence, il existe des bases $\mathcal{B}_1$ et $\mathcal{B}_2$ qui co-trigonalisent ces deux groupes. En concaténant ces bases, on obtient une base de co-trigonalisation de G .

Ainsi, la démonstration sera complète si l'on montre qu'un tel sous-espace **non-trivial** existe forcément. Pour cela, on procède par l'absurde : supposons que les seuls sous-espaces G -stables de E sont $\{0\}$ et E . On va montrer qu'alors d ne peut être égal qu'à 1, ce qui contredira l'hypothèse précédemment faite. On raisonne sur l'indice de résolubilité de G , c'est-à-dire le plus petit entier r tel que $D^r(G) = 1$. Notons celui-ci r .

On remarque que $r \geq 2$. En effet, si $r = 0$, G serait trivial, et si $r = 1$, G est abélien, donc co-trigonalisable d'après le lemme de co-trigonalisation. Dans ce cas, il existerait une droite stable par tous les éléments de G qui contredirait l'inexistence d'un sous-espace G -stable non-trivial. Considérons $H := D^{r-1}(G)$.

Posons alors :

$$V := \text{Vect}(\{x \in E \mid x \text{ est un vecteur propre de tous les éléments de } G\}) \quad (6)$$

Comme H est un groupe abélien non-trivial, car son indice de résolubilité est 1. Ainsi, par le lemme de co-trigonalisation, il existe une droite de E stable par tous les éléments de H , et donc V est non nul. Mais H est distingué dans G . Ainsi, si x est un vecteur propre commun à tout élément de H , on a :

$$\forall h \in H, \quad h(g(x)) = g \circ \underbrace{g^{-1} \circ h \circ g}_{\in H}(x) \quad (7)$$

$$= \lambda g(x) \quad (8)$$

où λ est la valeur propre associée à x pour l'endomorphisme $g^{-1} \circ h \circ g$. Ainsi $g(x) \in V$, et par linéarité, V est stable par tous les éléments de G . Ainsi, par hypothèse, $V = E$. Par théorème

de la base extraite, il existe une base $\mathcal{B}_D$ de E composée de vecteurs propres communs à tous les éléments de H . En d'autres termes, H est co-diagonalisable.

On fait alors agir G sur H par conjugaison (ce qui est possible puisque H est distingué dans G). Soit $h \in H$. L'application :

$$\begin{aligned} G &\rightarrow H \\ g &\mapsto g \circ h \circ g^{-1} \end{aligned} \tag{9}$$

est continue et surjective sur l'orbite de h . Celle-ci est donc connexe. Mais dans la base $\mathcal{B}_D$, tous les éléments de l'orbite de h sont diagonaux et ont les mêmes valeurs propres, car ils sont tous semblables. Ceci impose que cette orbite est finie. Or une partie connexe finie non-vide d'un espace métrique est réduite à un singleton. Donc l'orbite de h sous l'action de G est triviale : h appartient au centre de G , noté $\mathcal{Z}(G)$.

Mais on sait, par de l'algèbre linéaire élémentaire, si u et v sont deux endomorphismes de E qui commutent, tout espace propre de u est stable par v . Donc les espaces propres de h sont stables par tous les éléments de G . Ainsi, si λ est une valeur propre de h et E_λ est l'espace propre associé, E_λ est non nul et stable par tous les éléments de G , donc par hypothèse $E_\lambda = E$ et h est l'homothétie de rapport λ , et ce raisonnement vaut pour tous les éléments de H .

On a supposé que r était supérieur ou égal à 2. Donc $r - 1 \geq 1$ et on a :

$$H = D^{r-1}(G) < D(G) < D(GL(E)) = SL(E) \tag{10}$$

Donc le déterminant des éléments de H est 1. Puisque H est composé uniquement d'homothéties, le rapport de celles-ci doit être une racine d -ième de l'unité. En particulier, H est un groupe fini. Il ne reste qu'à démontrer ceci pour conclure :

Lemme 64. *Le groupe dérivé de G est connexe.*

Démonstration. $D(G)$ est, par définition, le sous-groupe de G engendré par les commutateurs. On peut utiliser son écriture explicite :

$$D(G) = \bigcup_{n=0}^{+\infty} \left\{ \prod_{i=1}^n [g_i, h_i], g_i, h_i \in G \right\} \tag{11}$$

puisque l'inverse d'un commutateur est encore un commutateur. Or l'application :

$$G^2 \rightarrow D(G) \tag{12}$$

$$(g, h) \mapsto [g, h] \tag{13}$$

est continue. D'où $D(G)$ est une union croissante de parties connexes : c'est encore un connexe. $\square$

Par récurrence immédiate, H est un groupe connexe fini : c'est le groupe trivial. Ceci contredit le fait que r est l'indice de résolubilité de G , et nie notre hypothèse initiale : il existe un sous-espace non trivial G -stable, ce qui permet de conclure la preuve.

Remarque : Pour celles et ceux qui aiment ce point de vue, je glisse ici une petite remarque : l'essentiel de la preuve est en fait une démonstration d'existence d'un sous-espace G -stable. Ceci peut se reformuler en utilisant d'autres langages, qui s'éloignent de l'algèbre linéaire classique

pour mieux refléter la structure qu'on étudie. Ainsi, si vous aimez les représentations de groupes, vous apprécierez peut-être de savoir qu'on a démontré que si G est un sous-groupe connexe résoluble de $GL(E)$, alors la représentation évidente qu'on a de G par inclusion dans $GL(E)$ n'est irréductible que si $d = 1$. Si vous préférez le langage des modules, on peut aussi dire que le $\mathbb{C}[G]$ -module E muni de la multiplication externe $g \cdot x := g(x)$ pour tout $g \in G$ et $x \in E$ est un module simple si, et seulement si, $d = 1$. $\blacklozenge$

Et à quoi ça sert ?

Personnellement, je dois avouer n'avoir jamais utilisé un résultat de trigonalisation simultanée pour faire quoi que ce soit d'autre qu'un exercice gratuit. Comme le jury a l'air d'apprécier ces résultats (d'après le rapport de la leçon 156), j'ai farfouillé un peu dans le livre de Chambert-Loir, et il semble que ces résultats, en particulier le théorème de Lie-Kolchin soit utilisé en théorie de Galois différentielle. Sans chercher à rentrer dans les détails (d'autant que je n'y connais goutte), c'est une théorie qui cherche entre autres à appliquer des méthodes similaires à celles de la théorie de Galois classiques pour déterminer la résolubilité d'équations différentielles linéaires par quadrature. Pour ce faire, on définit la notion de dérivation abstraite, de corps différentiel et d'extension de corps différentiel, puis on s'intéresse au groupe d'automorphismes de ces extensions. On s'intéresse souvent, paraît-il, à la composante connexe de l'identité dans ces groupes, et c'est là qu'intervient le théorème de Lie-Kolchin.

1.14 Théorème de Kronecker ★★ ★

Recasages possibles : 102, 144

On va dans ce développement démontrer un théorème attribué à Kronecker qui montre une certaine rigidité sur le lieu des racines d'un polynôme à coefficients entiers. J'ai ajouté deux applications, une qu'on trouve sur beaucoup d'autres versions et une autre moins connue. Selon votre vitesse, vous pouvez choisir d'en traiter une à l'oral. Je pense que celle sur les matrices apportera un vrai plus au développement, l'autre apparaissant plus comme une curiosité amusante.

C'est le développement sur lequel je suis tombé lors du vrai oral, et j'ai obtenu 19,25 avec celui-ci (bien que le développement ne fasse pas tout) !

Etant donné un polynôme P à coefficients complexes, on notera $\mathcal{Z}(P)$ un $\deg(P)$ -uplet composé de toutes ses racines complexes avec multiplicité. Pour $n \in \mathbb{N}^*$ et $1 \leq k \leq n$, on désigne par $\sigma_{n,k}$ le k -ième polynôme symétrique élémentaire à n indéterminées, soit :

$$\sigma_{n,k} := \sum_{I \in \mathcal{P}_k(\llbracket 1, n \rrbracket)} \prod_{i \in I} T_i \quad (1)$$

avec $\mathcal{P}_k(\llbracket 1, n \rrbracket)$ l'ensemble des parties à k éléments de $\llbracket 1, n \rrbracket$.

Notre objectif est de démontrer :

Théorème 65 (Kronecker ([CP22], exercice 3.17)). *Soit $P \in \mathbb{Z}[X]$ unitaire dont toutes les racines complexes sont de module au plus 1. Toutes les racines non nulles de P sont des racines de l'unité.*

Le théorème

Commençons par remarquer qu'on peut écrire $P = X^p Q$ avec $Q(0) \neq 0$. Dans ce cas, Q est nécessairement à coefficients entiers, unitaire, et ses racines sont les racines non nulles de P . Quitte à travailler sur Q , on supposera donc dans toute la suite que $P(0) \neq 0$.

On va montrer que toutes les racines de P sont d'ordre fini dans $\mathbb{C}^\times$. Notons n le degré de P . On introduit $\Omega_n := \{Q \in \mathbb{Z}[X] \text{ unitaire} \mid \deg(Q) = n \wedge \mathcal{Z}(Q) \in \bar{B}(0, 1)^n\}$. C'est un ensemble fini. En effet, prenons $Q \in \Omega_n$. Alors, en notant $(r_1, \dots, r_n) := \mathcal{Z}(Q)$:

$$Q = X^n + \sum_{k=0}^{n-1} (-1)^{n-k} \sigma_{n,k}(r_1, \dots, r_n) X^k \quad (2)$$

par relations racines-coefficients. Or on a :

$$|\sigma_{n,k}(r_1, \dots, r_n)| \leq \sum_{I \in \mathcal{P}_k(\llbracket 1, n \rrbracket)} \prod_{i \in I} |z_i| \leq \binom{n}{k} \quad (3)$$

Comme ces coefficients sont entiers, ils ne peuvent prendre qu'un nombre fini de valeurs et Ω_n est fini.

Il découle alors immédiatement que l'ensemble des racines des polynômes de Ω_n est fini. Notons $(z_1, \dots, z_n) = \mathcal{Z}(P)$. Les (z_i) sont non nuls car 0 n'est pas une racine de P par hypothèse.

Posons :

$$\forall l \in \mathbb{N}^*, Q_l := \prod_{i=1}^n (X - T_i^l) \in \mathbb{Z}[X][T_1, \dots, T_n] \quad (4)$$

Pour tout $l \in \mathbb{N}$, Q_l est symétrique en les indéterminées (T_i) . Par théorème de structure des polynômes symétriques, il existe $R_l \in \mathbb{Z}[X][S_1, \dots, S_n]$ tel que :

$$Q_l = R_l(\sigma_{n,1}, \dots, \sigma_{n,n}) \quad (5)$$

Ainsi, le polynôme $Q_l(z_1, \dots, z_n) = R_l(\sigma_{n,1}(z_1, \dots, z_n), \dots, \sigma_{n,n}(z_1, \dots, z_n))$ est à coefficients entiers. Comme il est trivialement unitaire et que ses racines sont de module au plus 1, les (z_i^l) sont racines d'un polynôme de Ω_n . Il résulte que le morphisme de groupes $l \mapsto z_i^l$ ne peut pas être injectif, et donc z_i est d'ordre fini, ce qui achève la démonstration du théorème! $\square$

Deux applications

Commençons par une première application surprenante de ce théorème, qui complètera agréablement ce développement un peu court :

Application 66. (Sous-groupes finis de $GL_n(\mathbb{Z})$ ([CP22], remarque 3.20)) Soit $n \in \mathbb{N}^*$. Si G est un sous-groupe fini de $GL_n(\mathbb{Z})$, alors l'ordre de G satisfait l'inégalité :

$$|G| \leq \prod_{k=0}^{n-1} (3^n - 3^k) \quad (6)$$

Démonstration. Commençons par donner du sens à ce majorant : il s'agit du cardinal de $GL_n(\mathbb{F}_3)$, où $\mathbb{F}_3$ désigne le corps à 3 éléments. Ceci nous aiguille pour amorcer la preuve : on va construire une injection de G dans $GL_n(\mathbb{F}_3)$. Le candidat est naturel : la projection naturelle $\mathbb{Z} \twoheadrightarrow \mathbb{F}_3$ s'étend en un morphisme d'anneaux $\mathbb{M}_n(\mathbb{Z}) \twoheadrightarrow \mathbb{M}_n(\mathbb{F}_3)$ en l'appliquant coordonnée par coordonnée. Comme un morphisme d'anneaux envoie toujours les inversibles sur les inversibles, celui-ci se restreint en un morphisme de groupes $GL_n(\mathbb{Z}) \twoheadrightarrow GL_n(\mathbb{F}_3)$ dont on notera π la restriction à G . Montrons que π est injectif.

Soit $M \in \text{Ker}(\pi)$. Comme M appartient à G , c'est une matrice d'ordre fini, donc toutes ses valeurs propres complexes sont des racines de l'unité (car le polynôme $X^{|G|} - 1$ annule M). Par ailleurs, $\pi(M) = \bar{I}_n$ la matrice identité de $GL_n(\mathbb{F}_3)$. Ainsi, il existe une matrice A à coefficients entiers telle que :

$$M = I_n + 3A \quad (7)$$

On aura montré notre résultat si on montre que $A = 0$. On isole A dans l'équation :

$$A = \frac{1}{3}(M - I_n) \quad (8)$$

On déduit deux choses de cette écriture :

1. A est diagonalisable dans $\mathbb{C}$. En effet, M est annulée par un polynôme scindé à racines simples, donc est diagonalisable, et la translation par l'identité n'y change rien.
2. Les valeurs propres de A sont exactement les éléments de la forme $\frac{\lambda-1}{3}$ avec λ une valeur propre de M .

Or :

$$\forall \lambda \in Sp(M), \left| \frac{\lambda - 1}{3} \right| \leq \frac{|\lambda| + 1}{3} < 1 \quad (9)$$

car le spectre de M est inclus dans $\mathbb{U}$. On va ainsi chercher à appliquer le théorème de Kronecker au polynôme caractéristique de A , qui est unitaire et à coefficients entiers. Ses racines non nulles sont nécessairement des racines de l'unité, or on a montré que celles-ci sont de module inférieur strict à 1. Donc χ_A n'a pas de racine non nulle et A est nilpotente. Comme elle est également disagonalisable, elle est nulle, et π est injectif. On a bien :

$$|G| \leq |GL_n(\mathbb{F}_3)| = \prod_{k=0}^{n-1} (3^n - 3^k) \quad (10)$$

□

Remarque : Pourquoi $\mathbb{F}_3$, et pas un autre corps fini ? En fait, pour n'importe quel nombre premier p , on peut construire pareillement un morphisme $\pi_p : G \rightarrow GL_n(\mathbb{F}_p)$ et dérouler la démonstration. Toutefois, pour $p = 2$, on n'aura plus la majoration stricte du module des valeurs propres de A , et donc on ne pourra pas conclure que $A = 0$. Pour les autres nombres premiers, tout fonctionnent pareil ; le choix $p = 3$ est simplement celui qui donne la majoration la plus fine. ♦

Voici une deuxième application, beaucoup plus rapide à traiter, qui demande toutefois de connaître quelques résultats sur la cyclotomie.

Application 67. Soit $P \in \mathbb{Z}[X]$ unitaire et irréductible sur $\mathbb{Q}$. Si toutes les racines complexes de P sont de module au plus 1, alors P est X ou un polynôme cyclotomique.

Démonstration. Si $P(0) = 0$, $X|P$ et par irréductibilité de P , $P = X$. Dans le cas contraire, toutes les racines de P sont des racines de l'unité d'après le théorème de Kronecker. Comme P est irréductible sur $\mathbb{Q}$, c'est le polynôme minimal d'une racine de l'unité, c'est-à-dire que c'est un polynôme cyclotomique. □

Chapitre 2

Développements d'analyse

But : trouver tous les 🍍 tels que :

$$\sum_{k=1}^{+\infty} \frac{1}{k^{\text{🍍}}} = 0$$

* « Développement présenté avec pédagogie »

2.1 Calcul d'une transformée de Fourier par le théorème de résidus ★★

Le calcul explicite de transformées de Fourier intervient dans de nombreuses situations assez variées ; on comprend que développer des méthodes pour franchir les problèmes posés par les calculs d'intégrales est important. On va ici explorer ces calculs sous l'angle de l'analyse complexe et obtenir des expressions de transformées de Fourier de fonctions assez régulières (par exemple certaines fractions rationnelles) à partir de calcul de résidus. On donnera une application de ces calculs pour obtenir la fonction caractéristique d'une variable aléatoire de Cauchy.

Il s'agit d'un développement plutôt facile, et surtout un peu court. Je vous encourage, lors de sa présentation, à mettre l'accent sur le côté pédagogique et la clarté de l'exposé. J'ai entendu un retour de quelqu'un qui a eu une excellente note en présentant ce développement à l'oral : son apparente simplicité n'est pas une faiblesse !

Etant donnée g une fonction de classe L^1 sur $\mathbb{R}$, on notera

$$\mathcal{F}[g] : \xi \mapsto \int_{\mathbb{R}} g(x) e^{-i\xi x} dx \quad (1)$$

sa transformée de Fourier. Si h est une fonction méromorphe sur un ouvert de $\mathbb{C}$ et a un pôle de $\mathbb{C}$, on désignera par $\text{Res}(h, a)$ le résidu de h au point a , qu'on définit comme le coefficient d'indice -1 de h dans son développement en série de Laurent au voisinage de a . On désignera par $\mathbb{H}$ le demi-plan de Poincaré formé des complexes de partie imaginaire strictement positive, et $\bar{\mathbb{H}}$ sa fermeture composée des complexes de partie imaginaire positive.

Dans toute la suite, on fixe f une fonction de classe L^1 sur $\mathbb{R}$. On suppose :

1. f se prolonge en une fonction méromorphe sur un voisinage de $\bar{\mathbb{H}}$ qu'on notera encore f .
2. f a un nombre fini de pôles. On notera A l'ensemble de ses pôles contenus dans $\mathbb{H}$ ⁽ⁱ⁾
3. $|f(z)| \xrightarrow{|z| \rightarrow +\infty} 0$

Notre objectif est de donner une expression de la transformée de Fourier de f sur $\mathbb{R}_*$ en exploitant l'analyse complexe :

Théorème 68 ([AM04], §8.4.3). Soit $\xi > 0$. On a :

$$\mathcal{F}[f](-\xi) = \int_{\mathbb{R}} f(x) e^{ix\xi} dx = 2i\pi \sum_{a \in A} \text{Res}(f, a) \quad (2)$$

où f_ξ est la fonction $z \mapsto f(z) e^{iz\xi}$.

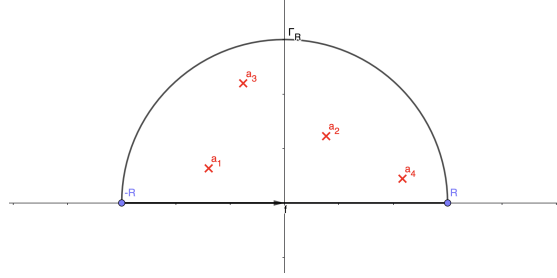
Remarque : Bien que ces hypothèses puissent avoir l'air très restrictives, il s'agit en réalité d'une situation assez fréquente. Par exemple, toute fraction rationnelle sans pôle réel de degré inférieur ou égal à -2 satisfait ces hypothèses. ♦

Démonstration. L'idée est d'intégrer sur un contours bien choisi tracé dans $\bar{\mathbb{H}}$ et d'y appliquer le théorème de résidus. Soit $R > 0$ tel que

$$\forall a \in A, |a| < R \quad (3)$$

(i). comme on a supposé $f|_{\mathbb{R}}$ de classe L^1 , f n'a pas de pôle réel.

(un tel R existe par finitude de A). On considère alors le contours suivant :



où Γ_R désigne le demi-cercle de centre 0 et de rayon R , orienté dans le sens trigonométrique. Une application du théorème de résidus fournit :

$$2i\pi \sum_{a \in A} \text{Res}(f_\xi, a) = \int_{-R}^R f(x)e^{i\xi x} dx + \int_{\Gamma_R} f(z)e^{i\xi z} dz \quad (\star)$$

Lorsque R tend vers $+\infty$, le premier terme intégral tend vers $\mathcal{F}[f](-\xi)$: on aura donc établi notre formule si l'on montre que l'intégrale de f_ξ sur Γ_R tend vers 0 lorsque R tend vers $+\infty$.

En paramétrisant Γ_R de la façon naturelle, on obtient :

$$\left| \int_{\Gamma_R} f(z)e^{i\xi z} dz \right| = \left| \int_0^\pi iRf(Re^{i\theta})e^{i\xi Re^{i\theta}} d\theta \right| \quad (4)$$

$$\leq R \int_0^\pi |f(Re^{i\theta})| e^{-R\xi \sin(\theta)} d\theta \quad (5)$$

$$\leq 2RM(R) \int_0^{\frac{\pi}{2}} e^{-R\xi \sin(\theta)} d\theta \quad (6)$$

avec

$$M(R) := \sup_{z \in \Gamma_R} |f(z)| \quad (7)$$

et en utilisant le fait que $\sin(\theta) = \sin(\pi - \theta)$ pour modifier les bornes de l'intégrale. On utilise maintenant la classique inégalité de concavité du sinus, valable pour $\theta \in [0, \pi/2]$:

$$\sin(\theta) \geq \frac{2}{\pi} \theta \quad (8)$$

qui permet d'écrire :

$$\left| \int_{\Gamma_R} f(z)e^{i\xi z} dz \right| \leq 2RM(R) \int_0^{\frac{\pi}{2}} e^{-\frac{R\xi}{\pi} \theta} d\theta \quad (9)$$

$$\leq 2RM(R) \left[-\frac{\pi}{R\xi} e^{-\frac{R\xi}{\pi} \theta} \right]_{\theta=0}^{\theta=\pi} \quad (10)$$

$$= O_{R \rightarrow +\infty}(M(R)) \quad (11)$$

$$= o(1)_{R \rightarrow +\infty} \quad (12)$$

car on a supposé $|f(z)| \xrightarrow{|z| \rightarrow +\infty} 0$. On peut donc passer à la limite $[R \rightarrow +\infty]$ dans $(\star)$ pour obtenir :

$$\mathcal{F}[f](-\xi) = \int_{\mathbb{R}} f(x)e^{i\xi x} dx = 2i\pi \sum_{a \in A} \text{Res}(f_\xi, a) \quad (13)$$

□

Remarque : En ayant des hypothèses similaires sur le demi-plan inférieur, on pourrait obtenir la même formule pour $\xi < 0$ en sommant cette fois sur les pôles de f inclus dans le demi-plan inférieur. ♦

Donnons un exemple pour illustrer cette formule :

Application 69. Soit $a > 0$ et X une variable aléatoire suivant la loi de Cauchy de paramètre a , c'est-à-dire admettant la densité :

$$f : x \mapsto \frac{a}{\pi(x^2 + a^2)} \quad (14)$$

La fonction caractéristique de X est donnée par :

$$\varphi_X(\xi) = e^{-a|\xi|} \quad (15)$$

Démonstration. La fonction caractéristique de X est donnée par l'expression :

$$\varphi_X(\xi) := \mathbb{E}[e^{i\xi X}] = \int_{\mathbb{R}} \frac{ae^{ix\xi}}{\pi(x^2 + a^2)} dx = \mathcal{F}[f](-\xi) \quad (16)$$

La fonction f est une fraction rationnelle de degré -2 sans pôles réels : elle satisfait les hypothèses du théorème. Elle possède seulement deux pôles : ai et $-ai$. Fixons $\xi > 0$. Alors :

$$\mathcal{F}[f](\xi) = 2i\pi \operatorname{Res}(f_\xi, ai) \quad (17)$$

or f_ξ s'exprime comme un quotient de fonctions holomorphes dont le numérateur ne s'annule pas en i et dont i est un zéro d'ordre 1 du dénominateur :

$$f_\xi(z) = \frac{ae^{i\xi z}}{a^2 + z^2} \quad (18)$$

On peut donc calculer son résidu en ai en évaluant le numérateur en ai et en dérivant le dénominateur puis en l'évaluant en ai :

$$\operatorname{Res}(f_\xi, i) = \frac{[z \mapsto ae^{i\xi z}](ai)}{[z \mapsto \pi(a^2 + z^2)]'(ai)} = \frac{ae^{-a\xi}}{2a\pi i} \quad (19)$$

On trouve donc :

$$\mathcal{F}[f](-\xi) = e^{-a\xi} = e^{-a|\xi|} \quad (20)$$

Comme f est paire, sa transformée de Fourier l'est également et donc la formule reste valide pour $\xi < 0$. Enfin, comme une transformée de Fourier de fonction L^1 est continue, la formule est encore vraie pour $\xi = 0$ par continuité en 0. □

Remarque : Il est bon d'être conscient qu'on pourrait obtenir cette expression par des moyens plus élémentaires, par exemple à l'aide d'une décomposition en éléments simples. Cette application doit plus être vue comme un exemple qui illustre bien l'efficacité de cette technique sans comporter de calculs lourds, et donc est agréable à suivre lors d'un oral, mais notre méthode a

plutôt vocation à être utilisée sur des fonctions plus compliquées (surtout que malgré sa simplicité apparente, il ne faut pas oublier que le théorème de résidus reste un élément d’artillerie lourde). ♦

2.2 Différentielle du flot d'une équation différentielle autonome ★ ★ ★ ★ ★ ★

Recasages possibles : 214, 215, 220, 221

Bon. Je sais que vous l'avez remarqué. Le nombre anormal d'étoiles. Je ne vais pas y aller par quatre chemins : ce développement est très difficile. C'est de loin le plus dur que j'ai préparé. Au programme : du calcul différentiel en dimension infinie, des résolvantes de systèmes différentiels linéaires et tout ça avec à peine mieux que les étapes de calculs disponibles dans [GT98]. Le développement est très long, il est impossible de tout présenter proprement. Il faut donc se l'approprier longtemps en avance, savoir exactement ce que l'on présente et ce que l'on étudie à l'oral, et être prêt.e à riposter quand vous venez les méchantes questions de calcul diff. Dans ce poly, j'ai essayé de détailler le plus possible pour rendre aussi clair que faire se peut la nature des objets manipulés, mais ce que j'ai rédigé n'est pas représentatif de ce qu'il est humainement possible de faire en quinze minutes.

Si vous vous sentez prêt.e à l'affronter, je vous promets qu'il peut tenir en quinze minutes, et qu'il se recase très bien dans les quatre leçons qui sont annoncées. Il offre l'avantage de donner une vraie application au niveau de l'agreg du calcul diff en dimension infinie, permettant de placer toute votre leçon dans le cadre des espaces de Banach. Si comme moi vous aimez beaucoup le calcul diff, vous saurez l'apprécier.

La partie sur la résolvante est intégralement hors programme, mais une fois maîtrisée, elle remplira à merveille le plan de la leçon 221. Un ami est tombé en oral blanc sur cette leçon, a présenté ce développement et a obtenu 19,5/20 : le dev est faisable ! Notez qu'il faudra sûrement remanier l'ordre des propriétés et le choix des démonstrations pour cette leçon.

Allons, trêve de dramatisation ! On s'intéresse ici au temps d'existence des solutions d'une équation différentielle non-linéaire autonome. On va montrer un résultat de régularité sur le temps d'existence par rapport à la condition initiale en exploitant le calcul différentiel en dimension infinie. On obtiendra au passage que le flot est de classe C^1 , ainsi qu'une expression implicite de sa différentielle.

Dans ce développement, on considère $d \in \mathbb{N}^*$ et $f : \mathbb{R}^d \rightarrow \mathbb{R}^d$ un champ de vecteurs de classe C^1 . On s'intéresse à l'équation différentielle autonome :

$$x' = f(x) \quad (\star)$$

Comme f est C^1 , c'est un champ de vecteur localement lipschitzien en la variable d'état : le théorème de Cauchy-Lipschitz indique que pour toute condition initiale $x_0 \in \mathbb{R}^d$, il existe une solution locale à $(\star)$ prenant la valeur x_0 et $t = 0$, solution qu'on notera $x(-; x_0)$. Notre objectif est de démontrer qu'on peut avoir un certain contrôle sur le temps d'existence des solutions :

Théorème 70 ([GT98], II.2.5). *Soit $\alpha \in \mathbb{R}^d$. Il existe $T^* > 0$ tel que pour tout $0 < T < T^*$, il existe un voisinage V de α dans $\mathbb{R}^d$ tel que pour tout x_0 , le domaine de définition de la solution maximale à $(\star)$ issue de x_0 contienne $[-T, T]$. De plus, l'application :*

$$V \ni x_0 \mapsto x(-; x_0)|_{[-T, T]} \quad (1)$$

est de classe C^1 .

La démonstration consistera en une utilisation astucieuse du théorème des fonctions implicites permettant d'exprimer localement une solution du problème comme une fonction C^1 de la condition initiale. Introduisons quelques notations. Fixons $T > 0$ tel que $x(-, \alpha)$ soit définie au moins sur $[-T, T]$ On définit :

$$E := (C^0([-T, T], \mathbb{R}^d), \|\cdot\|_\infty) \quad (2)$$

On reformule le fait d'être solution de $(\star)$ comme le fait de satisfaire une équation intégrale de la façon suivante :

$$F : E \times \mathbb{R}^d \rightarrow E \quad (3)$$

$$(x, x_0) \mapsto \left(t \mapsto x(t) - x_0 - \int_0^t f(x(s)) ds \right)$$

Il est alors évident que $x \in E$ est solution de $(\star)$ sur $[-T, T]$ avec la condition initiale $x(0) = x_0$ si, et seulement si, $F(x, x_0) = 0$. Cherchons à appliquer le théorème des fonctions implicites à F .

Calcul des différentielles partielles de F

On va dans un premier temps calculer les deux différentielles partielles de F , dans le but d'une part de montrer que la fonction est C^1 , puis de montrer que la seconde différentielle partielle est inversible.

Différentielle en x_0

Celle-ci est aisée à calculer, car des trois composantes de F , une seule dépend de x_0 . Celle-ci correspond à l'application :

$$\begin{aligned} \mathbb{R}^d &\rightarrow E \\ x_0 &\mapsto (t \mapsto x_0) \end{aligned} \quad (4)$$

Cette application est clairement linéaire et continue, aussi :

$$\forall (x_0, h_0) \in (\mathbb{R}^d)^2, \forall x \in E, \quad d_2 F(x_0, x) \cdot h_0 = (t \mapsto h_0) \quad (5)$$

Remarquons que l'application $(x_0, x) \mapsto d_2 F(x_0, x)$ est continue (elle ne dépend même pas de x et $x_0 \dots$).

Différentielle en x

Cette partie est nettement plus délicate. Commençons par le lemme suivant :

Lemme 71. *On considère :*

$$\begin{aligned} G : E &\rightarrow E \\ x &\mapsto \left(t \mapsto \int_0^t f(x(s)) ds \right) \end{aligned} \quad (6)$$

G est différentiable, et sa différentielle est donnée par :

$$\forall (x, h) \in E^2, \quad dG(x) \cdot h = \left(t \mapsto \int_0^t df(x(s)) \cdot h(s) ds \right) \quad (7)$$

Démonstration. On est en dimension infinie, aucun espoir d'aide du côté des dérivées partielles ! Qu'à cela ne tienne, revenons à la définition de la différentielle. Prenons $(x, h) \in E^2$. On a alors, pour $t \in [-T, T]$:

$$G(x+h)(t) = \int_0^t f(x(s) + h(s)) ds \quad (8)$$

$$= \int_0^t f(x(s)) + df(x(s)) \cdot h(s) + o(\|h(s)\|) ds \quad (9)$$

$$= G(x)(t) + \int_0^t df(x(s)) \cdot h(s) ds + \int_0^t o(\|h(s)\|) ds \quad (10)$$

Ce premier calcul permet de conjecturer l'expression de la différentielle, mais ne résout pas le problème. On a deux choses à montrer : tout d'abord que le terme central est linéaire et continue en h , puis que le terme résiduel est un $o(\|h\|_\infty)$, ce qui n'est absolument pas évident, car d'un développement limité ponctuel de f , sous l'intégral qui plus est, on doit déduire un développement global.

La linéarité du terme central découle simplement de la linéarité de la différentielle et de l'intégrale. Montrons la continuité. Pour $t \in [-T, T]$, on a :

$$\left| \int_0^t df(x(s)) \cdot h(s) ds \right| \leq \left| \int_0^t \|df(x(s))\| \times \|h(s)\| ds \right| \quad (ii) \quad (11)$$

$$\leq T \sup_{s \in [-T, T]} \|df(x(s))\| \times \|h\|_\infty \quad (12)$$

Notons simplement que x étant continue et f étant C^1 , l'application $s \mapsto df(x(s))$ est continue et donc elle est bornée en norme d'opérateur sur le compact $[-T, T]$. Ainsi, le terme central est bien continu en h .

Occupons-nous maintenant du petit o . Posons :

$$\Delta(t) := \left| G(x+h)(t) - G(x)(t) - \int_0^t df(x(s)) \cdot h(s) ds \right| \quad (13)$$

Une petite réécriture des termes donne :

$$\Delta(t) = \left| \int_0^t f(x(s) + h(s)) - f(x(s)) - df(x(s)) \cdot h(s) ds \right| \quad (14)$$

Cherchons à faire apparaître une différence de différentielles sous l'intégrale :

$$f(x(s) + h(s)) - f(x(s)) = \int_0^1 df(x(s) + \theta h(s)) \cdot h(s) d\theta \quad (iii) \quad (15)$$

. Notre intégrale se réexprime alors de la façon suivante :

$$\Delta(t) = \left| \int_0^t \int_0^1 df(x(s) + \theta h(s)) \cdot h(s) - df(x(s)) \cdot h(s) d\theta ds \right| \quad (16)$$

Or l'application df est continue : par théorème de Heine, elle est uniformément continue sur les compacts. De plus, l'image de $[-T, T]$ par x est un compact de $\mathbb{R}^d$: soit $R > 0$ telle que cette image soit incluse dans $B_{\mathbb{R}^d}(0, R)$. Soit $\varepsilon > 0$. Il existe, par uniforme continuité, $\eta > 0$ tel que :

$$\forall (x_0, y_0) \in \bar{B}_{\mathbb{R}^d}(0, R)^2, \|x_0 - y_0\| \leq \eta \implies \|df(x_0) - df(y_0)\| \leq \varepsilon \quad (17)$$

(ii). Attention à ne pas oublier les valeurs absolues : t pourrait être négatif !

(iii). Il s'agit de l'intégrale de $(\theta \mapsto f(x(s) + \theta h(s)))'$ le long du segment reliant $x(s) + h(s)$ à $x(s)$

Ainsi, pour h de norme infinie inférieure à η et telle que $x + h$ reste à valeurs dans $\bar{B}_{\mathbb{R}^d}(0, R)$ (ce qui est vrai dès qu'on prend h suffisamment petite en norme infinie), on obtient la majoration :

$$\Delta(t) \leq \left| \int_0^t \int_0^1 \|df(x(s) + \theta h(s)) - df(x(s))\| \|d\theta\| \|h(s)\| ds \right| \quad (18)$$

$$\leq T\varepsilon \|h\|_\infty \quad (19)$$

et comme ε peut être pris arbitrairement petit lorsque $\|h\|_\infty \rightarrow 0$, on trouve bien $\|\Delta\|_\infty = o(\|h\|_\infty)$, ce qui achève la démonstration du lemme. $\square$

Ce lemme technique en poche, on calcule facilement la différentielle en x de F :

$$d_1 F(x, x_0) \cdot h = \left(t \mapsto h(t) - \int_0^t df(x(s)) \cdot h(s) ds \right) \quad (20)$$

Il reste toutefois à montrer que l'application $(x, x_0) \mapsto d_1 F(x, x_0)$ est continue pour obtenir de le caractère C^1 de F . Cette application ne dépendant pas de x_0 , il suffit de considérer $z \in E$ et d'observer, pour $h \in E$ et $t \in [-T, T]$:

$$|d_1 F(x, x_0) \cdot h(t) - d_1 F(z, x_0) \cdot h(t)| \leq \left| \int_0^t \|df(x(s)) - df(z(s))\| \cdot h(s) ds \right| \quad (21)$$

On exploite encore l'uniforme continuité de df sur les compacts : si $B(0, R)$ est une boule ouverte qui contient $x([-T, T])$ et tel que la distance de $x([T, T])$ au complémentaire de cette boule soit au moins 1, il existe $1 > \eta > 0$ tel que :

$$\forall (v, w) \in (\mathbb{R}^d)^2, \|v - w\| \leq \eta \implies \|df(v) - df(w)\| \leq \varepsilon \quad (22)$$

On a alors, si $\|x - z\|_\infty \leq \eta$:

$$|d_1 F(x, x_0) \cdot h(t) - d_1 F(z, x_0) \cdot h(t)| \leq \left| \int_0^T \|df(x(s)) - df(z(s))\| ds \right| \|h\|_\infty \quad (23)$$

$$\leq T\varepsilon \|h\|_\infty \quad (24)$$

En passant à la norme d'opérateur :

$$\|x - z\|_\infty \leq \eta \implies \|d_1 F(x, x_0) - d_1 F(z, x_0)\| \leq T\varepsilon \quad (25)$$

et donc $d_1 F$ est bien continue.

Inversibilité de $d_1 F(x(-; \alpha), \alpha)$

Notons $y := x(-; \alpha)|_{[-T, T]} \in E$. On va montrer que $d_1 F(y, \alpha)$ est inversible, de sorte à pouvoir appliquer le théorème des fonctions implicites. Je propose deux façons de faire, à choisir selon la leçon présentée. Notez que la seconde est plus simple, mais fournit un résultat moins puissant.

Avec la résolvante

Pour avoir un bon recasage dans la leçon sur les équations différentielles linéaires, on va calculer explicitement l'inverse de $d_1F(y, \alpha)$ en résolvant une équation linéaire à coefficients variables. Attention, cette partie ne figure pas dans la référence et nécessite la notion hors-programme de résolvante ! Soit $z \in E$. On a, pour $h \in E$:

$$d_1F(y, \alpha) \cdot h = z \iff \forall t \in [-T, T], h(t) - \int_0^t df(y(s)) \cdot h(s) ds = z(t) \quad (26)$$

Commençons par supposer que z est C^1 . Il vient alors immédiatement que l'éventuel antécédant h doit être C^1 , ce qui permet de dériver l'équation ci-dessus :

$$d_1F(y, \alpha) \cdot h = z \iff \begin{cases} \forall t \in [-T, T], h'(t) - df(y(t)) \cdot h(t) = z'(t) \\ h(0) = z(0) \end{cases} \quad (27)$$

Puisque l'application $t \mapsto df(y(t))$ est continue à valeurs dans $\mathcal{L}(\mathbb{R}^d)$, il s'agit d'un système différentiel linéaire à coefficients variables, qu'on résout à l'aide de la formule de Duhamel. Introduisons R la résolvante du système, c'est-à-dire l'application :

$$\begin{aligned} [-T, T]^2 &\rightarrow \mathcal{L}(\mathbb{R}^d) \\ (t, s) &\mapsto R(t, s) \end{aligned} \quad (28)$$

telle que $\partial_1 R(t, s) = df(y(t)) \circ R(t, s)$ et $R(t, t) = I_d$ pour tout t et s . On a alors :

$$d_1F(y, \alpha) \cdot h = z \iff \forall t \in [-T, T], h(t) = R(t, 0) \cdot z(0) + \int_0^t R(t, s) \cdot z(s) ds \quad (29)$$

On va chercher à éliminer les occurrences de z' dans cette formule, afin d'obtenir une formule qui puisse être encore valable lorsque z est seulement supposée C^0 . Pour cela, une petite intégration par parties nous donnera la solution :

$$d_1F(y, \alpha) \cdot h = z \iff \forall t \in [-T, T], h(t) = R(t, 0) \cdot z(0) + [R(t, s) \cdot z(s)]_{s=0}^{s=t} - \int_0^t \partial_2 R(t, s) \cdot z(s) ds \quad (30)$$

Par ailleurs, on sait calculer les dérivées partielles de la résolvante :

$$\partial_2 R(t, s) = -R(t, s) \circ df(y(s)) \quad (31)$$

ce qui fournit, après simplification :

$$d_1F(y, \alpha) \cdot h = z \iff \forall t \in [-T, T], h(t) = z(t) + \int_0^t R(t, s) \circ df(y(s)) \cdot z(s) ds \quad (32)$$

Montrons maintenant que cette expression est encore celle de la solution lorsque z est seulement supposée continue. Posons :

$$\begin{aligned} S : E &\rightarrow E \\ z &\mapsto \left(t \mapsto z(t) + \int_0^t R(t, s) \circ df(y(s)) \cdot z(s) ds \right) \end{aligned} \quad (33)$$

S est linéaire et continue en z , puisque :

$$\|S(z)\|_\infty \leq \left(1 + \sup_{(t,s) \in [-T,T]^2} \|R(t,s)\| \times \sup_{s \in [-T,T]} \|df(y(s))\| \right) \|z\|_\infty \quad (34)$$

où les deux supremums sont finis car $s \mapsto df(y(s))$ et $(t,s) \mapsto R(t,s)$ sont continues et évaluées sur un compact.

On exploite maintenant la densité de $C^1([-T,T], \mathbb{R}^d)$ dans E (qu'on peut observer par exemple à l'aide du théorème de Weierstraß) : restreints à l'espace des fonctions C^1 , on a l'égalité :

$$d_1 F(y, \alpha) \circ S = S \circ d_1 F(y, \alpha) = Id_{C^1} \quad (35)$$

égalité qui se transporte sur E tout entier par uniforme continuité des opérateurs manipulés. Ainsi, $d_1 F(y, \alpha)$ est inversible en tant qu'opérateur de E .

Par la complétude

Il y a un moyen plus rapide de procéder, qui oblige cela dit à contraindre plus la valeur de T . On observe :

$$d_1 F(y, \alpha) = Id_E - H \quad (36)$$

où H est l'opérateur de E défini par :

$$\forall h \in E, H(h) = \left(t \mapsto \int_0^t df(y(s)) \cdot h(s) ds \right) \quad (37)$$

qui est linéaire et continu en h . Sa norme d'opérateur est majorée de la façon suivante :

$$\|H\| \leq T \sup_{s \in [-T,T]} \|df(y(s))\| \quad (38)$$

donc quitte à remplacer T par T' tel que

$$T' < \frac{1}{\sup_{s \in [-T,T]} \|df(y(s))\|} \quad (iv) \quad (39)$$

On obtient que $\|H\| \leq 1$. Un résultat général sur les algèbres de Banach prouve alors que $d_1 F(y, \alpha)$ est inversible, et son inverse est donné par :

$$d_1 F(y, \alpha)^{-1} = \sum_{n=0}^{+\infty} H^n \quad (40)$$

Conclusion de l'étude

Le théorème des fonctions implicites s'applique : il existe V un voisinage ouvert de α dans $\mathbb{R}^d$ et $\Phi : V \rightarrow E$ une fonction C^1 telle que :

$$\forall (x, x_0) \in E \times \mathbb{R}^d, (F(x, x_0) = 0 \wedge x_0 \in V) \iff (x = \Phi(x_0) \wedge x_0 \in V) \quad (41)$$

Or par définition de la fonction F , pour $x_0 \in V$, $\Phi(x_0) = x(-; x_0)|_{[-T,T]}$. Enfin, remarquons que notre preuve reste valable pour tout $T' < T$; la preuve du théorème est donc complète !

(iv). On doit supposer que f est non constant pour écrire ça, mais avouons que dans le cas contraire, le problème n'est pas bien intéressant... Attention par ailleurs aux dépendances : T a été fixé au préalable, et on prend T' qui dépend de T , mais dans la suite on écrira seulement T ...

Annexe : la résolvante

Dans cette section, je propose d'introduire les quelques outils concernant la résolvante qui seront nécessaires si vous choisissez le chemin de preuve avec l'équation différentielle linéaire. Je conseille de mettre les propriétés en questions dans le plan. J'utilise pour toute cette partie [Ben14], II-6-2.

Je pense que pour recaser ce développement dans la 221, il est préférable d'admettre l'expression des différentielles partielles pour à la place démontrer les propriétés de la résolvante.

On considère une équation différentielle linéaire à coefficients variables :

$$y' = A(t) \cdot y(t) \quad (42)$$

où $t \mapsto A(t)$ est une application continue à valeurs dans $\mathcal{L}(R^d)$ et définie sur un intervalle I de $\mathbb{R}$.

Définition 72 (Résolvante). On appelle résolvante du système au point $s \in I$, notée $R(-; s)$ la solution (globale) de l'équation linéaire **matricielle** :

$$\begin{cases} X'(t) = A(t)X(t) \\ X(s) = I_d \end{cases} \quad (43)$$

Remarque : Bien que la variable soit ici matricielle, il s'agit bien d'une équation linéaire, et donc le théorème de Cauchy-Lipschitz linéaire s'applique. $\blacklozenge$

La résolvante sert à généraliser l'exponentielle de matrice, qui fournit la solution des systèmes linéaires autonomes, au cas des coefficients variables :

Proposition 73. La solution du problème de Cauchy :

$$\begin{cases} y'(t) = A(t)y(t) \\ y(s) = y_0 \end{cases} \quad (44)$$

est donnée par $t \mapsto R(t; s)y_0$.

Démonstration. C'est un calcul immédiat ! Posons $y : t \mapsto R(t; s)y_0$. Alors on trouve :

$$y'(t) = \partial_1 R(t; s)y_0 = A(t)R(t; s)y_0 = A(t)y(t) \quad (45)$$

par définition de la résolvante. De plus :

$$y(s) = R(s; s)y_0 = I_d y_0 = y_0 \quad (46)$$

□

Ceci va nous permettre d'utiliser la clause d'unicité du théorème de Cauchy-Lipschitz pour montrer des propriétés sur la résolvante.

Proposition 74 (Equations de la résolvante). *La résolvante satisfait les équations suivantes :*

1. $\forall \tau \in I, \forall (t, s) \in I^2, R(t; \tau)R(\tau; s) = R(t, s)$
2. *L'application $(t, s) \mapsto R(t; s)$ est continue.*
3. $\forall (t, s) \in I^2, \partial_1 R(t, s) = A(t)R(t, s), \quad \partial_2 R(t, s) = -R(t, s)A(s)$

Démonstration.

1. Soit $y_0 \in \mathbb{R}^d$. On considère, pour $s \in I$ fixé, l'application $u : t \mapsto R(t; s)y_0$. On s'intéresse alors au problème de Cauchy :

$$\begin{cases} y'(t) = A(t) \cdot y(t) \\ y(\tau) = u(\tau) \end{cases} \quad (47)$$

Il est évident que u est solution de ce problème de Cauchy. Mais on sait également, par la proposition précédente, que la solution est donnée par :

$$y : t \mapsto R(t; \tau)u(\tau) = R(t; \tau)R(\tau; s)y_0 \quad (48)$$

Ainsi, par unicité de la solution :

$$\forall t \in I, R(t; \tau)R(\tau; s)y_0 = R(t; s)y_0 \quad (49)$$

Comme c'est vrai pour tout $y_0 \in \mathbb{R}^d$, cette égalité se reporte sur les matrices :

$$\forall (s, t) \in I^2, R(t; \tau)R(\tau; s) = R(t; s) \quad (50)$$

2. Comme cas particulier de la formule précédente, on trouve :

$$\forall (t, s) \in I^2, R(t; s)R(s; t) = I_d \quad (51)$$

Comme par construction, l'application $t \mapsto R(t; s)$ est continue à s fixé, on trouve que l'application $s \mapsto R(t; s) = R(s; t)^{-1}$ est continue par continuité de l'inverse matriciel. Ainsi, pour $(t_0, s_0) \in I^2$, et $(t, s) \in I^2$, on trouve :

$$R(t; s) = R(t; t_0)R(t_0; s) \xrightarrow{(t, s) \rightarrow (t_0, s_0)} R(t_0; t_0)R(t_0; s_0) = R(t_0; s_0) \quad (52)$$

par continuité du produit matriciel. Donc la résolvante est bien continue comme application de deux variables.

3. Enfin, par différentiabilité de l'inverse matriciel, l'application $s \mapsto R(t; s) = R(s; t)^{-1}$ est différentiable. Ainsi, en partant de l'équation $I_d = R(t; s)R(s; t)$, on trouve :

$$0 = \partial_s [R(t; s)R(s; t)] = \partial_2 R(t; s)R(s; t) + R(t; s)\partial_1 R(s; t) \quad (53)$$

et par définition de la résolvante, $\partial_1 R(s; t) = A(s)R(s; t)$. Finalement :

$$\partial_2 R(t; s)R(s; t) = -R(t; s)A(s)R(s; t) \quad (54)$$

Il ne reste qu'à simplifier par $R(s; t)$ (qui est inversible) pour obtenir l'équation souhaitée !

□

On conclut enfin avec la formule de Duhamel :

Théorème 75 (Formule de Duhamel). *La solution au problème de Cauchy avec terme source :*

$$\begin{cases} y'(t) = A(t)y(t) + b(t) \\ y(t_0) = y_0 \end{cases} \quad (55)$$

est donnée par :

$$t \mapsto R(t; t_0)y_0 + \int_{t_0}^t R(t; s)b(s)ds \quad (56)$$

Travail fortement inspiré d'un cours de Karine Beauchard

fait en collaboration avec Pierre Larivé

2.3 Echantillonnage de Schannon ★★

L'intégralité du développement est extrait de [Amr], chapitre III, section 2.21. Attention il y a des erreurs dans la preuve du livre !

On va étudier ici les propriétés des fonctions L^2 dont la transformée de Fourier est à support astreint à un compact. Suite à cette étude, on obtiendra le théorème d'échantillonnage de Shannon, qui est un théorème fondateur en théorie du signal, affirmant qu'on peut reconstruire un signal à spectre borné et l'échantillonnant seulement sur un nombre discret de valeurs.

On notera L^p l'ensemble des fonctions mesurables de puissance p -ième intégrable sur $\mathbb{R}$ modulo l'égalité presque partout et L^∞ celui des fonctions essentiellement bornées modulo l'égalité presque partout. On notera par $\mathcal{F}$ la transformée de Fourier, aussi bien de L^2 dans lui-même que dans L^1 dans L^∞ . Posons :

$$BL^2 := \{f \in L^2 \mid (1 - \mathbf{1}_{[-1,1]})\mathcal{F}[f] = 0\}^{(v)} \quad (1)$$

Etude de l'espace BL^2

On va montrer ici un certain nombre de propriétés de l'espace BL^2 .

Lemme 76. *Toute fonction de BL^2 admet un représentant continu qui tend vers 0 à l'infini. On a par ailleurs une constante réelle $C^{(vi)}$ telle que :*

$$\forall f \in BL^2, \|f\|_\infty \leq C\|f\|_2 \quad (2)$$

Démonstration. Commençons par remarquer que $\mathcal{F}$ envoie L^2 sur L^2 . Ainsi, pour $f \in BL^2$, $\mathcal{F}[f]$ est une fonction L^2 à support contenu dans $[-1, 1]$. Comme cet intervalle est de mesure finie, $\mathcal{F}[f]$ est donc également L^1 sur $[-1, 1]$ et il existe une constante C telle que $\|\mathcal{F}[f]\|_1 \leq C\|\mathcal{F}[f]\|_2$, constante ne dépendant pas de $\mathcal{F}$. Par ailleurs, la formule d'inversion de Fourier donne :

$$f = \frac{1}{2\pi} \mathcal{F}[\mathcal{F}[\check{f}]] \quad (3)$$

avec $\check{f} : x \mapsto f(-x)$. Donc f s'écrit comme la transformée de Fourier d'une fonction L^1 : elle est représentée par une fonction continue. De plus, comme la transformée de Fourier de L^1 dans L^∞ est continue de norme 1, il vient :

$$\|f\|_\infty \leq \frac{1}{2\pi} \|\mathcal{F}[f]\|_1 \leq \frac{C}{2\pi} \|\mathcal{F}[f]\|_2 = C\|f\|_2 \quad (4)$$

où la dernière égalité vient de la formule de Plancherel. $\square$

Dans toute la suite, ce lemme nous permettra d'identifier les fonctions de BL^2 avec leurs représentants continus.

(v). J'ai pris quelques pincettes pour écrire proprement l'intuition qui est que $\mathcal{F}[f]$ est nulle presque partout en-dehors de $[-1, 1]$ car rigoureusement, l'évaluation d'un élément de L^2 n'a pas de sens (on peut aussi s'en sortir en choisissant un représentant).

(vi). On peut calculer explicitement sa valeur via l'inégalité de Jensen, mais ça n'est pas nécessaire pour le développement.

Proposition 77. BL^2 est un sous-espace de Hilbert de L^2 .

Démonstration. BL^2 est clairement un sous-espace vectoriel de L^2 . Il suffit donc de prouver que c'est un sous-espace fermé. Soit $(f_n) \in (BL^2)^\mathbb{N}$ une suite convergeant au sens L^2 vers une certaine fonction $f \in L^2$. Par continuité de la transformée de Fourier, on a :

$$(1 - \mathbb{1}_{[-1,1]})\mathcal{F}[f_n] \xrightarrow[n \rightarrow +\infty]{(L^2)} (1 - \mathbb{1}_{[-1,1]})\mathcal{F}[f] \quad (5)$$

Or le terme de gauche est constant égal à 0. Donc $(1 - \mathbb{1}_{[-1,1]})\mathcal{F}[f] = 0$, c'est-à-dire $f \in BL^2$ et BL^2 est un fermé de L^2 . $\square$

On va maintenant calculer une base hilbertienne de BL^2 .

Définition 78 (Sinus cardinal). On définit la fonction sinus cardinal comme :

$$\begin{aligned} \text{sinc} : \mathbb{R} &\rightarrow \mathbb{R} \\ x &\mapsto \begin{cases} \frac{\sin(x)}{x} & \text{si } x \neq 0 \\ 1 & \text{sinon} \end{cases} \end{aligned}$$

Le sinus cardinal va nous fournir une base hilbertienne de BL^2 :

Théorème 79. La famille $\left(\frac{1}{\sqrt{\pi}} \tau_{n\pi} \text{sinc} \right)_{n \in \mathbb{Z}}$ est une base hilbertienne de BL^2 , où τ_x est l'opérateur de translation par x .

Démonstration. On commence par calculer la transformée de Fourier de la fonction porte :

$$\mathcal{F}[\mathbb{1}_{[-1,1]}](\xi) = \int_{-1}^1 e^{-i\xi x} dx \quad (6)$$

$$= \left[\frac{e^{-i\xi x}}{-i\xi} \right]_{x=-1}^{x=1} \quad (7)$$

$$= 2 \text{sinc}(\xi) \quad (8)$$

Ainsi, comme la transformée de Fourier transforme les facteurs exponentiels en translations :

$$\mathcal{F} \left[x \mapsto \frac{1}{2} e^{in\pi x} \mathbb{1}_{[-1,1]} \right] (\xi) = \tau_{n\pi} \text{sinc}(\xi) \quad (9)$$

On calcule donc à l'aide de la formule de Plancherel pour k et l des entiers :

$$2\pi \langle \tau_{n\pi} \text{sinc}, \tau_{k\pi} \text{sinc} \rangle = \frac{(2\pi)^2}{4} \langle x \mapsto e^{in\pi x} \mathbb{1}_{[-1,1]}, x \mapsto e^{ik\pi x} \mathbb{1}_{[-1,1]} \rangle \quad (10)$$

$$= \pi^2 \int_{-1}^1 e^{i(n-k)\pi x} dx \quad (11)$$

$$= 2\pi^2 \delta_{k,m} \quad (12)$$

Donc cette famille est orthogonale. Il suffit de la renormaliser pour obtenir le résultat souhaité. $\square$

Théorème d'échantillonnage de Shannon

Maintenant que l'espace BL^2 est bien compris, le théorème d'échantillonnage va tomber presque comme une évidence :

Théorème 80 (Echantillonnage de Shannon). *L'application :*

$$\begin{aligned} BL^2 &\rightarrow l^2(\mathbb{Z}) \\ f &\mapsto (f(n\pi))_{n \in \mathbb{Z}} \end{aligned}$$

est une isométrie. De plus, on peut reconstruire f à partir de son échantillonnage :

$$\forall f \in BL^2, f = \sum_{n \in \mathbb{Z}} f(n\pi) \tau_{n\pi} \text{sinc} \quad (13)$$

où cette série converge au sens L^2 et uniformément.

Démonstration. De manière générale, si H est un espace de Hilbert muni d'une base hilbertienne $(e_n)_{n \in \mathbb{Z}}$, l'application $H \ni f \mapsto (\langle e_n, f \rangle)_{n \in \mathbb{Z}} \in l^2(\mathbb{Z})$ est isométrique et f est donnée comme la somme de la série $\sum \langle e_n, f \rangle e_n$ où la convergence est au sens L^2 . On applique ce résultat à l'espace BL^2 muni de sa base hilbertienne composée des translatés du sinus cardinal. Grâce au lemme 76, la convergence L^2 implique la convergence L^∞ , c'est-à-dire la convergence uniforme. Enfin, remarquons que le sinus cardinal s'annule sur les points de la forme $n\pi$ où n est un entier non nul. Ainsi, par convergence uniforme, on a :

$$\forall x \in \mathbb{R}, f(x) = \sum_{n \in \mathbb{Z}} \langle \tau_{n\pi} \text{sinc}, f \rangle \tau_{n\pi} \text{sinc}(x) \quad (14)$$

et donc en évaluant cette égalité pour $x = n\pi$, on obtient exactement :

$$\langle \tau_{n\pi} \text{sinc}, f \rangle = f(n\pi) \quad (15)$$

et le résultat s'en suit immédiatement. $\square$

2.4 Formule sommatoire de Poisson et inversion de Fourier dans L^2 ★★

Recasages possibles : 209, 235, 241, 246, 250

L'objectif de ce développement est de démontrer la formule d'inversion de Fourier dans $L^2(\mathbb{R})$ à l'aide de la formule sommatoire de Poisson. La démonstration permettra de créer un lien entre la transformée de Fourier d'une fonction et la théorie des séries de Fourier, ce qui n'était pas évident a priori.

Etant donnée une fonction mesurable sur $\mathbb{R}$, on notera lorsque ça a du sens :

$$\mathcal{F}[f] : \xi \mapsto \int_{\mathbb{R}^d} f(x) e^{-2i\pi \xi x} dx \quad (1)$$

sa transformée de Fourier^(vii). Dans le cas où f est L^2 , on notera toujours $\mathcal{F}[f]$ sa transformée de Fourier, même lorsque la formulation intégrale n'a pas de sens. Si de plus f est supposée $T\mathbb{Z}$ -périodique, on écrira :

$$\forall n \in \mathbb{Z}, c_n^{(T)}(f) := \frac{1}{T} \int_0^T f(x) e^{-\frac{2i\pi n x}{T}} dx \quad (2)$$

ses coefficients de Fourier.

On désigne par $\mathcal{S}(\mathbb{R})$ l'espace de Schwartz, c'est-à-dire l'ensemble des fonctions C^∞ de $\mathbb{R}$ dans $\mathbb{C}$ dont toutes les dérivées sont à décroissance rapide (i.e. plus vite que n'importe quel polynôme).

Formule sommatoire de Poisson [QZ], p 95

Nous allons montrer le théorème suivant :

Théorème 81 (Formule sommatoire de Poisson). *Soit $f \in \mathcal{S}(\mathbb{R})$. On a alors :*

$$\forall x \in \mathbb{R}, \sum_{n \in \mathbb{Z}} f(x + Tn) = \frac{1}{T} \sum_{n \in \mathbb{Z}} \mathcal{F}[f] \left(\frac{n}{T} \right) e^{\frac{2i\pi n x}{T}} \quad \text{(viii)} \quad (3)$$

La démonstration va reposer sur la théorie des séries de Fourier. Comme f n'est pas périodique, la première étape va consister à la transformer en une fonction périodique « la plus proche possible » de f .

Lemme 82. *Soit $T > 0$. On pose :*

$$\forall x \in \mathbb{R}, f_T(x) := \sum_{n \in \mathbb{Z}} f(x + Tn) \quad (4)$$

Alors f_T est bien définie et C^∞ .

(vii). Attention, la constante 2π dans le noyau exponentielle n'est pas la convention que j'utilise dans le reste de ce document. On la préférera pour ce développement car les expressions obtenues dans la suite seront beaucoup plus élégantes (et ça limitera les erreurs de calcul).

(viii). La somme sur $\mathbb{Z}$ est au sens suivant : les sommes prises respectivement sur les entiers positifs et négatifs convergent absolument, et la somme sur $\mathbb{Z}$ est tout simplement leur somme.

Démonstration. Prenons $\alpha \in \mathbb{N}$ et $n \in \mathbb{Z}$. Soit $M \in \mathbb{R}$ tel que :

$$\forall x \in \mathbb{R}, (x^2 + 1)|f^{(\alpha)}(x)| \leq M \quad (5)$$

qui existe car $f \in \mathcal{S}(\mathbb{R})$. On a alors, en se plaçant sur un compact $K \subset \mathbb{R}$:

$$\|x \mapsto f^{(\alpha)}(x + Tn)\|_{K,\infty} \leq \sup_{x \in K} \frac{M}{(x + Tn)^2 + 1} \sim \frac{M}{T^2 n^2} \quad (6)$$

où l'équivalent est pris lorsque $[|n| \rightarrow +\infty]$. Ainsi, la série de fonctions de terme général $x \mapsto f^{(\alpha)}(x + Tn)$ est normalement convergente, ce qui implique en particulier que f_T est bien définie et C^∞ .

On remarque alors que f_T est $T\mathbb{Z}$ -périodique en calculant $f_T(x + kT)$ avec $k \in \mathbb{Z}$ ^(ix). $\square$

Comme f_T est T -périodique, on peut calculer ses coefficients $T\mathbb{Z}$ -périodiques.

Lemme 83. Soit $k \in \mathbb{Z}$. Le k -ième coefficient de Fourier $T\mathbb{Z}$ -périodique de f_T est donné par :

$$c_k^{(T)}(f) = \frac{1}{T} \mathcal{F}[f] \left(\frac{k}{T} \right) \quad (7)$$

Démonstration. C'est un calcul direct :

$$c_k^{(T)}(f_T) := \frac{1}{T} \int_0^T f_T(x) e^{-\frac{2i\pi kx}{T}} dx \quad (8)$$

$$= \frac{1}{T} \sum_{n \in \mathbb{Z}} \int_0^T f(x + Tn) e^{-\frac{2i\pi kx}{T}} dx \quad \text{par convergence normale de la série} \quad (9)$$

$$= \frac{1}{T} \sum_{n \in \mathbb{Z}} \int_{Tn}^{T(n+1)} f(x) e^{-\frac{2i\pi kx}{T}} \overset{e^{2i\pi kn} \rightarrow 1}{\phantom{e^{-\frac{2i\pi kx}{T}}}} \quad \text{avec } x \leftarrow x - Tn \quad (10)$$

$$= \frac{1}{T} \int_{\mathbb{R}} f(x) e^{\frac{2i\pi kx}{T}} dx \quad (11)$$

$$= \frac{1}{T} \mathcal{F}[f] \left(\frac{k}{T} \right) \quad (12)$$

$\square$

Remarque : L'entièreté du développement peut se faire en prenant $T = 1$, ce qui a le mérite de simplifier les calculs et les expressions obtenues. Toutefois, en prenant T quelconque, on peut remarquer un lien très fort qui apparaît entre les théories de la transformée de Fourier et des séries de Fourier. En effet, en connaissant les coefficients de Fourier de toutes les périodisées de

(ix). Le procédé de périodisation que l'on vient d'appliquer est en fait un cas particulier de quelque chose de plus général. Si un groupe G agit sur un ensemble, on peut faire une sorte de moyenne sous l'action du groupe pour fabriquer des fonctions G -invariantes. Dans le cas typique où G est fini, on pose souvent

$$f^{(G)} : x \mapsto \frac{1}{|G|} \sum_{g \in G} f(g \cdot x)$$

La fonction $f^{(G)}$ est alors insensible à l'action du groupe G

f , on connaît intégralement sa transformée de Fourier, et réciproquement, ce qui n'apparaît pas si on se contente du cas $T = 1$. $\blacklozenge$

On est maintenant en mesure de démontrer la formule sommatoire de Poisson. En effet, f_T étant C^1 , on peut appliquer le théorème de Dirichlet :

$$\forall x \in \mathbb{R}, f_T(x) = \sum_{k \in \mathbb{Z}} c_k^{(T)}(f) e^{\frac{2i\pi kx}{T}} = \frac{1}{T} \sum_{k \in \mathbb{Z}} \mathcal{F}[f] \left(\frac{k}{T} \right) e^{\frac{2i\pi kx}{T}} \quad (13)$$

D'où la formule sommatoire de Poisson :

$$\forall x \in \mathbb{R}, \sum_{n \in \mathbb{Z}} f(x + Tn) = \frac{1}{T} \sum_{n \in \mathbb{Z}} \mathcal{F}[f] \left(\frac{n}{T} \right) e^{\frac{2i\pi nx}{T}} \quad (14)$$

Remarque : Dans le cas $x = 0$, $T = 1$, on obtient la plus élégante formule :

$$\sum_{n \in \mathbb{Z}} f(n) = \sum_{n \in \mathbb{Z}} \mathcal{F}[f](n) \quad (15)$$

$\blacklozenge$

Inversion de Fourier dans L^2 [Les], p. 254

La formule sommatoire de Poisson va nous permettre de démontrer que la transformée de Fourier est un automorphisme continu de $L^2(\mathbb{R})$. Pour cela, nous allons commencer par démontrer la formule d'inversion de Fourier dans $\mathcal{S}(\mathbb{R})$:

Théorème 84 (Inversion de Fourier dans l'espace de Schwartz). Soit $f \in \mathcal{S}(\mathbb{R})$. On a l'égalité suivante :

$$f(x) = \int_{\mathbb{R}} \mathcal{F}[f](t) e^{2i\pi xt} dt \quad (16)$$

On commence par poser, pour tout $t \in \mathbb{R}$, $\psi_t : x \mapsto f(x) e^{2i\pi tx}$. ψ_t est dans $\mathcal{S}(\mathbb{R})$ et donc on peut lui appliquer deux fois la formule sommatoire de Poisson avec $T = 1$ et $x = 0$.

$$\sum_{k \in \mathbb{Z}} \psi_t(k) = \sum_{k \in \mathbb{Z}} \mathcal{F}[\psi_t](k) = \sum_{k \in \mathbb{Z}} \mathcal{F}[\mathcal{F}[\psi_t]](k) \quad (17)$$

On peut maintenant réécrire cette expression en exploitant le fait que la transformée de Fourier transforme les translations en phase pure :

$$\mathcal{F}[s \mapsto f(s) e^{2i\pi st}](x) = \mathcal{F}[f](x - t) = \mathcal{F}[\mathcal{F}[f]](x) e^{-2i\pi xt} \quad (18)$$

et donc notre formule devient :

$$\forall t \in \mathbb{R}, \sum_{k \in \mathbb{Z}} f(k) e^{2i\pi kt} = \sum_{k \in \mathbb{Z}} \mathcal{F}[\mathcal{F}[f]](k) e^{-2i\pi kt} \quad (19)$$

Comme ces deux séries sont normalement convergentes (la transformée de Fourier stabilisant l'espace de Schwartz), on peut intégrer en t et intervertir la somme et l'intégrale :

$$\sum_{k \in \mathbb{Z}} f(k) \int_0^1 e^{2i\pi kt} dt = \sum_{k \in \mathbb{Z}} \mathcal{F}[\mathcal{F}[f]](k) \int_0^1 e^{-2i\pi kt} dt \quad (20)$$

C'est-à-dire :

$$f(0) = \mathcal{F}[\mathcal{F}[f]](0) \quad (21)$$

Cette formule étant vraie pour toute fonction de $\mathcal{S}(\mathbb{R})$, on peut l'appliquer à la fonction $\zeta_x : t \mapsto f(x+t)$ pour un x quelconque, ce qui donne :

$$f(x) = \zeta_x(0) \quad (22)$$

$$= \mathcal{F}[\mathcal{F}[\zeta_x]](0) \quad (23)$$

$$= \int_{\mathbb{R}} \mathcal{F}[\zeta_x](t) dt \quad (24)$$

$$= \int_{\mathbb{R}} \mathcal{F}[f](t) e^{2i\pi xt} dt \quad (25)$$

C'est exactement la formule d'inversion de Fourier !

Remarque : On peut s'étonner qu'il n'y ait pas de constante devant l'intégrale. Toutefois, la définition que nous avons choisie de la transformée de Fourier n'est pas standard. On retrouve la formulation standard en faisant le changement de variable $u \leftarrow 2\pi t$ dans l'intégrale, et alors on retrouve l'expression usuelle avec $\frac{1}{2\pi}$ en facteur. $\blacklozenge$

On est maintenant en mesure de prouver l'important théorème suivant :

Théorème 85. *La transformée de Fourier est un automorphisme continu de $L^2(\mathbb{R})$.*

En effet, on peut maintenant poser l'opérateur :

$$\begin{aligned} \mathcal{G} : \mathcal{S}(\mathbb{R}) &\rightarrow L^2(\mathbb{R}) \\ f &\mapsto (x \mapsto \int_{\mathbb{R}} \mathcal{F}[f](\xi) e^{i\xi x} d\xi) \end{aligned}$$

Cet opérateur est continu pour la topologie L^2 , car il s'exprime facilement en fonction de la transformée de Fourier. Il existe donc un unique opérateur linéaire continu, que nous noterons toujours $\mathcal{G}$, allant de L^2 dans lui-même, et coïncidant avec l'expression intégrale ci-dessus sur $\mathcal{S}(\mathbb{R})$ par densité de $\mathcal{S}(\mathbb{R})$ dans $L^2(\mathbb{R})$. Alors, les opérateurs $\mathcal{F} \circ \mathcal{G}$ et $\mathcal{G} \circ \mathcal{F}$ sont linéaires continus de $L^2(\mathbb{R})$ dans lui-même et coïncident avec l'identité de $L^2(\mathbb{R})$ sur la partie dense qu'est $\mathcal{S}(\mathbb{R})$. Nécessairement, $\mathcal{F}$ est inversible et son inverse est donné par $\mathcal{G}$.

Remarque : La démonstration ci-dessus ne s'adapte pas à L^1 . En effet, la transformée de Fourier sur L^1 est à valeurs dans un sous-espace de L^∞ , et l'opérateur $\mathcal{F}$ y est continue pour les topologies L^1 au départ et L^∞ à l'arrivée. L'opérateur inverse $\mathcal{G}$ n'est alors pas trivialement continu (et d'ailleurs, il est déjà difficile de le définir correctement car $\mathcal{S}(\mathbb{R})$ n'est pas dense dans L^∞ (il est l'est dans $C_0^0(\mathbb{R})$ l'espace des fonctions continues qui tendent vers 0 à l'infini, qui contient l'image de la transformée de Fourier L^1)). $\blacklozenge$

2.5 Lemme de la grenouille ★

On sait, de façon générale, qu'une suite convergente à valeurs dans un espace métrique a nécessairement ses accroissements qui tendent vers 0. La réciproque est fausse, bien qu'on pourrait naïvement avoir l'idée du contraire : il suffit par exemple de considérer la suite $(\sqrt{n})$. En revanche, on va voir ici que si les accroissements d'une suite tendent vers 0, les valeurs d'adhérence de celle-ci sont contraintes sous certaines hypothèses de compacité.

On se placera dans toute la suite dans le cadre abstrait des espaces métriques. Pour pouvoir présenter ce développement dans la leçon 223, il est nécessaire de remplacer l'espace abstrait par $\mathbb{C}$, mais la preuve est identique et je ne crois pas qu'elle puisse être simplifiée.

On retrouvera ce développement dans [IP24].

Soit (X, d) un espace métrique compact et $x = (x_n)$ une suite d'éléments de X . On notera $\text{Adh}(x)$ l'ensemble (non-vide par compacité) des valeurs d'adhérence de x .

Théorème 86 (Grenouille). *On suppose que $d(x_{n+1}, x_n) \rightarrow 0$. Alors $\text{Adh}(x)$ est connexe.*

Démonstration. On va supposer par l'absurde que $\text{Adh}(x)$ n'est pas connexe. Dans ce cas, il existe F_1 et F_2 des fermés disjoints de $\text{Adh}(x)$ tels que $\text{Adh}(x) = F_1 \sqcup F_2$. Remarquons :

$$\text{Adh}(x) = \bigcap_{p \in \mathbb{N}} \overline{\{x_n, n \geq p\}} \quad (1)$$

C'est en particulier un fermé de X comme intersection de fermés, donc un compact, et par conséquent, les (F_i) sont compacts^(x). Notons :

$$\alpha := d(F_1, F_2) := \inf_{(y, z) \in F_1 \times F_2} d(y, z) \quad (2)$$

Comme les (F_i) sont fermés dans X , $\alpha > 0$.

Pour comprendre ce résultat, nous allons raconter une petite histoire qui permettra au passage de comprendre pourquoi on appelle ce résultat « lemme de la grenouille ». Imaginez une petite grenouille mathématique qui saute de positions en positions pendant un temps infini. (x_n) modélisera la suite des positions qu'elle occupe. Comme bondir une infinité de fois est fatigant, notre petite grenouille a de moins en moins d'énergie, ce qui se traduit par le fait que $d(x_{n+1}, x_n)$ tende vers 0. Notre petite grenouille n'éprouve d'intérêt que pour ses deux nénuphars métriques préférés : F_1 et F_2 . Or ceux-ci ne se touchent pas ! Le temps passant, notre grenouille n'aura plus assez d'énergie pour relier deux voisinages de F_1 et F_2 en un seul saut, et par conséquent, elle tombera à l'eau. Une infinité de fois. On montrera qu'alors sa trajectoire admet une valeur d'adhérence dans l'eau.

Cette histoire tragique contient toute la fin de la démonstration. On va « gonfler » un peu F_1 et F_2 pour se ménager une marge de manœuvre : posons

$$U_i := \bigcup_{z \in F_i} B(z, \alpha/3) \quad (3)$$

(x). Il ne faut pas négliger cette étape. La négation de la connexité de $\text{Adh}(x)$ permet de prendre les (F_i) fermés dans $\text{Adh}(x)$, pas dans X a priori. La fermeture de $\text{Adh}(x)$ dans X est essentielle ici (considérez par exemple l'ensemble $[-1, 1] \setminus \{0\}$ dans le compact $[-1, 1]$, qui n'est pas connexe et réunion des fermés relatifs $[-1, 0[$ et $]0, 1]$, deux parties qui ne sont pas fermées dans $[-1, 1]$).

U_i est un voisinage ouvert de F_i comme union d'ouverts. On remarque que ces voisinages sont distants :

$$d(U_1, U_2) \geq \frac{\alpha}{3} > 0 \quad (4)$$

On considère l'eau dans laquelle va tomber la grenouille :

$$K := X \setminus (U_1 \cup U_2) \neq \emptyset \quad (5)$$

K est compact comme compact privé d'un ouvert. Tout ce qu'il nous reste à faire, c'est de montrer que (x_n) a une infinité de ses termes dans K . Nécessairement, par compacité de K , x devra alors admettre une valeur d'adhérence dans K . Fixons $n_0 \in \mathbb{N}$ tel que :

$$\forall n \geq n_0, d(x_{n+1}, x_n) \leq \alpha/4 \quad (6)$$

et $N \geq n_0$. Comme F_1 est composé de valeurs d'adhérences de x , il existe $n_1 \geq N$ tel que $x_{n_1} \in U_1$. Le raisonnement valant également pour F_2 , on peut poser :

$$n_2 := \inf\{n \geq n_1 \mid x_n \in U_2\} \quad (7)$$

Mais alors, par construction, $x_{n_2-1} \notin U_2$. Par ailleurs, $d(x_{n_2-1}, x_{n_2}) \leq \alpha/4 < d(U_1, U_2)$. Donc $x_{n_2-1} \notin U_1$. Il suit que $x_{n_2-1} \in K$. Finalement :

$$\forall N \in \mathbb{N}, \exists n^* \geq N : x_{n^*} \in K \quad (8)$$

Il existe donc une suite extraite de x à valeurs dans le compact K , et par conséquent, $\text{Adh}(x) \cap K \neq \emptyset$. Absurde ! $\text{Adh}(x)$ est donc connexe. $\square$

Une application

Grâce à ce résultat, on va pouvoir donner une condition suffisante pour que la convergence vers 0 des écarts implique la convergence de la suite :

Proposition 87. *Soit $f : [0, 1] \rightarrow [0, 1]$ une application continue. Considérons (x_n) la suite définie par la relation de récurrence :*

$$\begin{cases} x_0 \in [0, 1] \\ x_{n+1} = f(x_n) \end{cases} \quad (9)$$

(x_n) converge si, et seulement si, $|x_{n+1} - x_n| \rightarrow 0$.

Démonstration. Le sens direct est immédiat. On suppose donc que $|x_{n+1} - x_n|$ tend vers 0. Dans ce cas, par le lemme de la grenouille, $\text{Adh}((x_n))$ est un connexe de $[0, 1]$, c'est-à-dire un intervalle. Supposons par l'absurde que (x_n) ne converge pas. Alors, comme $[0, 1]$ est compact, $\text{Adh}((x_n))$ est non réduit à un singleton, et donc d'intérieur non-vidé.

Pour $x^* \in \text{Adh}((x_n))$, il existe $\varphi : \mathbb{N} \rightarrow \mathbb{N}$ strictement croissante telle que $x_{\varphi(n)} \rightarrow x^*$. On observe alors, par continuité de f :

$$|f(x_{\varphi(n)}) - x_{\varphi(n)}| \rightarrow |f(x^*) - x^*| \quad (10)$$

Comme les écarts de (x_n) tendent vers 0, on déduit $f(x^*) = x^*$. Autrement dit, toutes les valeurs d'adhérence de (x_n) sont des points fixes de f .

Si maintenant x^* est un point intérieur de (x_n) , la suite $x_{\varphi(n)}$ est ultimement dans $\text{Adh}((x_n))$. Mais puisque $\text{Adh}((x_n))$ est composé de points fixes de f , on en déduit que la suite (x_n) est en réalité ultimement constante égale à x^* et $\text{Adh}(x_n)$ est réduit à x^* . Absurde! La suite (x_n) est donc convergente. $\square$

Remarque : Dans cette application, on voit apparaître le fait que dans $\mathbb{R}$, on obtient mieux que la connexité de $\text{Adh}((x_n))$: cet ensemble est convexe. On pourrait se demander si c'est toujours le cas lorsque (X, d) est un espace vectoriel normé. Il n'en est rien, comme on peut le voir par exemple avec la suite complexe $e^{i \log(n+1)}$ dont l'ensemble des valeurs d'adhérence est le cercle unité. $\blacklozenge$

Bonus : une idée d'application

Avouons-le, le corollaire précédent est amusant, mais il est difficile d'en imaginer une utilisation concrète. Toutefois, le lemme de la grenouille trouve bien des applications pratiques : on peut lire par exemple dans [IP24] qu'il sert à l'élaboration de la méthode de Jacobi, qui est une méthode numérique permettant d'obtenir les valeurs propres des matrices symétriques. Sans forcément en savoir plus, je pense qu'il est bon d'avoir cette application en tête pour le passage à l'oral

2.6 Holomorphie sous l'intégrale et problème des moments

★ ★ ★

Recasages possibles : 235, 239, 245, 261

Dans ce petit développement fait maison, on va mêler analyse complexe et probabilités pour répondre dans un cas particulier au problème des moments, qui consiste à savoir si une loi de probabilité sur $\mathbb{R}$ est caractérisée par ses moments. On va commencer par démontrer le théorème d'holomorphie sous l'intégrale, avec une preuve qu'on trouve assez rarement et qui exploite très fortement la théorie des fonctions holomorphes. On l'utilisera ensuite pour étudier les fonctions caractéristiques des lois à moments gaussiens.

J'ai eu l'occasion de présenter ce développement en oral blanc et il a été beaucoup apprécié pour son originalité. Par ailleurs, l'application que je propose justifie de se placer dans le cadre abstrait des espaces mesurés σ -finis. Si comme moi vous aimez la théorie de la mesure, ce genre de résultats vous permettra de faire des plans beaucoup plus abstraits sans qu'on vous reproche une abstraction gratuite sans application à la clef.

Remarquez qu'il est un peu exagéré de considérer la référence que j'ai indiqué comme une référence, puisque d'après les auteurs, le théorème d'holomorphie sous l'intégrale est une « conséquence immédiate du critère de Morera »...

Soit $(X, \mathfrak{A}, \mu)$ un espace probabilisé σ -fini. On note $L^1(\mu)$ l'ensemble des fonctions mesurables et intégrables de X dans $\mathbb{C}$, modulo l'égalité presque partout. Pour U un ouvert de $\mathbb{C}$, on désigne par $H(U)$ l'ensemble des fonctions holomorphes définies sur U .

Holomorphie sous l'intégrale

La première étape de notre développement consiste à démontrer le théorème suivant :

Théorème 88 (Holomorphie sous l'intégrale ([QQ17])). Soit U un ouvert de $\mathbb{C}$. On considère $f : X \times U \rightarrow \mathbb{C}$ un noyau satisfaisant les conditions suivantes :

1. $\forall z \in U, f(-, z) \in L^1(\mu)$
2. $\forall x \in X, f(x, -) \in \mathcal{H}(U)$
3. $\forall K \subset U$ compact, $\exists g_K \in L^1(\mu) : \forall (x, z) \in X \times K, |f(x, z)| \leq g_K(x)$

Alors la fonction :

$$F : z \mapsto \int_X f(x, z) d\mu(x) \quad (1)$$

est définie et holomorphe sur U . De plus, on peut intervertir la dérivation et l'intégrale :

$$\forall n \in \mathbb{N}, \forall z \in U, F^{(n)}(z) = \int_X \partial_z^n f(x, z) d\mu(x) \quad (2)$$

Remarque : Le cadre de travail est à la fois large et restreint. On énonce fréquemment ce théorème en remplaçant l'espace mesuré $(X, \mathfrak{A}, \mu)$ par $\mathbb{R}^d$ muni de la mesure de Lebesgue, et c'est effectivement le cadre le plus fréquent d'utilisation. Mais dans la suite du développement, on va vouloir l'utiliser pour des intégrales contre des mesures, sur $\mathbb{R}$ certes, mais beaucoup plus louches que la mesure de Lebesgue, d'où la nécessité de travailler dans un cadre un peu élargi. MAIS le

théorème d'holomorphic sous l'intégrale reste vrai en retirant l'hypothèse de σ -finitude sur X . La démonstration que je vais proposer nécessite le caractère σ -fini car on va se servir du théorème de Fubini. La démonstration générale n'est pas beaucoup plus difficile, mais celle que je propose utilise beaucoup plus d'outils d'analyse complexe, c'est donc celle que j'ai préféré développer. ♦

Démonstration. Remarquons tout d'abord que F est définie sur U en vertu de l'hypothèse 1.

Notre démonstration va utiliser le critère de Morera pour caractériser les fonctions holomorphes : F est holomorphe sur U si, et seulement si, elle est continue et son intégrale sur le bord de tout triangle plein contenu dans U est nulle.

Continuité de F : pas de mystère ici, on utilise le théorème de continuité sous le signe intégrale, dont les hypothèses sont strictement moins fortes que celles qu'on a faites sur f . Il s'applique donc sans rajouter quoi que ce soit et F est continue.

Intégrale sur le bord des triangles : Soit $T \subset U$ un triangle plein inclus dans U ^(xi). On va commencer par montrer que F est intégrable sur ∂T . Remarquons que F étant continue, elle est intégrable sur ∂T , qui est un compact de U . On a :

$$\int_{\partial T} F(z) = \int_{\partial T} \int_X f(x, z) d\mu(x) dz = \int_0^1 \int_X f(x, \gamma(t)) \gamma'(t) d\mu(x) dt \quad (3)$$

où $\gamma : [0, 1] \rightarrow U$ est une paramétrisation continue et C^1 par morceaux de ∂T . Passons donc au module sous les deux intégrales **dans la dernière expression** ^(xii) :

$$\int_0^1 \int_X |f(x, \gamma(t)) \gamma'(t)| d\mu(x) dt \leq \int_0^1 \int_X g_T(x) |\gamma'(t)| d\mu(x) dt \leq \|g_T\|_1 L(\gamma) < +\infty \quad (4)$$

avec $L(\gamma)$ la longueur de l'arc γ . Ainsi, $(x, t) \mapsto f(x, \gamma(t)) \gamma'(t)$ est intégrable sur $X \times [0, 1]$ et on peut utiliser le théorème de Fubini :

$$\int_{\partial T} F(z) dz = \int_X \underbrace{\int_{\partial T} f(x, z) dz}_{=0} d\mu(x) = 0 \quad (5)$$

car $f(x, -)$ est holomorphe

Par le lemme de Morera, F est holomorphe sur U .

On va utiliser la formule de Cauchy pour montrer que la dérivation commute avec l'intégrale. Soit $n \in \mathbb{N}$, soit $z_0 \in U$ et soit $r > 0$ tel que $\bar{B}(z_0, r) \subset U$. Comme F est holomorphe sur U

$$F^{(n)}(z_0) = \frac{n!}{2i\pi} \int_{\partial B(z_0, r)} \frac{F(z)}{(z - z_0)^{n+1}} dz \quad (6)$$

$$= \frac{n!}{2i\pi} \int_{\partial B(0, r)} \int_X \frac{f(x, z)}{(z - z_0)^{n+1}} d\mu(x) dz \quad (7)$$

$$= \int_X \frac{n!}{2i\pi} \int_{\partial B(0, r)} \frac{f(x, z)}{(z - z_0)^{n+1}} dz d\mu(x) \quad (8)$$

$$= \int_X \partial_z^n f(x, z) d\mu(x) \quad (9)$$

en appliquant une nouvelle fois le théorème de Fubini (la démonstration est identique que celle déjà faite plus haut) et en utilisant encore le fait que $f(x, -)$ est holomorphe pour tout x . □

(xi). Attention, l'intérieur du triangle doit être inclus dans U , sinon l'intégrale sur son bord peut ne pas faire 0.
(xii). Attention à ne pas se précipiter et à calculer $\int_{\partial T} \int_X |f(x, z)| d\mu(x) dz$, car la mesure dz est une mesure complexe ! Il faut d'abord faire apparaître une paramétrisation γ pour intégrer contre une mesure positive.

Le problème des moments

On va maintenant relier des propriétés d'holomorphicité au problème des moments en probabilités :

Application 89. (Problème des moments : intégrabilité gaussienne) Soient X et Y deux variables aléatoires réelles dont les lois respectives admettent un moment gaussien, c'est-à-dire qu'il existe $\varepsilon > 0$ tel que $e^{\varepsilon X^2}$ est intégrable^(xiii). On suppose :

$$\forall n \in \mathbb{N}, \mathbb{E}[X^n] = \mathbb{E}[Y^n] \quad (10)$$

Alors X et Y ont même loi.

Remarque : En réalité, on peut seulement supposer que X a un moment gaussien. Si Y a les mêmes moments que X , elle a alors automatiquement un moment gaussien (voir [QQ17] dans la section « Applications en probabilités »). ♦

Pour cela, on va travailler sur les fonctions caractéristiques en montrant le lemme suivant :

Lemme 90. Soit μ une mesure de probabilité sur $\mathbb{R}$ telle que $x \mapsto e^{\varepsilon x^2}$ est μ -intégrable pour un certain $\varepsilon > 0$. L'application :

$$\varphi : z \mapsto \int_{\mathbb{R}} e^{izx} d\mu(x) \quad (11)$$

est holomorphe sur $\mathbb{C}$.

Démonstration. Tout d'abord, notons que μ étant une mesure finie, $(\mathbb{R}, \mathcal{B}(\mathbb{R}), \mu)$ est un espace mesuré σ -fini. Soit $R > 0$. Pour z de module au plus R , on a :

$$|e^{izx}| = e^{-x \operatorname{Im}(z)} \leq e^{xR} \quad (12)$$

Or la fonction $x \mapsto e^{xR}$ est intégrable contre μ car dominée par $x \mapsto e^{\varepsilon x^2}$. Comme c'est vrai pour tout R et que $z \mapsto e^{izx}$ est holomorphe sur $\mathbb{C}$, le théorème d'holomorphicité sous l'intégrale indique que φ est holomorphe sur $\mathbb{C}$. □

On est alors en mesure de démontrer l'application : par le lemme, les fonctions caractéristiques

(xiii). Cette hypothèse implique que X et Y admettent des moments à tout ordre.

de X et Y se prolongent sur $\mathbb{C}$ en deux fonctions entières. Donc :

$$\forall z \in \mathbb{C}, \varphi_X(z) = \sum_{n=0}^{+\infty} \frac{\varphi_X^{(n)}(0)}{n!} z^n \quad (13)$$

$$= \sum_{n=0}^{+\infty} i^n \frac{\mathbb{E}[X^n]}{n!} z^n \quad (14)$$

$$= \sum_{n=0}^{+\infty} i^n \frac{\mathbb{E}[Y^n]}{n!} z^n \quad (15)$$

$$= \varphi_Y(z) \quad (16)$$

Ainsi, X et Y ont mêmes fonctions caractéristiques : puisque celle-ci caractérisent la loi, elles ont même loi !

Application 91. La loi normale est caractérisée par ses moments parmi les lois à moments gaussiens.

Remarque : De façon plus générale, on peut montrer qu'il suffit que la fonction caractéristique soit développable en série entière au voisinage de 0 pour qu'elles le soient au voisinage de tout réel (cf le poly de Clémentine), et donc il existe un prolongement holomorphe des fonctions caractéristiques à un voisinage de $\mathbb{R}$. Dans ce cas, le fait que X et Y aient les mêmes moments implique qu'elles coïncident sur un voisinage de 0, et donc sur tout un voisinage de $\mathbb{R}$ par connexité. ♦

2.7 Optimisation dans un espace de Hilbert et théorème de Banach-Alaoglu ★★ ★

Recasages possibles : 203, 205, 213, 219, 229, 253

L'objectif de ce développement est de généraliser un résultat d'optimisation de fonctionnelles convexes, valables dans les espaces de dimension finie par des arguments très simples de compacité, au cadres des espaces de Hilbert. La preuve exploitera un de mes théorèmes d'analyse fonctionnelle préférés : le théorème de Banach-Alaoglu, qu'on démontrera en première partie. C'est un théorème qui est vrai dans un cadre beaucoup plus général en travaillant avec la notion de convergence faible-*, mais qui s'obtient dans le cadre hilbertien en utilisant seulement des notions du programme d'agreg.

L'intégralité du développement se retrouve dans [IP24]. J'ai légèrement modifié l'organisation de la preuve pour mettre plus en valeur le théorème de Banach-Alaoglu, mais tous les ingrédients se retrouvent dans leur livre.

J'ai ajouté une annexe à la fin du document qui fournit un contre-exemple lorsque on sort du cadre hilbertien, et qui donne une application concrète du théorème (qui que très hors-programme pour l'agrégation).

Soit H un espace de Hilbert. On considère $J : H \rightarrow \mathbb{R}$ une fonction convexe, continue et coercive^(xiv). Le but est de démontrer :

Théorème 92. Il existe $x^* \in H$ tel que $J(x^*) = \inf_{x \in H} J(x)$.

Lorsque H est de dimension finie, ce théorème s'obtient très simplement de la façon suivante : on choisit $y \in \mathbb{R}$ une valeur prise par J , et $R > 0$ tel que $J(x) \geq y$ dès que $\|x\|$ est de norme plus grande que R , donné par coercivité de J . Comme J est continue, elle atteint un minimum sur le compact $B(0, R)$ et ce minimum est global de par le choix de R . Mais en dimension infinie, c'est une autre paire de manches !

Théorème de Banach-Alaoglu

La clef de la démonstration sera le théorème suivant :

Théorème 93 (Banach-Alaoglu). Soit $(x_n) \in H^{\mathbb{N}}$ une suite bornée de vecteurs de H . Il existe $\psi : \mathbb{N} \rightarrow \mathbb{N}$ strictement croissante et $x^* \in H$ tel que :

$$\forall y \in H, \langle x_{\psi(n)}, y \rangle \xrightarrow{n \rightarrow +\infty} \langle x^*, y \rangle \quad (1)$$

On dit alors que (x_n) converge faiblement vers x^* , ce qu'on note $x_n \rightharpoonup x^*$.

Démonstration. On considère le sous-espace vectoriel fermé de H :

$$F := \overline{\text{Vect}(x_n \mid n \in \mathbb{N})} \quad (2)$$

(xiv). C'est-à-dire telle que $|J(x)| \xrightarrow{\|x\| \rightarrow +\infty} +\infty$.

F est un espace séparable, c'est-à-dire qu'il contient une partie dénombrable et dense (par exemple le $\mathbb{Q}$ -espace vectoriel engendré par les (x_n)). Soit (y_n) une telle suite dense. On va alors construire itérativement une suite d'extractrice de la façon suivante :

- Comme (x_n) est bornée, la suite $(\langle x_n, y_0 \rangle)$ est bornée dans $\mathbb{R}$ par l'inégalité de Cauchy-Schwarz. Ainsi, il existe φ_0 une extractrice telle que $(\langle x_{\varphi_0(n)}, y_0 \rangle)$ converge.
- Si on suppose contruite une famille $\varphi_0, \dots, \varphi_k$ d'extractrices telle que pour tout $i \in \llbracket 0, k \rrbracket$, la suite $(\langle x_{\varphi_0 \circ \dots \circ \varphi_i(n)}, y_i \rangle)$ converge dans $\mathbb{R}$, il suffit de prendre φ_{k+1} une extractrice telle que $(\langle x_{\varphi_0 \circ \dots \circ \varphi_{k+1}(n)}, y_{k+1} \rangle)$ converge pour obtenir le terme suivant.

Sans surprise, on procède à une extraction diagonale : posons

$$\begin{aligned} \psi : \mathbb{N} &\rightarrow \mathbb{N} \\ k &\mapsto \varphi_0 \circ \dots \circ \varphi_k(k) \end{aligned} \quad (3)$$

Pour tout $k \in \mathbb{N}$, la suite $(\langle x_{\psi(n)}, y_k \rangle)$ est alors convergente dans $\mathbb{R}$.

Montrons maintenant que $(\langle x_{\psi(n)}, y \rangle)$ converge encore pour tout vecteur $y \in F$ en montrant que cette suite est de Cauchy. Soit $y \in F$. Par densité :

$$\exists k \in \mathbb{N} : \|y - y_k\| \leq \frac{\varepsilon}{2 \sup_{n \in \mathbb{N}} \|x_n\|} \quad (4)$$

De plus, grâce au travail précédent :

$$\exists N \in \mathbb{N}, : \forall n, p \geq N, |\langle x_{\psi(n)} - x_{\psi(p)}, y_k \rangle| \leq \frac{\varepsilon}{2} \quad (5)$$

Ainsi, pour tout $n, p \geq N$, on a :

$$|\langle x_{\psi(n)} - x_{\psi(p)}, y \rangle| \leq |\langle x_{\psi(n)} - x_{\psi(p)}, y - y_k \rangle| + |\langle x_{\psi(n)} - x_{\psi(p)}, y_k \rangle| \quad (6)$$

$$\leq (\|x_{\psi(n)}\| + \|x_{\psi(p)}\|) \|y - y_k\| + |\langle x_{\psi(n)} - x_{\psi(p)}, y_k \rangle| \quad (7)$$

$$\leq \varepsilon \quad (8)$$

Cette suite est donc de Cauchy dans $\mathbb{R}$: elle converge. On peut donc légitimement poser :

$$\begin{aligned} L : F &\rightarrow \mathbb{R} \\ y &\mapsto \lim_{n \rightarrow +\infty} \langle x_{\psi(n)}, y \rangle \end{aligned} \quad (9)$$

Par linéarité et unicité de la limite, ainsi que par bilinéarité du produit scalaire, L est linéaire. Montrons qu'elle est continue. Pour $y \in F$ et $n \in \mathbb{N}$, on a :

$$|\langle x_{\psi(n)}, y \rangle| \leq \sup_{n \in \mathbb{N}} \|x_n\| \|y\| \quad (10)$$

d'où en passant à la limite :

$$|L(y)| \leq \sup_{n \in \mathbb{N}} \|x_n\| \|y\| \quad (11)$$

et L est bien continue. On peut donc appliquer le théorème de représentation de Riesz :

$$\exists ! x^* \in F : L = \langle x^*, - \rangle \quad (12)$$

ce qui est exactement dire que $(x_{\psi(n)})$ converge faiblement vers x^* **dans F** . Montrons que x^* est encore valeur d'adhérence au sens faible dans H . Pour cela, on exploite la décomposition orthogonale :

$$H = F \oplus F^\perp \quad (13)$$

Pour $y \in H$, il existe une unique écriture $y = y_F + y_\perp$ avec $y_F \in F$ et $y_\perp \in F^\perp$. On a alors :

$$\langle x_{\psi(n)}, y \rangle = \langle x_{\psi(n)}, y_F \rangle \xrightarrow{n \rightarrow +\infty} \langle x^*, y_F \rangle = \langle x^*, y \rangle \quad (14)$$

car $x^* \in F$. Donc la convergence faible vers x^* a en fait lieu dans H . $\square$

Minimisation de J

On va maintenant exploiter le théorème de Banach-Alaoglu pour prouver que J admet un minimum sur H . Notons $\alpha := \inf_{x \in H} J(x) \in \mathbb{R} \cup \{-\infty\}$. Par définition de l'infimum, il existe (x_n) une suite d'éléments de H telle que :

$$J(x_n) \xrightarrow{n \rightarrow +\infty} \alpha \quad (15)$$

Par coercivité de J , (x_n) est bornée, puisque si (x_n) admettait une sous-suite dont la norme diverge, $J(x_n)$ tendrait vers $+\infty$ le long de cette sous-suite. Ainsi, le théorème de Banach-Alaoglu fournit un vecteur $x^* \in H$ et une extractrice φ telle que $x_{\varphi(n)} \rightharpoonup x^*$. Montrons que $J(x^*) = \alpha$ (et donc en particulier que α est fini).

On considère pour cela $\beta > \alpha$ et on définit :

$$C_\beta := J^{-1}(]-\infty, \beta]) \quad (16)$$

On va montrer que $x^* \in C_\beta$, c'est-à-dire $J(x^*) \leq \beta$. Si on a cette appartenance pour tout β , on aura nécessairement que $J(x^*) = \alpha$.

Comme J est continue, C_β est fermé dans H . De plus, par convexité de J , C_β est convexe. Enfin, la définition de β implique que C_β est non-vide. On peut donc appliquer le théorème de projection sur un convexe fermé : notons p_β la projection sur C_β .

Puisque $J(x_n) \rightarrow \alpha$, il existe un rang n_0 à partir duquel les termes de la suite (x_n) sont dans C_β . Par caractérisation de l'angle obtus, on a alors :

$$\Re(\langle x^* - p_\beta(x^*), x_{\varphi(n)} - p_\beta(x^*) \rangle) \leq 0 \quad (17)$$

Or en passant à la limite $[n \rightarrow +\infty]$, par convergence faible de $(x_{\varphi(n)})$:

$$\Re(\langle x^* - p_\beta(x^*), x^* - p_\beta(x^*) \rangle) = \|x^* - p_\beta(x^*)\|^2 \leq 0 \quad (18)$$

et donc $x^* = p_\beta(x^*)$: on a bien prouvé $J(x^*) \leq \beta$! Ainsi, $J(x^*)$ ne peut-être que α , et α est fini ! $\square$

Annexe : à quoi ça sert ?

Bonne question en effet, le genre qu'il faut s'être posée avant de passer à l'oral. Le fait est que ça n'est pas évident de trouver une « vraie » application de ce résultat.

Remarquons tout d'abord que le résultat qu'on vient de montrer est plus fort que le théorème de Lax-Milgram : on s'attend à pouvoir montrer le même genre de résultat avec. En l'occurrence, il peut servir à montrer que certains types d'équations différentielles non linéaires admettent des solutions. Cela dit, le cadre habituel de travail pour faire des équations diff est l'ensemble des fonctions de classe C^k selon la régularité dont on a besoin, et... ça n'est pas un espace de Hilbert ! On va devoir introduire certains espaces qui vont permettre de chercher des solutions au sens faible :

Définition 94 (Espace de Sobolev ([HL09], chap X.1)). Pour $k \in \mathbb{N}$, on note $H^k(I)$ l'ensemble des fonctions de classe L^2 sur I dont les k premières dérivées au sens des distributions sont représentées par une fonction L^2 . En d'autres termes :

$$H^k(I) = \left\{ f \in L^2 \mid \forall j \in \llbracket 0, k \rrbracket, \exists g_j \in L^2(I) : \forall \varphi \in C_c^\infty(I), (-1)^k \int_I f \varphi^{(j)} = \int_I g_j \varphi \right\} \quad (19)$$

On muni $H^k(I)$ du produit scalaire :

$$\langle f, g \rangle_{H^k} := \sum_{j=0}^k \int_I f^{(j)} g^{(j)} \quad (20)$$

C'est alors un espace de Hilbert.

On note H_0^k l'adhérence des fonctions C^∞ à support compact dans H^k .

Ces espaces offrent un cadre de travail pour étudier des équations différentielles tout en profitant des outils hilbertiens. Notre théorème permet notamment de montrer :

Théorème 95. Soient f une fonction L^2 et p un entier non nul. Il existe une unique fonction $u \in H^2(\rceil 0, 1 \rceil) \cap H_0^1(\rceil 0, 1 \rceil)$ telle que :

$$\begin{cases} (u'^2)' + u|u|^{p-1} = f \\ u(0) = u(1) = 0 \end{cases} \quad (21)$$

La démonstration est TRÈS longue et technique (elle est faite intégralement dans la version de Thomas Courant). Je pense que c'est complètement déraisonnable de vouloir placer ce genre de choses dans son plan (sauf si vous une machine de guerre sur ces sujets), mais à mon avis c'est toujours bon d'avoir une vague idée d'à quoi sert le résultat.

Bonus : un contre-exemple dans le cas non hilbertien

Pour pouvoir affirmer que ce théorème a de la valeur, il m'a semblé important de trouver un contre-exemple lorsque H n'est plus supposé être un espace de Hilbert (mettons un espace de Banach par exemple). Je vais donner ici un exemple de fonctionnelle convexe, continue et coercive qui n'atteint pas de minimum, mais il est très possible que mon exemple soit beaucoup trop compliqué pour ce que c'est. Si jamais vous trouvez quelque chose de beaucoup plus simple, n'hésitez pas à me contacter !

On se place dans l'espace E des suites réelles qui convergent vers 0 qu'on munit de la norme infinie. Soit (a_n) une suite l^1 de réels strictement positifs dont la somme vaut 1. On pose alors :

$$\begin{aligned} \varphi : E &\rightarrow \mathbb{R} \\ (x_n) &\mapsto \sum_{n=0}^{+\infty} a_n x_n \end{aligned} \quad (22)$$

φ est une forme linéaire continue de norme d'opérateur 1. Remarquons que φ n'atteint pas sa norme d'opérateur, dans le sens où pour toute suite non nulle (x_n) de E , $|\varphi((x_n))| < \|(x_n)\|_\infty$, puisque l'égalité est atteinte si, et seulement si, (x_n) est constante (et n'est donc pas un élément de E). On pose alors :

$$\begin{aligned} J : E &\rightarrow \mathbb{R} \\ (x_n) &\mapsto \max\{1, \|(x_n)\|_\infty\} + |1 - \varphi((x_n))| \end{aligned} \quad (23)$$

J est clairement continue et coercive. Montrons qu'elle est de plus convexe. Soit $(x, y) \in E^2$ et soit $\lambda \in [0, 1]$. On a :

$$|1 - \varphi(\lambda x + (1 - \lambda)y)| \leq \lambda |1 - \varphi(x)| + (1 - \lambda) |1 - \varphi(y)| \quad (24)$$

par inégalité triangulaire. De plus :

$$\max\{1, \|\lambda x + (1 - \lambda)y\|_\infty\} \leq \max\{1, \lambda \|x\|_\infty + (1 - \lambda) \|y\|_\infty\} \quad (25)$$

$$\leq \max\{1, \lambda \|x\|_\infty\} + \max\{1, (1 - \lambda) \|y\|_\infty\} \quad (26)$$

$$\leq \lambda \max\{1, \|x\|_\infty\} + (1 - \lambda) \max\{1, \|y\|_\infty\} \quad (27)$$

Il est clair que $J \geq 1$ et que si $((x_n^{(p)}))_{p \in \mathbb{N}}$ est une suite de suites de E de norme 1 telle que $\varphi((x_n^{(p)})) \xrightarrow{p \rightarrow +\infty} 1$, alors $J((x_n^{(p)})) \rightarrow 1$. Donc 1 est l'infimum de J . Mais pourtant :

$$J(x_n) = 1 \iff \max\{1, \|(x_n)\|_\infty\} = 1 \wedge |1 - \varphi((x_n))| = 0 \quad (28)$$

$$\iff \|(x_n)\|_\infty \leq 1 \wedge \varphi((x_n)) = 1 \quad (29)$$

ce qui est impossible puisque φ n'atteint pas sa norme. Donc J n'a pas de minimum sur E .

2.8 Probabilité(s) invariante(s) d'une chaîne de Markov ★★

★★

Recasages possibles : 203, 206, 226, 262, 264

Dans ce développement, on va s'intéresser à l'existence et à l'unicité de mesures de probabilités invariantes pour des chaînes de Markov à espace d'états fini en fonction des propriétés de leurs noyaux de transition. C'est un développement qui s'adresse majoritairement aux options A. Tous les outils qui y seront utiles sont au programme de l'option, mais une bonne partie est hors-programme au tronc commun.

Soit $\mathcal{X}$ un ensemble fini de cardinal d . On considère $X = (X_n)$ une chaîne de Markov à valeurs dans $\mathcal{X}$ de noyau de transition $P \in \mathbb{M}_d(\mathbb{R})$ défini par :

$$\forall (x, y) \in \mathcal{X}^2, P(x, y) = \mathbb{P}(X_1 = y \mid X_0 = x) \quad (1)$$

Si π est une mesure de probabilité sur $\mathcal{X}$, représentée par un vecteur ligne, on note (X_n^π) la chaîne de noyau P et de loi initiale π . Dans ce cas, la loi de X_n^π est donnée par $\pi \cdot P^n$. Une telle probabilité est dite P -invariante si pour tout $n \in \mathbb{N}$, la loi de X_n^π est π . En d'autres termes, π est un vecteur propre à gauche de P à composantes positives et tel que $\|\pi\|_1 = 1$.

On va s'intéresser ici à des conditions d'existence et d'unicité d'une telle probabilité invariante, et à une manière algorithmique de les approcher.

Existence d'une probabilité invariante

Lorsque l'espace d'état est fini, il existe toujours une mesure **de probabilité** invariante^(xv).

Théorème 96 (Existence d'une mesure de probabilité invariante ([BE07], 2.2.1)).

Toute chaîne de Markov X sur $\mathcal{X}$ admet une probabilité invariante π .

Démonstration. Comme $\mathcal{X}$ est fini, une mesure de probabilité sur $\mathcal{X}$ est complètement décrite par les probabilités atomiques, aussi l'ensemble $\mathcal{P}(\mathcal{X})$ des mesures de probabilité sur $\mathcal{X}$ s'identifie à un sous-ensemble de $\mathbb{R}^{\mathcal{X}}$, décrit de la façon suivante :

$$\mathcal{P}(\mathcal{X}) \cong \{p : \mathcal{X} \rightarrow \mathbb{R}_+ \mid \sum_{x \in \mathcal{X}} p(x) = 1\} \quad (2)$$

Il s'agit d'un sous-ensemble convexe, fermé et borné de $\mathbb{R}^{\mathcal{X}}$. Comme cet espace est de dimension finie, c'est un compact.

Soit $\mu \in \mathcal{P}(\mathcal{X})$. On considère la suite $(\tilde{\mu}_n)_{n \in \mathbb{N}^*}$ définie par :

$$\tilde{\mu}_n := \frac{1}{n} \sum_{k=1}^n \mu \cdot P^k \quad (3)$$

Comme P est une matrice stochastique, $\mu \cdot P^k \in \mathcal{P}(\mathcal{X})$ pour tout k . De plus, par convexité de $\mathcal{P}(\mathcal{X})$, cette suite est à valeurs dans $\mathcal{P}(\mathcal{X})$. Elle est en particulier bornée, et donc :

$$\tilde{\mu}_n \cdot P - \mu = \frac{1}{n} (\mu \cdot P^{n+1} - \mu) \xrightarrow{n \rightarrow +\infty} 0 \quad (\star)$$

(xv). Il est important de bien préciser mesure de probabilité, car la mesure nulle est toujours une mesure invariante, même lorsque l'espace d'états n'est pas fini.

De plus, par compacité de $\mathcal{P}(\mathcal{X})$, $(\tilde{\mu}_n)$ admet une valeur d'adhérence notée π . En passant à la limite le long d'une sous-suite adéquate dans $(\star)$, on trouve :

$$\pi \cdot P - \pi = 0 \quad (4)$$

et donc π est une probabilité invariante. $\square$

Remarque : L'hypothèse de finitude de $\mathcal{X}$ est cruciale ici : le raisonnement précédent exploite la description des compacts de $\mathcal{P}(\mathcal{X})$, or cet espace est de dimension infinie si $\mathcal{X}$ est infini. On trouve facilement des contre-exemples lorsque notre chaîne de Markov est à espace d'états infini. Par exemple, la suite (déterministe) $X_{n+1} = X_n + 1$, $X_0 = 0$ est une chaîne de Markov définie sur $\mathbb{N}$, dont la seule mesure invariante est la mesure nulle (qui n'est pas une mesure de probabilité). $\blacklozenge$

Condition suffisante d'unicité ([BE07], 2.2.2)

On va maintenant s'intéresser à une condition simple permettant d'affirmer l'unicité de la mesure invariante.

Définition 97 (Chaîne irréductible). Une chaîne de Markov (X_n) est dite irréductible lorsque :

$$\forall (x, y) \in \mathcal{X}^2, \exists k \in \mathbb{N}^* : \mathbb{P}(X_k = y, X_0 = x) > 0 \quad (5)$$

Remarque : De façon équivalente, une chaîne de Markov est irréductible si son graphe est fortement connexe. On peut encore caractériser l'irréductibilité d'une chaîne par son noyau de transition :

$$\forall (x, y) \in \mathcal{X}^2, \exists k \in \mathbb{N}^* : P^k(x, y) > 0 \quad (6)$$

On dit dans ce cas que la matrice P est irréductible. $\blacklozenge$

Théorème 98. Soit (X_n) une chaîne de Markov à valeurs dans $\mathcal{X}$. Si (X_n) est irréductible, elle admet une unique mesure de probabilité invariante π . De plus :

$$\forall x \in \mathcal{X}, \pi(\{x\}) > 0 \quad (7)$$

lorsque cette écriture a du sens.

Avant de passer à la démonstration du théorème, on va introduire la notion de fonction P -invariante.

Définition 99 (Fonction P -invariante). Une fonction $f : \mathcal{X} \rightarrow \mathbb{R}$ est dite P -invariante si :

$$\forall x \in \mathcal{X}, \mathbb{E}[f(X_1) \mid X_0 = x] = f(x) \quad (8)$$

Démonstration. Soit π une mesure de probabilités invariante pour $\mathcal{X}$, fournie par le théorème précédent. On va commencer par montrer que π charge tous les points de $\mathcal{X}$. Soit $x \in \mathcal{X}$.

Comme π est une mesure de probabilité, il existe $y \in \mathcal{X}$ tel que $\pi(\{y\}) > 0$. Par irréductibilité de la chaîne de Markov, il existe $k > 0$ tel que $P^k(y, x) > 0$. Ainsi, en exploitant le fait que pour tout $n \in \mathbb{N}$, X_n^π est de loi π , on a :

$$\pi(\{x\}) = \mathbb{P}(X_k^\pi = x) \quad (9)$$

$$= \sum_{z \in \mathcal{X}} \mathbb{P}(X_k^\pi = x \mid X_0^\pi = z) \mathbb{P}(X_0^\pi = z) \quad (10)$$

$$= \sum_{z \in \mathcal{X}} P^k(z, x) \pi(\{z\}) \quad (11)$$

$$\geq P^k(y, x) \pi(\{y\}) \quad (12)$$

$$> 0 \quad (13)$$

Donc π charge tous les points.

On va maintenant travailler sur les fonctions P -invariantes :

Lemme 100. *Les fonctions P -invariantes sont constantes.*

Démonstration. Soit f une fonction P -invariante. On va montrer que la quantité :

$$\mathcal{E}(f) := \mathbb{E}[(f(X_1^\pi) - f(X_0^\pi))^2] \quad (14)$$

est nulle. Pour cela, on développe directement :

$$\mathcal{E}(f) = \mathbb{E}[f(X_1^\pi)^2] + \mathbb{E}[f(X_0^\pi)^2] - 2\mathbb{E}[f(X_1^\pi)f(X_0^\pi)] \quad (15)$$

Mais X_0^π et X_1^π ont même loi (π). Donc :

$$\mathbb{E}[f(X_1^\pi)^2] = \mathbb{E}[f(X_0^\pi)^2] \quad (16)$$

Pour simplifier le terme croisé, on va utiliser l'espérance conditionnelle ^(xvi) :

$$\mathbb{E}[f(X_1^\pi)f(X_0^\pi)] = \mathbb{E}[\mathbb{E}[f(X_1^\pi)f(X_0^\pi) \mid X_0^\pi]] \quad (17)$$

$$= \mathbb{E}[\mathbb{E}[f(X_1^\pi) \mid X_0^\pi]f(X_0^\pi)] \quad (18)$$

$$= \mathbb{E}[f(X_0^\pi)^2] \quad (19)$$

en utilisant le fait que $f(X_0^\pi)$ est une variable aléatoire $\sigma(X_0^\pi)$ -mesurable et bornée et que f est P -invariante. En combinant ces égalités, on trouve bien :

$$\mathcal{E}(f) = 0 \quad (20)$$

De cela, on déduit que f est constante, puisqu'on a :

$$\mathbb{E}[(f(X_1^\pi) - f(X_0^\pi))^2] = \sum_{(x,y) \in \mathcal{X}^2} (f(y) - f(x))^2 \mathbb{P}(X_1^\pi = y, X_0^\pi = x) \quad (21)$$

$$= \sum_{(x,y) \in \mathcal{X}^2} (f(y) - f(x))^2 \mathbb{P}(X_1^\pi = y \mid X_0^\pi = x) \mathbb{P}(X_0^\pi = x) \quad (22)$$

$$= \sum_{(x,y) \in \mathcal{X}^2} (f(y) - f(x))^2 P(x, y) \pi(\{x\}) \quad (23)$$

(xvi). Quitte à faire joujou avec l'artillerie de l'option A, autant y aller à fond...

et comme π charge tous les points, on a :

$$\forall (x, y) \in \mathcal{X}^2, (f(x) - f(y))^2 P(x, y) = 0 \quad (24)$$

Or par irréductibilité de P , il existe nécessairement $x_1, \dots, x_k$ un chemin dans $\mathcal{X}$ telle que $P(x_i, x_{i+1}) > 0$ pour tout i ^(xvii). On a alors $f(x_i) = f(x_{i+1})$ pour tout i et donc f est constante sur $\mathcal{X}$. $\square$

Soit ν une autre mesure de probabilité invariante. On considère alors la fonction :

$$\begin{aligned} f : \mathcal{X} &\rightarrow \mathbb{R} \\ x &\mapsto \frac{\nu(\{x\})}{\pi(\{x\})} \end{aligned} \quad (25)$$

bien définie puisque ν charge tous les points en tant que mesure de probabilité invariante. On définit la matrice :

$$\forall (x, y) \in \mathcal{X}^2, Q(x, y) := \frac{\pi(\{y\})}{\pi(\{x\})} P(y, x) \quad (26)$$

C'est une matrice stochastique. En effet, à x fixé :

$$\sum_{y \in \mathcal{X}} Q(x, y) = \frac{1}{\pi(\{x\})} \sum_{y \in \mathcal{X}} \pi(\{y\}) P(y, x) = 1 \quad (27)$$

car π est P -invariante. Elle est de plus irréductible puisque π charge tous les points et P est irréductible. Montrons que f est Q -invariante. Soit (Y_n) une chaîne de Markov de noyau de transition Q et soit $x \in \mathcal{X}$. On a :

$$\mathbb{E}[f(Y_1) \mid Y_0 = x] = \sum_{y \in \mathcal{X}} f(y) \mathbb{P}(Y_1 = y \mid Y_0 = x) \quad (28)$$

$$= \sum_{y \in \mathcal{X}} \frac{\nu(\{y\})}{\pi(\{y\})} Q(x, y) \quad (29)$$

$$= \frac{1}{\pi(\{x\})} \sum_{y \in \mathcal{X}} \nu(\{y\}) P(y, x) \quad (30)$$

$$= f(x) \quad (31)$$

car ν est P -invariante. Ainsi, f est une fonction constante, et nécessairement, $\nu = \pi$. $\square$

Remarque : Si on retire l'hypothèse d'irréductibilité, on n'a pas l'unicité de la mesure de probabilité invariante. Par exemple, on peut considérer le modèle de Wright-Fisher : sur $\mathcal{X} = \llbracket 1, d \rrbracket$, on définit une chaîne de Markov (X_n) de sorte à ce que la loi de X_{n+1} sachant X_n soit $\mathcal{B}(d, X_n/d)$ (ou de façon plus formelle, $\mathbb{P}(X_{n+1} = y \mid X_n = x) = \binom{d}{y} (x/d)^y (1 - x/d)^{d-y}$). Alors δ_1 et δ_d sont deux probabilités invariantes pour cette chaîne. $\blacklozenge$

(xvii). Cela se voit bien en utilisant le fait que le graphe associé à (X_n) est fortement connexe, et donc il existe un chemin parcourant tous ses sommets, qui sont les éléments de $\mathcal{X}$. Comme deux sommets x et y sont reliés si, et seulement si $P(x, y) > 0$, on peut construire un tel chemin.

Application 101. Si (X_n) est irréductible, alors pour toute mesure de probabilité μ sur $\mathcal{X}$, la suite géométrique $(\mu \cdot P^n)$ converge au sens de Cesàro vers la probabilité invariante.

Démonstration. On a vu dans la première partie que la suite :

$$\tilde{\mu}_n := \frac{1}{n} \sum_{k=1}^n \mu \cdot P^k \tag{32}$$

est à valeur dans le compact $\mathcal{P}(\mathcal{X})$ et que ses valeurs d'adhérence sont toutes des probabilités invariantes. Puisqu'il n'existe qu'une seule telle probabilité invariante, $(\tilde{\mu}_n)$ n'a qu'une valeur d'adhérence, donc elle converge par compacité. $\square$

2.9 Quadrature de Gauss-Legendre ★

Recasages possibles : 209, 218, 236

Pour tout $n \in \mathbb{N}$, on note $\mathbb{R}_n[X]$ l'ensemble des polynômes à coefficients réels de degré au plus n . On confondra allégrement les polynômes avec leurs fonctions polynomiales associées. On pose

$$\mathcal{J} : L^1([-1, 1], \mathbb{R}) \rightarrow \mathbb{R} \quad (1)$$

$$f \mapsto \int_{-1}^1 f(x) dx \quad (2)$$

L'objectif de ce développement est d'étudier une méthode de quadrature simple, dite de Gauss-Legendre, c'est-à-dire une façon d'approcher la forme linéaire $\mathcal{J}$ par une combinaison pondérée d'évaluation de f en certains points bien choisis de $[-1, 1]$, c'est-à-dire de pouvoir écrire :

$$\forall f \in C^0, \int_{-1}^1 f(x) dx \approx \sum_{i=1}^p \omega_i f(\alpha_i) \quad (3)$$

Tout le jeu d'une telle méthode est de choisir de bons paramètres p , (ω_i) et (α_i) . La méthode de Gauss-Legendre va exploiter une famille de polynômes orthogonaux : les polynômes de Legendre.

Préliminaires : Polynômes de Legendre

Cette partie sert de prémices au développement, mais il est préférable de l'indiquer dans le plan pour gagner du temps, le développement étant bien trop long sinon.

On muni l'espace $\mathbb{R}[X]$ du produit scalaire :

$$\langle P, Q \rangle := \int_{-1}^1 P(x)Q(x) dx \quad (4)$$

On définit $(P_n)_{n \in \mathbb{N}}$ comme l'orthonormalisé de Gram-Schmidt de la base canonique pour ce produit scalaire. On appelle P_n le n -ième polynôme de Legendre.

On va donner quelques unes des propriétés intéressantes des polynômes de Legendre, qu'on utilisera par la suite. Fixons n un entier naturel.

Lemme 102. P_n est un polynôme de degré n admettant exactement n -racines distinctes dans $[-1, 1]$.

Démonstration. L'orthonormalisé de Gram-Schmidt de la base canonique de $\mathbb{R}[X]$ est toujours échelonné en degrés, donc $\deg(P_n) = n$. Etant donné un polynôme P , on note $Z(P)$ l'ensemble de ses racines réelles. Posons alors :

$$\Lambda = \{\alpha \in Z(P_n) \cap [-1, 1] \mid \alpha \text{ est de multiplicité impaire}\} \quad (5)$$

Par l'absurde, on suppose $|\Lambda| < n$. On pose alors :

$$Q_n := \prod_{\alpha \in \Lambda} (X - \alpha) \quad (6)$$

On a donc par hypothèse $\deg(Q_n) < n$, et par construction, $P_n \perp Q_n$. Notons, pour $\alpha \in \Lambda$, m_α la multiplicité de α comme racine de P_n . On a alors :

$$0 = \langle P_n, Q_n \rangle \quad (7)$$

$$= \int_{-1}^1 \prod_{\alpha \in \Lambda} (x - \alpha)^{m_\alpha+1} \prod_{\alpha \in Z(P_n) \setminus \Lambda} (x - \alpha)^{m_\alpha} dx \quad (8)$$

Or le produit $\prod_{\alpha \in Z(P_n) \setminus \Lambda} (x - \alpha)^{m_\alpha}$ est de signe constant sur $[-1, 1]$, car pour $\alpha \in \Lambda \setminus Z(P_n)$, ou bien m_α est pair, ou bien $\alpha \notin [-1, 1]$ et donc le terme $(x - \alpha)^{m_\alpha}$ est de signe constant sur $[-1, 1]$. Donc cette intégrale ne peut pas être nulle ! On a une contradiction, ce qui montre le lemme. $\square$

Dans toute la suite, on notera $\alpha_1 < \dots < \alpha_n$ les n racines de P_n .

Lemme 103. *Pour tout $n \in \mathbb{N}$, le coefficient dominant de P_n est strictement positif.*

Démonstration. Cela découle juste de la construction par Gram-Schmidt : on peut écrire :

$$P_n = \frac{1}{C} X^n - Q \quad (9)$$

où Q est un polynôme de $\mathbb{R}_{n-1}[X]$ et C est la norme d'un polynôme non nul, donc est strictement positive. $\square$

Une méthode de quadrature

Les racines des polynômes de Legendre nous fournissent un candidat pour les paramètres de notre quadrature : on propose comme approximation :

$$\forall f \in C^0, \mathcal{J}(f) \approx \sum_{i=1}^n \omega_i f(\alpha_i) \quad (10)$$

Le tout va être de trouver des poids ω_i qui maximisent l'ordre de la méthode.

Remarquons que les (α_i) étant distincts, la famille (ev_{α_i}) d'évaluation en les (α_i) est une base de $\mathbb{R}_{n-1}[X]^*$. On note (L_i) sa base antéduale. L_i est alors le polynôme d'interpolation de Lagrange de degré $n-1$ valant 1 en α_i et 0 sur les autres racines de P_n .

Proposition 104. *La méthode de quadrature proposée est d'ordre au moins $n-1$ si, et seulement si, $\omega_i = I(L_i)$.*

Démonstration.

On peut écrire :

$$\mathcal{J} = \sum_{i=1}^n \lambda_i ev_{\alpha_i} \quad (11)$$

où les λ_i sont des réels. Puisque (L_i) est une base de $\mathbb{R}_{n-1}[X]$, on a l'équivalence :

$$\forall P \in \mathbb{R}_{n-1}[X], \mathcal{J}(P) = \sum_{i=1}^n \omega_i P(\alpha_i) \iff \forall i \in \llbracket 1, n \rrbracket, \mathcal{J}(L_i) = \sum_{i=1}^n \omega_i L_i(\alpha_i) \quad (12)$$

$$\iff \forall i \in \llbracket 1, n \rrbracket, \int_{-1}^1 L_i(x) dx = \omega_i \quad (13)$$

ce qui donne l'équivalence souhaitée. $\square$

Définition 105 (Quadrature de Gauss-Legendre). On définit :

$$\begin{aligned} \mathcal{Q}_n : C^0(\mathbb{R}) &\rightarrow \mathbb{R} \\ f &\mapsto \sum_{i=1}^n I(L_i) f(\alpha_i) \end{aligned}$$

Pour toute la suite, on fixera $(\omega_i) = (I(L_i))$.

Un petit miracle se produit ici : on a choisit nos paramètres de sorte à ce que $\mathcal{Q}_n$ définisse une méthode de quadrature d'ordre $n - 1$, mais on a d'un seul coup un résultat bien meilleur :

Proposition 106. La quadrature de Gauss-Legendre est d'ordre $2n - 1$.

Démonstration. Soit $P \in \mathbb{R}_{2n-1}[X]$. On effectue la division euclidienne de P par P_n :

$$P = P_n Q + R, \quad \deg(R) < n \quad (14)$$

On a alors $\deg(Q) \leq n - 1$ et donc $P_n \perp Q$. Il vient :

$$\mathcal{J}(P) = \langle P_n, Q \rangle + \mathcal{J}(R) = \mathcal{Q}_n(R) = \mathcal{Q}_n(P) \quad (15)$$

De plus, $(P_n)^2$ est un polynôme de degré $2n$ partageant toutes ses racines avec P , donc $\mathcal{Q}_n(P) = 0$, mais $\mathcal{J}(Q) \neq 0$ car $(P_n)^2$ induit une fonction positive, continue et non nulle sur $[-1, 1]$. La méthode est donc bel et bien d'ordre $2n - 1$! $\square$

Corollaire 107. Les ω_i sont strictement positifs.

Démonstration. Considérons L_i le polynôme d'interpolation de Lagrange tel que $L_i(\alpha_j) = \delta_{i,j}$. Ce polynôme est de degré $n - 1$ et donc L_i^2 est de degré $2n - 2 \leq 2n - 1$. Ainsi :

$$\omega_i = \mathcal{Q}_n(L_i^2) = \mathcal{J}(L_i^2) > 0 \quad (16)$$

$\square$

Majoration de l'erreur commise

Bien évidemment, l'intérêt de la méthode de quadrature n'est pas de calculer l'intégrale de fonction polynômiales, qu'on aurait pu calculer explicitement à l'aide du théorème fondamental de l'analyse. On va dans cette dernière partie donner une majoration de l'erreur commise dans l'approximation $\mathcal{J} \approx \mathcal{Q}_n$, sous réserve que f soit suffisamment régulière.

Théorème 108. Soit f une fonction de classe C^{2n} sur $] -1, 1[$. Notons $M := \sup_{x \in]-1, 1[} |f^{(2n)}(x)|$. On a alors la majoration :

$$|\mathcal{J}(f) - \mathcal{Q}_n(f)| \leq 4 \frac{M}{(2n+1)!} \quad (17)$$

Démonstration. On écrit la formule de Taylor-Lagrange à l'ordre $2n-1$ pour f :

$$\forall x \in]-1, 1[, \exists c_x \in]-1, 1[: f(x) = \underbrace{\sum_{k=0}^{2n-1} \frac{f^{(k)}(0)}{k!} x^k}_{=:T(x)} + \underbrace{\frac{f^{(2n)}(c_x)}{(2n)!} x^{2n}}_{=:u(x)} \quad (18)$$

On a alors :

$$|\mathcal{J}(f) - \mathcal{Q}_n(f)| = |(\underbrace{\mathcal{J}(T) - \mathcal{Q}_n(T)}_{=0}) + (\mathcal{J}(u) - \mathcal{Q}_n(u))| \leq |\mathcal{J}(u)| + |\mathcal{Q}_n(u)| \quad (19)$$

On va majorer séparément ces deux termes.

Une majoration brutale fournit :

$$|\mathcal{J}(u)| \leq 2 \frac{M}{(2n+1)!} \quad (20)$$

Pour l'autre terme, c'est un peu plus délicat. On commence par exploiter le fait que nos poids ω_i sont strictement positifs, de sorte que :

$$|\mathcal{Q}_n(u)| \leq \frac{1}{(2n)!} \sum_{i=1}^n |f^{(2n)}(c_{\alpha_i})| \omega_i \alpha_i^{2n} \leq \frac{M}{(2n)!} \underbrace{\mathcal{Q}_n(X^{2n})}_{>0} \quad (21)$$

Par ailleurs, en écrivant deux divisions euclidiennes successives par P_n , on obtient :

$$X^{2n} = P_n D_1 + S_1 = \frac{1}{\gamma^2} P_n^2 + P_n S_2 + S_1 \quad (22)$$

où $\gamma > 0$ est le coefficient dominant de P_n . En effet, D_1 est de degré exactement n et son coefficient dominant est nécessairement $1/\gamma$, donc sa division euclidienne par P_n s'écrit nécessairement :

$$D_1 = \frac{1}{\gamma^2} P_n + S_2 \quad (23)$$

avec $\deg(S_1) < 2n$. Cela permet d'écrire :

$$\mathcal{Q}_n(X^{2n}) = \mathcal{Q}_n(S_1) \quad (P_n(\alpha_i) = 0 \text{ pour tout } i) \quad (24)$$

$$= \mathcal{J}(S_1) \quad (\text{la méthode est d'ordre } 2n - 1) \quad (25)$$

$$= \mathcal{J}(X^{2n}) - \mathcal{J}(\gamma^{-2}P_n^2) - \mathcal{J}(P_n S_2) \quad (26)$$

$$= \mathcal{J}(X^{2n}) - \underbrace{\mathcal{J}(\gamma^{-2}P_n^2)}_{>0} - \underbrace{\mathcal{Q}_n(P_n S_2)}_{=0} \quad (\deg(P_n S_2) \leq 2n - 1) \quad (27)$$

$$\leq \mathcal{J}(X^{2n}) \quad (28)$$

$$\leq \frac{2}{2n + 1} \quad (29)$$

En recombinaison toutes ces inégalités, on obtient le résultat souhaité! $\square$

Remarque : Il ne faut pas se laisser tromper par ce majorant factoriel idyllique. La constante au numérateur est la norme infinie de $f^{(2n)}$, qui ne se contrôle pas à partir de la norme de f . Dans le cas particulier où f est analytique, les formules de Cauchy permettent de majorer cette norme, mais le majorant fait apparaître $(2n)! \|f\|_\infty \dots$. Ainsi notre merveilleux dénominateur factoriel s'est changé en $1/(2n + 1)$. Il s'agit cela dit du pire cas. Cette méthode ayant réellement été utilisée pour calculer des intégrales de façon approchée, il est à parier qu'en pratique, la croissance est bien plus rapide. $\blacklozenge$

2.10 Séries de Fourier divergentes par le théorème de Banach-Steinhaus ★★☆☆

Recasages possibles : 205, 230, 235, 246

Le théorème de Dirichlet affirme qu'une fonction continue, sous certaines hypothèses plus fortes de régularité, peut être reconstruite à partir de ses coefficients de Fourier. Si l'on retire ces hypothèses, c'est le théorème de Fejér qui prend le relais, mais la série de Fourier ne converge plus qu'au sens de Césaro. On est en droit de se demander si une fonction continue est toujours limite de sa série de Fourier, au moins au sens de la convergence simple. On va démontrer ici que la réponse est non : il existe des fonctions continues dont la série de Fourier diverge en 0. Pire, ces fonctions sont «nombreuses» : elles forment un ensemble dense dans les fonctions continues, ce qui peut paraître étonnant car il n'est pas évident de fabriquer explicitement un contre-exemple !

Notre démonstration va exploiter un corollaire classique du lemme de Baire^(xviii) : le théorème de Banach-Steinhaus.

Notons $C_{2\pi}^0$ l'ensemble des fonctions continues 2π -périodiques à valeurs dans $\mathbb{C}$ qu'on munit de la norme uniforme. On désigne par e_n la fonction $t \mapsto e^{int}$. Étant données deux fonctions $f, g \in C_{2\pi}^0$, on définit leur convolution 2π -périodique par :

$$f * g : x \mapsto \frac{1}{2\pi} \int_0^{2\pi} f(t)g(x-t)dt \quad (1)$$

de sorte que les coefficients de Fourier de f soient donnés par :

$$c_n(f) = f * e_n(0) \quad (2)$$

Théorème de Banach-Steinhaus ([Goub], Annexe A, exercice 7)

Théorème 109 (Banach-Steinhaus^(xix)). Soient $(E, \|\cdot\|_E)$ un espace de Banach et $(F, \|\cdot\|)$ un espace vectoriel normé^(xx). On considère une famille quelconque d'applications linéaires continues $(\varphi_i)_{i \in I} \in \mathcal{L}_c(E, F)^I$. Alors l'un est vérifié :

1. La famille $(\|\varphi_i\|)_{i \in I}$ est bornée dans $\mathbb{R}$.
2. L'ensemble $\{x \in E \mid (\|\varphi_i(x)\|_F)_{i \in I} \text{ est non-bornée dans } \mathbb{R}\}$ est dense dans E .

Posons, pour $n \in \mathbb{N}$:

$$\Omega_n := \{x \in E \mid \forall i \in I, \|\varphi_i(x)\|_F \leq n\} \quad (3)$$

Ω_n est un fermé de E , car on a :

$$\Omega_n = \bigcap_{i \in I} \|\varphi_i(\cdot)\|_F^{-1}([0, n]) \quad (4)$$

On distingue alors deux cas :

(xviii). Oui, j'ai bien écrit le lemme de Baire. C'est la vie !

(xx). J'énonce une version un peu plus forte du théorème que ce qu'on trouve dans le Gourdon, mais la démonstration est quasiment identique.

(xx). Certains livres (le Brézis par exemple) demandent aux deux espaces d'être complets, mais seul la complétude de l'espace de départ importe.

Si l'un des (Ω_n) est d'intérieur non-vidé, on a alors :

$$\exists n_0 \in \mathbb{N}, \exists x \in E, \exists r > 0 : B_E(x, r) \subset \Omega_{n_0} \quad (5)$$

Ainsi, pour tout $y \in B_E(0, r)$, on a :

$$\forall i \in I, \|\varphi_i(y)\|_F \leq \|\varphi_i(x+y)\|_F + \|\varphi_i(x)\|_F \leq 2n_0 \quad (6)$$

Et donc :

$$\forall i \in I, \|\varphi_i\| \leq \frac{2n_0}{r} < +\infty \quad (7)$$

ce qui prouve notre premier point.

Si au contraire tous les (Ω_n) sont d'intérieur vidé, alors en appliquant le lemme de Baire dans E , on obtient que $\bigcup_{n \in \mathbb{N}} \Omega_n$ est d'intérieur vidé. En passant au complémentaire :

$$\left(\bigcup_{n \in \mathbb{N}} \Omega_n \right)^c = \bigcap_{n \in \mathbb{N}} \{x \in E \mid \exists i \in I : \|\varphi_i(x)\|_F > n\} \quad (8)$$

est dense dans E , ce qui prouve notre deuxième point.

Puisque nécessairement, on tombe dans l'un de ces deux cas, on prouve le théorème de Banach-Steinhaus !

Série de Fourier divergente ([Goub], Annexe A, exercice 8)

Introduisons quelques notations. On pose :

$$l_n : C_{2\pi}^0 \rightarrow \mathbb{C} \\ f \mapsto \sum_{k=-n}^n c_f(f)$$

l'application qui à f associe sa n -ième somme partielle de Fourier évaluée en 0. On se servira dans la suite du n -ième noyau de Dirichlet :

$$D_n = \sum_{k=-n}^n e_k \quad (9)$$

dont les propriétés peuvent être trouvées dans [Amr], chapitre IV. On rappelle seulement :

$$\forall f \in C_{2\pi}^0, l_n(f) = f * D_n(0) \quad (10)$$

On va appliquer le théorème de Banach-Steinhaus pour démontrer l'existence de fonctions pathologiques, continues sans pour autant être sommes de leurs séries de Fourier.

Application 110. L'ensemble des fonctions de $C_{2\pi}^0$ dont la série de Fourier diverge en 0 est dense dans $C_{2\pi}^0$.

Il est clair que la série de Fourier d'une fonction f diverge en 0 si, et seulement si, la famille $(l_n(f))_{n \in \mathbb{N}}$ est non bornée. Pour montrer que cet ensemble est dense, on va chercher à appliquer le théorème de Banach-Steinhaus et montrer que la famille (l_n) est non-bornée en norme d'opérateur.

Les (l_n) sont linéaires continues. En effet, il est évident qu'il s'agit d'une famille d'applications linéaires. De plus :

$$\forall f \in C_{2\pi}^0, |l_n(f)| = \frac{1}{2\pi} \left| \int_{-\pi}^{\pi} f(t) D_n(-t) dt \right| \leq \frac{1}{2\pi} \|f\|_{\infty} \|D_n\|_1 \quad (11)$$

Donc l_n est linéaire continue et $\|l_n\| \leq \|D_n\|_1$.

Cette inégalité est une égalité. Pour cela, on va chercher à approcher le cas d'égalité. Intuitivement, on voit que si f désigne le signe de D_n (qui n'est pas continu), on aurait :

$$\|f\|_{\infty} = 1 \wedge |l_n(f)| = \frac{1}{2\pi} \int |D_n(t)| dt \quad (12)$$

où on utilise le fait que D_n est paire. Une idée est donc d'approcher le signe de D_n par des fonctions continues. Fixons $n \in \mathbb{N}$ et posons :

$$\forall \varepsilon > 0, f_{\varepsilon} : x \mapsto \frac{D_n(x)}{|D_n(x)| + \varepsilon} \quad (13)$$

f_{ε} est bien continue et 2π -périodique. Par ailleurs, on a clairement :

$$\forall \varepsilon > 0, \|f_{\varepsilon}\|_{\infty} \leq 1 \quad (14)$$

Enfin :

$$|l_n(f_{\varepsilon})| = \frac{1}{2\pi} \left| \int_{-\pi}^{\pi} \frac{D_n(t)^2}{|D_n(t)| + \varepsilon} dt \right| \quad (15)$$

Or l'intégrande converge simplement en croissant vers $|D_n|$. Par théorème de convergence monotone, on a :

$$\lim_{\varepsilon \rightarrow 0} |l_n(f_{\varepsilon})| = \frac{1}{2\pi} \|D_n\|_1 \geq \liminf_{\varepsilon \rightarrow 0} \frac{1}{2\pi} \|f_{\varepsilon}\|_{\infty} \|D_n\|_1 \quad (16)$$

Ce qui permet de conclure :

$$\|l_n\| = \frac{1}{2\pi} \|D_n\|_1 \quad (17)$$

Estimation de $\|D_n\|_1$. D_n est égale presque partout à :

$$x \mapsto \frac{\sin\left(\frac{2n+1}{2}x\right)}{\sin\left(\frac{x}{2}\right)} \quad (18)$$

Ainsi on a, en exploitant l'inégalité classique $\sin(x) \leq x$:

$$\|D_n\|_1 \geq \int_{-\pi}^{\pi} 2 \left| \frac{\sin\left(\frac{2n+1}{2}t\right)}{t} \right| dt \quad (19)$$

$$= 2 \int_{-\frac{(2n+1)\pi}{2}}^{\frac{(2n+1)\pi}{2}} \left| \frac{\sin(u)}{u} \right| du \quad (20)$$

$$(21)$$

Or le sinus cardinal n'est pas intégrable sur $\mathbb{R}$. Donc cette norme tend vers $+\infty$ quand n tend vers l'infini.

Ainsi, la famille (l_n) n'est pas bornée en norme d'opérateur. Comme $C_{2\pi}^0$ est complet muni de la norme uniforme, le théorème de Banach-Steinhaus indique que l'ensemble de ses points de divergence est dense, ce qui est exactement ce que l'on voulait montrer !

Bonus : un exemple explicite ([Goub], IV.5, exercice 4)

Le lemme de Baire (et donc ses conséquences) repose sur une formulation faible de l'axiome du choix, et qui dit axiome du choix dit procédé hautement non-constructif. On est donc en droit de se demander : est-ce que l'existence de ces fonctions pathologiques est lié à l'axiome du choix ? Peut-on le démontrer dans ZF ? La réponse est oui, ce qui peut faire dire que le théorème de Banach-Steinhaus est un joujou un peu violent pour nos besoins. Je pense qu'il est bon d'avoir en tête un exemple explicite, à défaut de savoir le traiter par cœur. Heureusement, Gourdon nous en fournit un dans son livre !

Proposition 111. *La fonction :*

$$f : x \mapsto \sum_{p=1}^{+\infty} \frac{1}{p^2} \sin \left[(2^{p^3} + 1) \frac{x}{2} \right] \quad (22)$$

est bien définie, continue et 2π périodique, mais sa série de Fourier diverge en 0.

2.11 Théorème de Hahn-Banach - forme géométrique ★★

Recasages possibles : 213, 229, 253

On va, dans ce développement, établir le théorème de Hahn-Banach dans sa forme géométrique dans le cas particulier des espaces de Hilbert. Ce théorème permet une vision très géométrique de la notion de partie convexe qui trouve des applications très concrètes. On détaillera d'ailleurs l'une d'entre-elles, utile notamment pour prouver l'un de mes théorèmes préférés en analyse (rien que ça).

On considère $(H, \langle \cdot, \cdot \rangle)$ un espace de Hilbert réel, dont on note H' le dual topologique. Pour toute partie $C \subset H$ convexe, fermée et non-vide, on notera p_C la projection orthogonale sur C fournie par le théorème de projection. On utilisera les définitions suivantes :

Définition 112 (Hyperplan affine). Une partie F de H est appelée hyperplan affine (fermé) de H lorsqu'il existe une forme linéaire continue $\varphi : H \rightarrow \mathbb{R}$ et un réel α tels que :

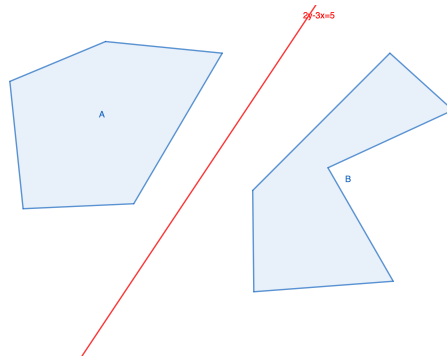
$$F = \varphi^{-1}(\{\alpha\}) \quad (1)$$

Définition 113 (Hyperplan séparateur). On dit que deux parties disjointes A et B de H sont séparées au sens strict (resp. au sens large) par un hyperplan affine lorsque :

$$\exists \varphi \in H', \exists \alpha \in \mathbb{R} : \sup_{x \in A} \varphi(x) < \alpha < \inf_{y \in B} \varphi(y) \quad (2)$$

(resp. avec une inégalité large)

Remarque : Géométriquement, cela se traduit bien par le fait que l'hyperplan affine d'équation $\varphi(x) = \alpha$ sépare strictement A et B , comme sur le dessin :



Nous allons démontrer ici le théorème suivant :

Théorème 114 (Hahn-Banach (forme géométrique) ([BMP05], application 3.16)). Soient A et B deux parties disjointes de H . On suppose :

1. A convexe et fermée.
2. B convexe, non vide et compacte.

Alors A et B sont séparées au sens strict par un hyperplan affine de H .

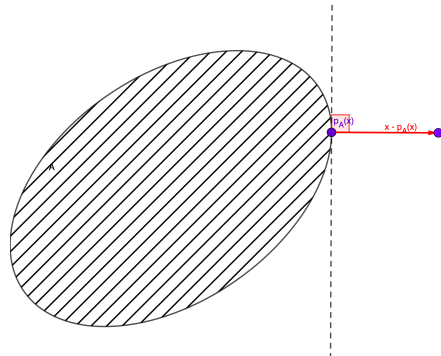
Théorème de Hahn-Banach

Remarque : Le cas où A est vide est vrai par vacuité, car toute forme linéaire continue de H est majorée sur B par compacité. On supposera donc dans la suite que A est également non vide (ce qui, admettons-le, est nettement plus intéressant).

Attention toutefois, le théorème est faux si B est vide : on pourrait alors prendre $A = H$... ♦

Démonstration.

Supposons d'abord que B est un singleton. On note alors $B = \{x\}$, avec $x \notin A$. La clef de la démonstration réside dans le théorème de projection sur un convexe fermé : on considère p_A la projection sur A . Le dessin suivant, en dimension 2, contient alors à lui seul toutes les informations nécessaires au bon déroulement de la preuve :



On part de l'intuition géométrique suivante : l'hyperplan séparateur recherché peut être pris orthogonal à $x - p_A(x)$. Ainsi, on cherche un hyperplan de la forme :

$$F := \text{Ker}(\langle x - p_A(x), - \rangle) + \alpha \quad (3)$$

Il ne nous reste qu'à choisir correctement α . Notons $(\varphi : y \mapsto \langle x - p_A(x), y \rangle) \in H'$. On remarque alors, par caractérisation de l'angle obtus, et parce que $x \notin A$:

$$\forall y \in A, \varphi(y - p_A(x)) \leq 0 < \|x - p_A(x)\|^2 = \varphi(x - p_A(x)) \quad (4)$$

Ainsi, par linéarité :

$$\forall y \in A, \varphi(y) \leq \varphi(p_A(x)) < \varphi(x) \quad (5)$$

On prend alors $\alpha := \frac{\varphi(x) + \varphi(p_A(x))}{2}$ et le tour est joué ! L'inégalité précédente donne :

$$\forall y \in A, \varphi(y) < \alpha < \varphi(x) \quad (6)$$

Revenons au cas général. On considère :

$$C := B - A := \{x - y, (x, y) \in B \times A\} \quad (7)$$

Il s'agit d'une partie non vide, convexe et fermée de H . En effet, le produit $A \times B$ est convexe dans H^2 par convexité de A et de B . Comme C est son image par l'application continue $(y, x) \mapsto x - y$, c'est un convexe. Si de plus $(x_n - y_n)$ est une suite d'éléments de C convergeant dans H vers l , avec $(x_n) \in B^{\mathbb{N}}$ et $(y_n) \in A^{\mathbb{N}}$, par compacité de B , on peut extraire une sous-suite le long de laquelle (x_n) converge vers $x \in B$. En passant à la limite le long de cette sous suite dans l'expression $y_n = x_n - (x_n - y_n)$, on trouve que y_n admet $x - l$ comme valeur d'adhérence, qui se trouve dans A puisque cette partie est fermée. Il vient alors $l = x - (x - l) \in C$. Donc C est bien fermée.

Puisque A et B sont disjointes, C ne contient pas 0. Par le travail précédemment effectué, il existe un hyperplan affine qui sépare C et $\{0\}$. En d'autres termes, on dispose d'une forme linéaire continue $\varphi \in H'$ et de $\alpha \in \mathbb{R}$ tels que :

$$\forall (y, x) \in A \times B, \varphi(y - x) < \alpha < \varphi(0) = 0 \quad (8)$$

et donc :

$$\forall (y, x) \in A \times B, \varphi(y) < \alpha + \varphi(x) < \varphi(x) \quad (9)$$

enfin, comme c'est vrai pour tout x et y , et comme α est strictement négatif :

$$\sup_{y \in A} \varphi(y) \leq \inf_{x \in B} \varphi(x) + \alpha < \inf_{x \in B} \varphi(x) \quad (10)$$

Quitte à modifier la valeur de α , on peut supposer qu'on a l'inégalité stricte, ce qui est une exacte reformulation du fait que A et B sont séparés au sens strict par un hyperplan affine. $\square$

Remarque : Ce théorème reste vrai dans le cas plus général des espaces de Banach, mais la preuve est très différente car on n'a plus accès aux commodités hilbertiens. Celle-ci requiert en particulier l'axiome du choix et la notion de jauge d'un convexe. Vous pouvez la trouver dans [Bre87]. A titre personnel, je pense que ce cadre élargi n'a pas d'intérêt au niveau de l'agreg, voire même risque de desservir à se placer dans un cadre plus abstrait, qui repose sur plus d'outils hors-programme, à moins d'avoir de vraies applications à en proposer. En plus, ça se recase moins-bien... $\blacklozenge$

Application à l'étude des fonctions convexes

On va maintenant fournir à l'aide du théorème de Hahn-Banach un critère pratique de convexité d'une fonction définie sur un espace de Hilbert réel.

Théorème 115 ([Li], chap. VI, exercice 3). Soit C une partie convexe fermée et non-vide de H et $f : C \rightarrow \mathbb{R}$ une application continue. On a alors l'équivalence suivante :

1. f est convexe.
2. Il existe $(g_i)_{i \in I}$ une famille d'applications affines de H telle que :

$$\forall x \in C, f(x) = \sup_{i \in I} g_i(x) \quad (11)$$

Démonstration.

Dans le sens indirect , il suffit de vérifier que le supremum d'une famille d'applications convexes est convexe, ce qui est élémentaire. Comme les applications affines sont convexes, f est bien convexe.

Dans le sens direct , on va exploiter le théorème de Hahn-Banach géométrique. On considère l'épigraphe de f , c'est-à-dire la partie de $H \times \mathbb{R}$:

$$\mathcal{E}(f) := \{(x, s) \in C \times \mathbb{R} \mid s \geq f(x)\} \quad (12)$$

Par convexité de f , $\mathcal{E}(f)$ est un convexe de $H \times \mathbb{R}$. Il est par ailleurs fermé par continuité de f et fermeture de C . Fixons alors $x_0 \in C$ et un réel $r < f(x_0)$, de façon à ce que (x_0, r) n'appartienne pas à $\mathcal{E}(f)$. Puisqu'on a une structure d'espace de Hilbert sur $H \times \mathbb{R}$ donnée par le produit scalaire

$$\langle (x_1, s_1), (x_2, s_2) \rangle := \langle x_1, x_2 \rangle_H + s_1 s_2 \quad (13)$$

on peut séparer strictement $\mathcal{E}(f)$ et (x_0, r) par un hyperplan affine d'équation $\Phi(x, s) = \alpha$, avec $\Phi \in (H \times \mathbb{R})'$ et $\alpha \in \mathbb{R}$. Posons $(\psi : x \mapsto \Phi(x, 0)) \in H'$ et $\lambda \in \mathbb{R}$ tel que :

$$\forall s \in \mathbb{R}, \Phi(0, s) = \lambda s \quad (14)$$

de sorte à pouvoir écrire :

$$\forall (x, s) \in H \times \mathbb{R}, \Phi(x, s) = \psi(x) + \lambda s \quad (15)$$

On a alors, par définition :

$$\forall (x, s) \in \mathcal{E}(f), \psi(x_0) + \lambda r < \alpha < \psi(x) + \lambda s \quad (16)$$

En particulier, $(x, s) = (x_0, f(x_0)) \in \mathcal{E}(f)$, on trouve en basculant tous les termes à gauche de l'inégalité :

$$\lambda(r - f(x_0)) < 0 \quad (17)$$

et comme r est choisi inférieur strict à $f(x_0)$, il vient que λ est strictement positif. Autrement dit, l'hyperplan d'équation $\Phi(x, s) = \alpha$ est le graphe de l'application affine :

$$\begin{aligned} g_{r, x_0} : H &\rightarrow \mathbb{R} \\ x &\mapsto \frac{1}{\lambda}(\alpha - \psi(x)) \end{aligned} \quad (18)$$

et par construction, on a :

$$\forall x \in C, \psi(x_0) + \lambda r < \underbrace{\psi(x) + \lambda g_{x_0, r}(x)}_{=\alpha} < \psi(x) + \lambda f(x) \quad (19)$$

D'où :

$$\forall x \in C, g_{x_0, r}(x) < f(x) \quad (20)$$

et donc, puisque ceci vaut quelque soit x_0 et r , on peut passer au sup :

$$\sup_{\substack{x_0 \in C \\ r < f(x_0)}} g_{x_0, r}(x) \leq f(x) \quad (21)$$

Mais pour x_0 fixé, on trouve :

$$\forall r < f(x_0), \quad r < g_{x_0,r}(x_0) < f(x_0) \quad (22)$$

et donc par passage au sup :

$$\sup_{r < x_0} g_{x_0,r}(x_0) = f(x_0) \quad (23)$$

d'où finalement l'égalité :

$$\sup_{\substack{x_0 \in C \\ r < f(x_0)}} g_{x_0,r}(x) = f(x) \quad (24)$$

□

Inégalité de Jensen et applications

Je ne traite pas cette partie lors du développement par manque de temps, mais il est bon de l'avoir en tête comme corollaire du résultat précédent. J'ai l'impression que l'inégalité de Jensen passe généralement un peu inaperçue lors des cours de théorie de la mesure, et pourtant, elle peut permettre de simplifier un bon nombre de preuves usuelles qui sont souvent menées de façon plus calculatoire. Je me suis permis de mettre l'inégalité de Hölder en exemple.

Théorème 116 (Inégalité de Jensen). Soient $d \in \mathbb{N}^*$ et $(X, \mathfrak{F}, \mu)$ un espace probabilisé. On considère $f : X \rightarrow \mathbb{R}^d$ une fonction intégrable et $\varphi : \mathbb{R}^d \rightarrow \mathbb{R}$ une fonction convexe telle que $\varphi \circ f$ est intégrable. Alors :

$$\varphi \left(\int_X f(x) d\mu(x) \right) \leq \int_X \varphi \circ f(x) d\mu(x) \quad (25)$$

Démonstration. On se donne $(g_i)_{i \in I}$ une famille de fonctions affines de $\mathbb{R}^d$ dans $\mathbb{R}$ telle que :

$$\varphi = \sup_{i \in I} g_i \quad (26)$$

Chaque fonction affine g_i peut se décomposer sous la forme $g_i = l_i + c_i$, avec l_i une forme linéaire et c_i une constante réelle. Or par définition de l'intégrale, l'opérateur d'intégration commute avec les formes linéaires continues. De plus, comme μ est une mesure de probabilités, les constantes sont invariantes par intégrale contre μ . Autrement dit :

$$\forall i \in I, \quad g_i \left(\int_X f(x) d\mu(x) \right) = \int_X g_i \circ f(x) d\mu(x) \quad (27)$$

Par croissance de l'intégrale, on peut passer au sup sous l'intégrale de droite :

$$\forall i \in I, \quad g_i \left(\int_X f(x) d\mu(x) \right) \leq \int_X \underbrace{\sup_{j \in I} g_j \circ f(x)}_{=\varphi} d\mu(x) \quad (28)$$

et comme c'est vrai pour tout $i \in I$, on obtient par passage au sup :

$$\varphi \left(\int_X f(x) d\mu(x) \right) \leq \int_X \varphi \circ f(x) d\mu(x) \quad (29)$$

□

Remarque : Dans le cas particulier $d = 1$, il existe une preuve beaucoup plus élémentaire de l'inégalité de Jensen (voir par exemple [BP]). Ce cas particulier suffit notamment à montrer bon nombre de résultats usuels, comme l'inégalité de Hölder, ou le théorème d'approximation par convolution. ♦

Corollaire 117 (Inégalité de Hölder). *Soit $(X, \mathfrak{F}, \mu)$ un espace mesuré. Soit $p \in [1, +\infty]$. Soit q son exposant conjugué, c'est-à-dire l'unique réel tel que $\frac{1}{p} + \frac{1}{q} = 1$. Pour toute fonction $f \in L^p(\mu)$ et toute fonction $g \in L^q(\mu)$, le produit fg est intégrable et on a :*

$$\|fg\|_1 \leq \|f\|_p \|g\|_q \quad (30)$$

Démonstration. Le cas $p \in \{1, +\infty\}$ est immédiat et résulte de la croissance de l'intégrale. On supposera donc que ni p ni q n'est égal à $+\infty$. On écrit alors :

$$\int_X |f(x)| |g(x)| d\mu(x) = \|f\|_p^p \left(\left(\int_X |f(x)|^{1-p} |g(x)| \frac{|f(x)|^p}{\|f\|_p^p} d\mu(x) \right)^q \right)^{1/q} \quad (31)$$

Mais la fonction $x \mapsto \frac{|f(x)|^p}{\|f\|_p^p}$ est intégrable, positive et d'intégrale 1. Ainsi, la mesure $\frac{|f|^p}{\|f\|_p^p} d\mu$ est une mesure de probabilités sur X . Comme la fonction $x \mapsto |x|^q$ est convexe sur $\mathbb{R}$, on peut appliquer l'inégalité de Jensen et pousser la puissance q -ième sous l'intégrale :

$$\int_X |f(x)| |g(x)| d\mu(x) \leq \|f\|_p^p \left(\int_X |f(x)|^{q(1-p)} |g(x)|^q \frac{|f(x)|^p}{\|f\|_p^p} d\mu(x) \right)^{1/q} \quad (32)$$

Il ne reste qu'à remarquer que $q = \frac{p}{p-1}$ pour conclure, car on a alors :

$$\forall x \in X, |f(x)|^{q(1-p)} |f(x)|^p = 1 \quad (33)$$

et donc :

$$\int_X |f(x)| |g(x)| d\mu(x) \leq \|f\|_p^p \times \|f\|_p^{-p/q} \left(\int_X |g(x)|^q d\mu(x) \right)^{1/q} \quad (34)$$

$$= \|f\|_p \|g\|_q \quad (35)$$

□

2.12 Théorème de Furstenberg-Kesten ★★

Recasages possibles : 223^(xxi), 262, 266

Nous allons démontrer un théorème attribué à Hillel Furstenberg et Harry Kesten, démontré dans les années 60, portant sur le comportement asymptotique d'un produit de matrices aléatoires. Ce théorème trouve des applications en physique, dans l'étude des systèmes dynamiques, par exemple un gaz. Il est fréquent de modéliser l'évolution d'un tel système par des produits successifs d'un vecteur avec des matrices aléatoires. La stabilité d'un tel système est liée à l'évolution de la norme du produit de ces matrices (voir [Kar10]). Ce théorème permet par ailleurs d'étendre d'une certaine manière la loi faible des grands nombres aux produits de matrices aléatoires.

J'ai à la base monté ce développement comme une revanche personnelle contre un sujet de l'ENS Ulm qui m'avait complètement bloqué en prépa. J'ai beaucoup plus tard dans l'année trouvé, par pure sérendipité, découvert que la preuve est faite dans [FGN22], mais j'avais déjà préparé et appris ma preuve, aussi la suite de ce document s'éloigne-t-elle pas mal de la référence.

Soit $d \in \mathbb{N}$. On considère $||| \cdot |||$ une norme d'opérateur sur $GL_d(\mathbb{C})$ et S une partie compacte de $GL_d(\mathbb{C})$. Soit μ une mesure de probabilité sur $GL_d(\mathbb{C})$ de support contenu dans S . L'objectif est de démontrer :

Théorème 118 (Furstenberg-Kesten, [21]). Soit $(M_n)_{n \in \mathbb{N}^*}$ une suite de variables aléatoires indépendantes et identiquement distribuées de matrices aléatoires de loi μ . On pose alors :

$$\Psi_n := \prod_{i=1}^n M_i \quad (1)$$

Il existe une constante déterministe $l(\mu) \in \mathbb{R}$ telle que :

$$\lim_{n \rightarrow +\infty} \frac{1}{n} \log |||\Psi_n||| \stackrel{(\text{xxii})}{=} l(\mu) \quad (2)$$

au sens L^1 .

Remarquons que lorsque $d = 1$, une matrice aléatoire est exactement une variable aléatoire complexe. La compacité du support de S permet d'assimiler les $(|||M_i|||)$ à une suite de variables aléatoires L^∞ et alors, la loi faible des grands nombres fournit :

$$\frac{1}{n} \log |||\Psi_n||| = \frac{1}{n} \sum_{i=1}^n \log |M_i| \xrightarrow{(\mathbb{P})} \mathbb{E}[\log |M_1|] \quad (3)$$

ce log étant bien défini car les matrices sont prises inversibles. Ainsi, dans ce cas, $l(\mu) = \mathbb{E}[\log |M_1|]$. Ce théorème est donc une extension de la loi faible des grands nombres au cas de matrices aléatoires.

Par ailleurs, dans le cas limite où S est un singleton $\{M\}$ (c'est-à-dire que Ψ_n peut être assimilé au produit déterministe M^n), le théorème de Gelfand affirme que $l(\mu)$ est le logarithme du rayon spectral de M .

(xxi). A condition de modifier la preuve, en insistant plus sur le lemme de Fekete et autres arguments de suites numériques que sur les arguments probabilistes.

(xxii). Remarquons que ce terme est presque sûrement bien défini car les (M_i) sont inversibles.

Un premier lemme probabiliste

Toute notre démonstration va s'axer autour du lemme suivant :

Lemme 119. Soit (X_n) une suite de variables aléatoires réelles presque sûrement dans $[a, b]$. On suppose qu'il existe $l \in \mathbb{R}$ tel que :

1. $\mathbb{E}[X_n] \rightarrow l$
2. $\forall \varepsilon > 0, \mathbb{P}(X_n \geq l + \varepsilon) \rightarrow 0$

Alors (X_n) converge vers l dans L^1 .

Démonstration. Soit $\varepsilon > 0$. Pour $n \in \mathbb{N}$, on pose :

$$A_n := (|X_n - l| \leq \varepsilon) \quad (4)$$

On a alors :

$$\mathbb{E}[|X_n - l|] = \mathbb{E}[|X_n - l| \mathbb{1}_{A_n}] + \mathbb{E}[|X_n - l|(1 - \mathbb{1}_{A_n})] \quad (5)$$

$$\leq \varepsilon + \mathbb{E}[(X_n - l)_+(1 - \mathbb{1}_{A_n})] + \mathbb{E}[(X_n - l)_-(1 - \mathbb{1}_{A_n})] \quad (\text{xxiii}) \quad (6)$$

Or, pour $\omega \in \Omega$, si $(X_n(\omega) - l)_+(1 - \mathbb{1}_{A_n}(\omega)) \neq 0$, on a $\omega \in \bar{A}_n \cap (X_n \geq l)$, c'est-à-dire $X_n(\omega) \geq l + \varepsilon$. On peut donc majorer :

$$0 \leq \mathbb{E}[(X_n - l)_+(1 - \mathbb{1}_{A_n})] \leq (b - l)_+ \mathbb{P}(X_n \geq l + \varepsilon) \quad (7)$$

et donc par hypothèse ce terme tend vers 0. Enfin, observons :

$$0 \leq \mathbb{E}[(X_n - l)_-(1 - \mathbb{1}_{A_n})] = \mathbb{E}[(X_n - l)_+(1 - \mathbb{1}_{A_n})] - \mathbb{E}[(X_n - l)] + \mathbb{E}[(X_n - l) \mathbb{1}_{A_n}] \quad (8)$$

$$\leq \varepsilon + o(1) \quad (9)$$

en majorant $(X_n - l) \mathbb{1}_{A_n}$ par ε et en utilisant l'hypothèse de convergence en espérance et la domination précédemment faite. En rassemblant tous les morceaux, on obtient :

$$\mathbb{E}[|X_n - l|] \leq (1 + (b - l)_+) \varepsilon + o(1) \quad (10)$$

et donc en passant à la limite supérieure quand $[\varepsilon \rightarrow 0]$, on obtient bien :

$$\lim_{n \rightarrow +\infty} \mathbb{E}[|X_n - l|] = 0 \quad (11)$$

ce qui achève la démonstration. $\square$

Remarque : Je ne démontre pas ce lemme à l'oral, la preuve étant technique et trop longue au vu de la suite. $\blacklozenge$

Il suffit donc de démontrer séparément la convergence en espérance vers une constante réelle $l(\mu)$ puis que pour tout $\varepsilon > 0$, on a :

$$\mathbb{P}\left(\frac{1}{n} \log |||\psi_n||| \geq l(\mu) + \varepsilon\right) \xrightarrow{n \rightarrow +\infty} 0 \quad (12)$$

(xxiii). Lorsque f est une fonction réelle, on note respectivement par f_+ et f_- ses parties positives et négatives.

Convergence en espérance

Dans un premier temps, on va montrer que les variables aléatoires $\left(\frac{1}{n}|||\Psi_n|||\right)$ sont presque sûrement astreintes à un intervalle compact.

Lemme 120. $\exists \alpha, \beta > 0 : \forall M \in S, \forall x \in \mathbb{C}^d, \alpha \|x\| \leq \|Mx\| \leq \beta \|x\|$

Démonstration. C'est une conséquence de la compacité du support S . En effet, en notant $S(0, 1)$ la sphère unité de $\mathbb{C}^d$ (compact), on peut considérer l'application :

$$\begin{aligned} S \times S(0, 1) &\rightarrow \mathbb{R}_+^* \\ (M, x) &\mapsto \|Mx\| \end{aligned}$$

qui est clairement continue. Son image est donc un compact de $\mathbb{R}_+^*$, de la forme $[\alpha, \beta]$. C'est-à-dire que pour tout $M \in S$ et pour tout $x \in \mathbb{C}^d$, on a :

$$\alpha \|x\| \leq \|Mx\| \leq \beta \|x\| \quad (13)$$

La stricte positivité de α et β provient du fait que S est un compact de $GL_d(\mathbb{C})$. $\square$

On déduit de suite, par récurrence immédiate, que $\alpha^n \leq |||\Psi_n||| \leq \beta^n$ et donc

$$\alpha \leq \frac{1}{n} |||\Psi_n||| \leq \beta \quad (14)$$

Montrons désormais l'existence d'une constante $l(\mu) \in \mathbb{R}$ telle que :

$$\frac{1}{n} \mathbb{E}[\log |||\Psi_n|||] \rightarrow l(\mu) \quad (15)$$

Pour cela, nous utiliserons le classique lemme sous-additif de Fekete.

Définition 121 (Suite sous-additive). Soit $(u_n)_{n \in \mathbb{N}}$ une suite de nombres réels. On dit que (u_n) est sous-additive si :

$$\forall (n, m) \in \mathbb{N}^2, u_{n+m} \leq u_n + u_m \quad (16)$$

Lemme 122 (Fekete). Soit (u_n) une suite sous-additive. Alors, la suite $\left(\frac{u_n}{n}\right)$ est convergente dans $\mathbb{R} \cup \{-\infty\}$ et :

$$\lim_{n \rightarrow +\infty} \frac{u_n}{n} = \inf_{n \in \mathbb{N}^*} \frac{u_n}{n} \quad (17)$$

Ainsi, il suffit de montrer que la suite $(\mathbb{E}[\log |||\Psi_n|||])$ est sous-additive.

Soit $n, m \in \mathbb{N}^{*2}$. Alors :

$$\mathbb{E}[\log |||\Psi_{n+m}|||] \leq \mathbb{E}[\log |||\prod_{i=1}^n M_i|||] + \mathbb{E}[\log |||\prod_{i=n+1}^{n+m} M_i|||] \quad (18)$$

$$= \mathbb{E}[\log |||\Psi_n|||] + \mathbb{E}[\log |||\Psi_m|||] \quad (19)$$

La majoration exploite la sous-multiplicativité des normes d'opérateur et la croissance du log et l'espérance. La deuxième égalité n'est pas triviale ! Elle vient du fait que les (M_i) sont i.i.d., et donc qu'on a l'égalité en loi :

$$\prod_{i=n+1}^{n+m} M_i \stackrel{(\mathcal{L})}{=} \Psi_m \quad (20)$$

La suite $(\mathbb{E}[\log |||\Psi_n|||])$ est donc sous-additive. En vertu du lemme de Fekete, $\left(\frac{1}{n}\mathbb{E}[\log |||\Psi_n|||]\right)$ converge vers son infimum.

On peut désormais appliquer le lemme :

$$\forall x \in \mathbb{C}^d, \quad ||\Psi_n x|| = \left\| \prod_{i=1}^n M_i x \right\| \quad (21)$$

$$\geq \alpha^n ||x|| \quad (22)$$

et donc :

$$|||\Psi_n||| \geq \alpha^n \quad (23)$$

Il vient donc immédiatement :

$$\frac{1}{n} \log |||\Psi_n||| \geq \log(\alpha) > 0 \quad (24)$$

Convergence des probabilités

La dernière chose à faire pour obtenir le théorème est de montrer que :

$$\forall \varepsilon > 0, \quad \mathbb{P} \left(\frac{1}{n} \log |||\Psi_n||| \geq l(\mu) + \varepsilon \right) \rightarrow 0 \quad (25)$$

et d'ensuite lui appliquer le lemme 119.

On commence par découper le produit Ψ_n pour faire apparaître des variables aléatoires réelles i.i.d. Soit $k_0 \in \mathbb{N}^*$. Pour tout $n \in \mathbb{N}$, on peut faire la division euclidienne :

$$n = qk_0 + r \quad (26)$$

$$0 \leq r < k_0 \quad (27)$$

On peut alors écrire :

$$\frac{1}{n} \log |||\Psi_n||| \leq \frac{1}{n} \left[\sum_{i=0}^{q-1} \log ||| \prod_{j=ik_0+1}^{(i+1)k_0} M_j ||| + \log ||| \prod_{j=qk_0+1}^n M_j ||| \right] \quad (28)$$

par sous-multiplicativité de la norme d'opérateur. Posons alors :

$$Q_i^{(k_0)} := \log ||| \prod_{j=ik_0+1}^{(i+1)k_0} M_j ||| \quad (29)$$

$$R^{(k_0)} := \log ||| \prod_{j=qk_0+1}^n M_j ||| \quad (30)$$

Par indépendance des (M_i) et égalité en loi, on a :

$$\forall i \in \llbracket 0, q-1 \rrbracket, Q_i^{(k_0)} \stackrel{(\mathcal{L})}{=} \log |||\Psi_{k_0}||| \quad (31)$$

Or les $(Q_i^{(k_0)})$ sont des variables L^1 car bornées, indépendantes et identiquement distribuées. D'après la loi forte des grands nombres :

$$\frac{1}{n} \sum_{i=0}^{q-1} Q_i^{(k_0)} = \frac{1}{q} \sum_{i=0}^{q-1} \frac{q}{n} Q_i^{(k_0)} \xrightarrow[n \rightarrow +\infty]{p.s.} \frac{1}{k_0} \mathbb{E}[\log |||\Psi_{k_0}|||] \quad (32)$$

Il est par ailleurs clair que $\frac{1}{n} R^{(k_0)}$ converge presque sûrement vers 0 quand n tend vers $+\infty$. Ainsi, pour tout $\varepsilon > 0$, on a :

$$\mathbb{P} \left(\frac{1}{n} \log |||\Psi_n||| \geq l(\mu) + \varepsilon \right) \leq \mathbb{P} \left(\frac{1}{n} \sum_{i=0}^{q-1} Q_i^{(k_0)} + \frac{1}{n} R^{(k_0)} \geq l(\mu) + \varepsilon \right) \quad (33)$$

$$\leq \mathbb{P} \left(\left[\frac{1}{n} \sum_{i=0}^{q-1} Q_i^{(k_0)} - \frac{1}{k_0} \mathbb{E}[\log |||\Psi_{k_0}|||] \right] + \left[\frac{1}{k_0} \mathbb{E}[\log |||\Psi_{k_0}|||] - l(\mu) \right] + \frac{1}{n} R^{(k_0)} \geq \varepsilon \right) \quad (34)$$

Il suffit maintenant de choisir k_0 de sorte à ce que :

$$\frac{1}{k_0} \mathbb{E}[\log |||\Psi_{k_0}|||] - l(\mu) \leq \frac{\varepsilon}{2} \quad (35)$$

ce qui est possible grâce à la convergence en espérance précédemment établie. On domine alors notre probabilité par :

$$\mathbb{P} \left(\left[\frac{1}{n} \sum_{i=0}^{q-1} Q_i^{(k_0)} - \frac{1}{k_0} \mathbb{E}[\log |||\Psi_{k_0}|||] \right] + \frac{1}{n} R^{(k_0)} \geq \frac{\varepsilon}{2} \right) \quad (36)$$

et ce dernier terme tend vers 0 lorsque n tend vers $+\infty$ car le terme de gauche tend presque sûrement vers 0.

On peut donc appliquer le lemme 119 pour obtenir le théorème !

Remarque : Le résultat original montre la convergence L^1 et la convergence presque sûre, comme dans la loi forte des grands nombres. Depuis, un théorème central limite a été montré. On pourra lire [Kar10] pour trouver des informations supplémentaires sur le théorème. ♦

Développement librement inspiré du sujet « Mathématiques D » du concours d'entrée 2021 de l'ENS Paris.

2.13 Théorème de Lévy et théorème central limite ★★ ★

Recasages possibles : 209, 234, 235, 239, 241, 261, 262, 266

Nous allons démontrer ici le célèbre théorème central limite dont l'utilité n'est pas plus à vanter. En chemin, nous aurons à montrer divers résultats sur la convergence en loi, en particulier le très important théorème de Lévy qui permet d'étudier la convergence en loi à l'aide de la fonction caractéristique.

C'est un résultat exigeant contre lequel j'aimerais mettre en garde. J'ai vu un certain nombre de personnes se lancer dedans, attirées par ses excellents recasages, mais croyez-moi : **le TCL n'est pas votre ami, et il ne se laissera pas démontrer comme ça**. On utilisera au cours de la preuve une assez grande variété de résultats non triviaux issus de domaines variés des maths, parmi lesquels :

1. Des théorèmes d'intégration contre des mesures abstraites.
2. L'espace de Schwartz, des résultats de densité à son propos et surtout le fait que la transformée de Fourier induit une bijection de cet espace dans lui-même.
3. Le logarithme holomorphe (notez que ce dernier point peut-être contourné).

Il faut évidemment en plus avoir une bonne aisance en probas, en particulier avec tout ce qui concerne la convergence en loi. Enfin, je pense qu'on peut difficilement présenter le TCL sans faire des stats, même un tout petit peu. Je suis d'avis que si vous n'évoquez même pas les intervalles de confiance, vous aurez forcément des questions dessus à l'oral. Ayez bien tout ceci en tête avant de vous lancer dans ce développement dont la difficulté est un peu occultée par la rédaction de Zwilly et Quéffélec.

Pour une suite de variables aléatoires (X_n) convergeant en loi vers une variable aléatoire X de loi $\mathcal{L}$, on notera :

$$X_n \xrightarrow[n \rightarrow +\infty]{(\mathcal{L})} X \quad \text{ou} \quad X_n \xrightarrow[n \rightarrow +\infty]{(\mathcal{L})} \mathcal{L} \quad (1)$$

Lorsque X est une variable aléatoire à valeurs dans $\mathbb{R}$, on notera φ_X sa fonction caractéristique et $\mathbb{P}_X$ sa loi.

On notera C_b^0 l'ensemble des fonctions continues bornées de $\mathbb{R}$ dans $\mathbb{C}$, C_0^0 celui des fonctions continues tendant vers 0 à l'infini et C_c^∞ celui des fonctions C^∞ à support compact.

Lorsque f est L^1 , on notera :

$$\mathcal{F}[f] : \xi \mapsto \int_{\mathbb{R}} f(x) e^{-i\xi x} dx \quad (2)$$

sa transformée de Fourier.

Théorème de Lévy [QZ], p.536

Le résultat suivant est dû à Paul Lévy. Il montre que la convergence de $\mathbb{E}[f(X_n)]$ vers $\mathbb{E}[f(X)]$ pour tout $f \in C_b^0$ est obtenue simplement par la convergence contre la famille à un paramètre $(x \mapsto e^{-itx})_{t \in \mathbb{R}}$, ce qui est loin d'être trivial !

Théorème 123 (Lévy). Soit (X_n) une suite de variables aléatoires réelles et X une autre variable aléatoire réelle. Les assertions suivantes sont équivalentes :

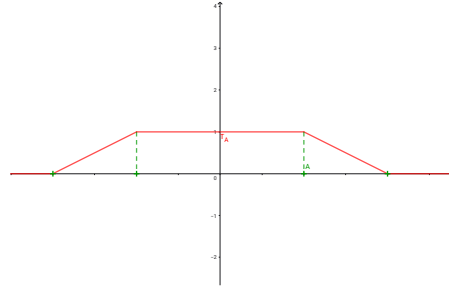
1. $X_n \xrightarrow[n \rightarrow +\infty]{(\mathcal{L})} X$
2. $\forall t \in \mathbb{R}, \varphi_{X_n}(t) \xrightarrow[n \rightarrow +\infty]{} \varphi_X(t)$

La démonstration du théorème de Lévy passera par le lemme suivant :

Lemme 124 ([QZ], II.18, p. 533). Les assertions suivantes sont équivalentes :

1. $\forall f \in C_b^0, \mathbb{E}[f(X_n)] \xrightarrow[n \rightarrow +\infty]{} \mathbb{E}[f(X)]$ (i.e. $X_n \xrightarrow[n \rightarrow +\infty]{(\mathcal{L})} X$)
2. $\forall f \in C_0^0, \mathbb{E}[f(X_n)] \xrightarrow[n \rightarrow +\infty]{} \mathbb{E}[f(X)]$

Démonstration. Il est évident que (1) $\implies$ (2) puisque $C_0^0 \subset C_b^0$. Soit donc $f \in C_b^0$. On va approcher f par des fonctions de C_0^0 à l'aide de fonctions trapèze. Etant donné un réel $A > 0$, on pose T_A la fonction trapézoïdale valant 1 sur $[-A, A]$ et 0 sur $\mathbb{R} \setminus [-2A, 2A]$.



Fixons $\varepsilon > 0$. Il existe alors $A \in \mathbb{R}$ tel que :

$$\mathbb{P}(|X| \geq A) = \int_{\mathbb{R}} \mathbf{1}_{|x| \geq A} d\mathbb{P}_X(x) \leq \varepsilon \quad (3)$$

On a alors :

$$|\mathbb{E}[f(X) - \mathbb{E}[f(X_n)]]| \leq \underbrace{|\mathbb{E}[f(X) - T_A f(X)]|}_{\alpha} + \underbrace{|\mathbb{E}[T_A f(X) - T_A f(X_n)]|}_{\beta_n} + \underbrace{|\mathbb{E}[T_A f(X_n) - f(X_n)]|}_{\gamma_n} \quad (4)$$

Contrôlons séparément chacun de ces trois termes.

•

$$\alpha = \left| \int_{\mathbb{R}} (1 - T_A(x)) f(x) d\mathbb{P}_X(x) \right| \quad (5)$$

$$\leq \|f\|_{\infty} \int_{\mathbb{R}} \mathbf{1}_{|x| \geq A} d\mathbb{P}_X(x) \quad (6)$$

$$\leq \|f\|_{\infty} \varepsilon \quad (7)$$

- $T_A f \in C_0^0$, donc par hypothèse :

$$\limsup_{n \rightarrow +\infty} \beta_n = 0 \quad (8)$$

•

$$\gamma_n = \left| \int_{\mathbb{R}} (T_A(x) - 1) f(x) d\mathbb{P}_{X_n}(x) \right| \quad (9)$$

$$\leq \|f\|_{\infty} \int_{\mathbb{R}} 1 - T_A(x) d\mathbb{P}_{X_n}(x) \quad (1 - T_A \geq 0) \quad (10)$$

$$= \|f\|_{\infty} \left(1 - \int_{\mathbb{R}} T_A(x) d\mathbb{P}_{X_n}(x) \right) \quad (11)$$

$$\xrightarrow{n \rightarrow +\infty} \|f\|_{\infty} \left(1 - \int_{\mathbb{R}} T_A(x) d\mathbb{P}_X(x) \right) \quad \text{car } T_A \in C_0^0 \quad (12)$$

$$\leq \|f\|_{\infty} \varepsilon \quad (13)$$

d'où on déduit :

$$\limsup_{n \rightarrow +\infty} \gamma_n \leq \|f\|_{\infty} \varepsilon \quad (14)$$

Finalement, en recombinaison ces inégalités, on obtient :

$$0 \leq \liminf_{n \rightarrow +\infty} |\mathbb{E}[f(X)] - \mathbb{E}[f(X_n)]| \leq \limsup_{n \rightarrow +\infty} |\mathbb{E}[f(X)] - \mathbb{E}[f(X_n)]| \leq 2\|f\|_{\infty} \varepsilon \quad (15)$$

et donc, en laissant tendre ε vers 0, on obtient bien $\mathbb{E}[f(X_n)] \xrightarrow{n \rightarrow +\infty} \mathbb{E}[f(X)]$. $\square$

On va maintenant démontrer le théorème de Lévy. Comme il est évident que (1) $\implies$ (2), on se contente du sens réciproque. Commençons par montrer la convergence lorsque $f \in \mathcal{F}(L^1(\mathbb{R}))$, c'est-à-dire lorsque f est la transformée de Fourier d'une fonction $\psi \in L^1$, c'est-à-dire :

$$f(x) = \int_{\mathbb{R}} \psi(t) e^{ixt} dt \quad (16)$$

On a alors :

$$\mathbb{E}[f(X_n)] = \int_{\mathbb{R}} \int_{\mathbb{R}} \psi(t) \exp(itx) dt d\mathbb{P}_{X_n}(x) \quad (17)$$

$$= \int_{\mathbb{R}} \psi(t) \int_{\mathbb{R}} \exp(itx) d\mathbb{P}_{X_n}(x) dt \quad (\text{théorème de Fubini}) \quad (18)$$

$$= \int_{\mathbb{R}} \psi(t) \varphi_{X_n}(t) dt \quad (\text{par définition de l'espérance}) \quad (19)$$

$$\xrightarrow{n \rightarrow +\infty} \int_{\mathbb{R}} \psi(t) \varphi_X(t) dt \quad (\text{par convergence dominée}) \quad (20)$$

$$= \mathbb{E} \left[\int_{\mathbb{R}} \psi(t) \exp(itX) dt \right] \quad (\text{en appliquant encore Fubini}) \quad (21)$$

$$= \mathbb{E}[f(X)] \quad (22)$$

Or, la transformée de Fourier est bijective sur l'espace de Schwartz^(xxiv). Donc toutes les fonctions Schwartz sont dans $\mathcal{F}(L^1)$, et en particulier, pour toute fonction $g \in C_c^{\infty}$, on a $\mathbb{E}[g(X_n)] \xrightarrow{n \rightarrow +\infty} \mathbb{E}[g(X)]$. Mais C_c^{∞} est dense dans C_0^0 pour la topologie de la norme ∞ ^(xxv).

(xxiv). Attention, ça n'est pas du tout évident à démontrer, et je pense qu'il faut au moins avoir une bonne idée de la preuve. Ça découle de la stabilité de l'espace de Schwartz par transformation de Fourier et de la formule d'inversion dans L^1 . Alternativement, on peut le démontrer à partir de la formule sommatoire de Poisson.

(xxv). Voir annexe.

Donc si f est une fonction de C_0^0 , pour tout $\varepsilon > 0$, il existe $g \in C_c^\infty$ telle que $\|f - g\|_\infty \leq \varepsilon$. On a alors :

$$|\mathbb{E}[f(X)] - \mathbb{E}[f(X_n)]| \leq |\mathbb{E}[f(X) - g(X)]| + |\mathbb{E}[g(X) - g(X_n)]| + |\mathbb{E}[g(X_n) - f(X_n)]| \quad (23)$$

$$\leq 2\|f - g\|_\infty + |\mathbb{E}[g(X) - g(X_n)]| \quad (24)$$

et donc

$$0 \leq \limsup_{n \rightarrow +\infty} |\mathbb{E}[f(X)] - \mathbb{E}[f(X_n)]| \leq 2\varepsilon \quad (25)$$

ce quelque soit ε . Il ne reste qu'à appliquer le lemme pour montrer le théorème de Lévy !

Remarque : Pour appuyer le recasage dans la 261, on peut observer ce corollaire immédiat :

Théorème 125. *Soient X, Y deux variables aléatoires réelles. Alors les assertions suivantes sont équivalentes :*

1. $\mathbb{P}_X = \mathbb{P}_Y$
2. $\varphi_X = \varphi_Y$

Démonstration. Il est clair que (i) $\implies$ (ii). Supposons donc que $\varphi_X = \varphi_Y$. On peut alors considérer la suite de variables aléatoires (X_n) constante égale à X . Alors φ_{X_n} converge simplement vers φ_Y et donc $X_n \xrightarrow[n \rightarrow +\infty]{(\mathcal{L})} Y$. Par unicité de la limite pour la topologie de la convergence en loi, on a alors $\mathbb{P}_X = \mathbb{P}_Y$. $\square$ $\blacklozenge$

Théorème central limite

On va maintenant démontrer le théorème central limite :

Théorème 126 (Central limite). *Soit (X_n) une suite de variables aléatoires réelles i.i.d. admettant une variance, notée $\sigma^2 \in \mathbb{R}_+$. Alors :*

$$\frac{1}{\sqrt{n\sigma^2}} \sum_{i=1}^n (X_i - \mathbb{E}[X_i]) \xrightarrow[n \rightarrow +\infty]{(\mathcal{L})} \mathcal{N}(0, 1) \quad (26)$$

où $\mathcal{N}(0, 1)$ désigne la loi normale standard.

Pour cela, on aura besoin du petit lemme :

Lemme 127. *Soit $z \in \mathbb{C}$. Alors :*

$$\left(1 + \frac{z}{n} + o\left(\frac{1}{n}\right)\right)^n \xrightarrow[n \rightarrow +\infty]{} e^z \quad (27)$$

Démonstration. Pour n suffisamment grand, $\left|\frac{z}{n} + o\left(\frac{1}{n}\right)\right| \leq \frac{1}{2}$. Désignons par $\log$ la détermination

principale du logarithme holomorphe^(xxvi). A partir d'un certain rang, on a donc :

$$\left(1 + \frac{z}{n} + o\left(\frac{1}{n}\right)\right)^n = \exp\left(n \log\left(1 + \frac{z}{n} + o\left(\frac{1}{n}\right)\right)\right) \quad (28)$$

$$= \exp(z + o(1)) \quad (29)$$

$$= e^z(1 + o(1)) \quad (30)$$

ce qui prouve le lemme. $\square$

Tous les ingrédients sont maintenant réunis pour prouver le théorème central limite. Notons φ la fonction caractéristique commune des (X_i) et S_n le terme de gauche dans (26). Sans perte de généralité, on suppose les (X_i) centrées et réduites.

$$\forall t \in \mathbb{R}, \varphi_{S_n}(t) = \varphi_{\frac{1}{\sqrt{n}} \sum_{i=1}^n X_i}(t) \quad (31)$$

$$= \prod_{i=1}^n \varphi_{X_i}\left(\frac{t}{\sqrt{n}}\right) \quad (32)$$

$$= \left(1 + i\mathbb{E}[X_1] \frac{t}{\sqrt{n}} - \mathbb{E}[(X_1)^2] \frac{t^2}{2n} + o\left(\frac{t}{n}\right)\right)^n \quad (33)$$

$$= \left(1 - \frac{t^2}{2n} + o\left(\frac{t}{n}\right)\right)^n \quad (34)$$

$$\xrightarrow[n \rightarrow +\infty]{} e^{-\frac{t^2}{2}} \quad (35)$$

en appliquant le lemme^(xxvii). Ainsi, la fonction caractéristique de S_n converge simplement vers celle de la loi normale centrée réduite. Par le théorème de Lévy, on en déduit le théorème central limite.

Annexe

On va montrer ici un théorème utilisé plus tôt :

Théorème 128. C_c^∞ est dense dans C_0^0 pour la topologie de la norme ∞ .

Je n'ai pas trouvé (en toute honnêteté, je n'ai pas cherché) de référence qui fasse la démonstration. Dans [QZ], on ne trouve qu'un flegmatique "Stone-Weierstraß". Je ne doute pas que le résultat découle directement de Stone-Weierstraß, mais pour être honnête je préfère éviter autant que possible cet énorme théorème pour me rabattre sur ses variantes affaiblies, comme Weierstraßpolynomial ou Féjèr. C'est moins dangereux que de parler d'algèbres séparantes, et j'ai toujours eu l'impression que les preuves au cas par cas m'apprenaient quelque chose.

Démonstration. Soit $f \in C_0^0$, soit $\varepsilon > 0$. Soit $A > 0$ tel que :

$$\forall x \in \mathbb{R}, |x| \geq A \implies |f(x)| \leq \varepsilon \quad (36)$$

(xxvi). Il est possible de faire une preuve plus élémentaire de ce résultat à l'aide d'une simple inégalité. Je trouve personnellement que la preuve avec le log holomorphe est plus élégante et efficace, et puisque celui-ci est au programme, on aurait tort de se priver. Toutefois, il faut bien avoir conscience que ce résultat, très simple à prouver dans le cas réel, nécessite des arguments plus raffinés pour fonctionner dans le cas complexe.

(xxvii). Ici, la variable t est réelle, et donc on pourrait être tenté de penser qu'on pourrait se contenter du lemme dans le cas réel, et donc évitant ainsi le log holomorphe. Attention ici, car c'est le o qui est un complexe ! On a donc vraiment besoin du lemme dans le cas complexe, le cas réel ne suffit pas.

(qui existe car f tend vers 0 en $\pm\infty$). On fixe alors χ une fonction plateau telle que :

1. $\chi \in C_c^\infty$
2. $0 \leq \chi \leq 1$
3. $\forall x \in \mathbb{R}, (|x| \leq A \implies \chi(x) = 1) \wedge (|x| \geq A+1 \implies \chi(x) = 0)$

□

On a alors :

$$\|f - \chi f\|_\infty \leq \varepsilon \quad (37)$$

Par ailleurs, le théorème d'approximation de Stone-Weierstraß indique qu'il existe P une fonction polynomiale telle que :

$$\|(f - P)\mathbb{1}_{[-A-1, A+1]}\|_\infty \leq \varepsilon \quad (38)$$

ce qui donne :

$$\|\chi f - \chi P\|_\infty \leq \varepsilon \quad (39)$$

Finalement, on a bien :

$$\|f - \chi P\|_\infty \leq 2\varepsilon \quad (40)$$

avec $\chi P \in C_c^\infty$.

2.14 Théorème de Liapounov ★★

Recasages possibles : 204, 215, 220, 221

L'objectif de notre développement est de démontrer une partie du théorème de linéarisation de Liapounov, qui permet d'étudier les points d'équilibre d'un système différentiel autonome en «linéarisant» le système au point d'équilibre, ce qui consiste à remplacer le champ de vecteur par sa différentielle en l'équilibre. C'est un intéressant résultat d'équa diff, qui se recase dans les deux leçons 220 et 221, tout en utilisant d'intéressants résultats d'algèbre linéaire, particulièrement en lien avec les exponentielles de matrices, et de connexité. Notez que les preuves qu'on trouve généralement dans les livres d'équations différentielles procèdent différemment. L'argument de connexité apparaît dans [QZ], livre que je trouve personnellement atroce, mais eh, le recasage en vaut la peine.

On considère d un entier, U un ouvert de $\mathbb{R}^d$ contenant 0 et f un champ de vecteur sur U , de classe C^1 , s'annulant en 0. D'après le théorème de Cauchy-Lipschitz, le système différentiel :

$$y' = f(y) \quad (\star)$$

admet une unique solution maximale issue de n'importe quel point de U . Pour alléger les notations, nous noterons A la matrice jacobienne de f en 0, c'est-à-dire la matrice dans la base canonique de l'application linéaire $df(0)$. On va montrer le résultat suivant :

Théorème 129 (Liapounov ([QZ], chap X, th IV.1)). *On suppose que la matrice A a toutes ses valeurs propres (complexes) de partie réelle strictement négative. Alors le point d'équilibre 0 est asymptotiquement stable.*

Un lemme d'algèbre linéaire

Lemme 130. *Soit $M \in \mathbb{M}_d(\mathbb{C})$. On suppose qu'il existe $\mu > 0$ tel que :*

$$\forall \lambda \in Sp(M), \Re(\lambda) < -\mu \quad (1)$$

Alors il existe une constante C_μ telle que :

$$\forall t \geq 0, |||\exp(tM)||| \leq C_\mu e^{-t\mu} \quad (2)$$

Démonstration. Par équivalence des normes en dimension finie, on peut travailler avec la norme infinie. Toute la démonstration repose sur une utilisation astucieuse du lemme des noyaux. Considérons le polynôme caractéristique χ_M de M qu'on écrit sous forme factorisée :

$$\chi_M = \prod_{i=1}^p (X - \lambda_i)^{\alpha_i} \quad (3)$$

Par le théorème de Cayley-Hamilton et le lemme des noyaux, on obtient la décomposition de $\mathbb{R}^d$

en somme des espaces caractéristiques de M :

$$\mathbb{R}^d = \bigoplus_{i=1}^p \text{Ker}((\lambda_i I_d - M)^{\alpha_i}) \quad (4)$$

Soit $x \in \mathbb{R}^d$. x admet une unique décomposition $x = x_1 + \dots + x_p$ avec $x_i \in \text{Ker}((\lambda_i I_d - M)^{\alpha_i})$. Ainsi, on a :

$$\exp(tM) \cdot x = \sum_{i=1}^p \exp(t(\lambda_i M - \lambda_i I_d + \lambda_i I_d)) \cdot x_i \quad (5)$$

$$= \sum_{i=1}^p \sum_{n=0}^{\alpha_i-1} t^n e^{\lambda_i t} \frac{(M - \lambda_i I_d)^n}{n!} \cdot x_i \quad (6)$$

Or, comme $\lambda_i < -\mu$, on a $t^n e^{\lambda_i t} = O_{t \rightarrow +\infty}(e^{-\mu t})$, ce pour tout $i \in \llbracket 1, p \rrbracket$ et $n \in \mathbb{N}$. Ainsi, il existe $C > 0$ tel que :

$$\forall i \in \llbracket 1, d \rrbracket, \forall t \in \mathbb{R}_+, \forall n \leq \max_{0 \leq i \leq d} \alpha_i, t^n e^{\lambda_i t} \leq C e^{-\mu t} \quad (7)$$

Il vient donc :

$$\|\exp(tM) \cdot x\|_\infty \leq C e^{-t\mu} \sum_{i=1}^p \sum_{n=0}^{\alpha_i-1} \frac{\|M - \lambda_i I_d\|}{n!} \|x_i\|_\infty \quad (8)$$

$$\leq \|x\|_\infty \times e^{-t\mu} \times \underbrace{C \sum_{i=1}^p \sum_{n=0}^{\alpha_i-1} \frac{\|M - \lambda_i I_d\|}{n!}}_{=: C_\mu} \quad (9)$$

Comme ceci vaut pour tout $x \in \mathbb{R}^d$, on trouve bien le résultat souhaité. $\square$

Le théorème de Liapounov

Ce lemme en poche, on va démontrer le théorème de Liapounov. Pour cela, il y a deux choses à faire : montrer que les solutions partant d'un point au voisinage de l'équilibre sont globales, et qu'elles tendent vers 0 à l'infini. Faisons d'une pierre de coup en démontrant le lemme suivant qui contient tout ce qu'il faut pour conclure :

Lemme 131. *Soit $\mu_0 > 0$ tel que pour toute valeur propre λ de A , $\Re(\lambda) < \mu_0$. Alors :*

$$\forall 0 < \mu < \mu_0, \forall C > 0, \exists 0 < \eta < C : \forall \|x_0\| \leq \eta, \forall t \in [0, T^*[, \|x(t)\| \leq C e^{-\mu t} \quad (10)$$

avec x la solution maximale issue de x_0 et T^ le supremum de son intervalle de définition.*

Remarque : Vous le sentez, cette démonstration va être lourde en notations, avec des paramètres dans tous les sens qui dépendent les uns des autres. Accrochez-vous, et si besoin est, n'hésitez pas à marquer explicitement les dépendances des paramètres introduits! $\blacklozenge$

Démonstration. On va travailler sur le lien entre notre système différentiel et le système linéarisé. On considère le développement limité de f en 0 :

$$\forall y \in \mathbb{R}^d, f(y) = f(0) + df(0) \cdot y + g(y) = A \cdot y + g(y) \quad (\star)$$

avec $g : \mathbb{R}^d \rightarrow \mathbb{R}^d$ une fonction telle que :

$$\frac{g(y)}{\|y\|} \xrightarrow{y \rightarrow 0} 0 \quad (11)$$

Il existe un seuil $\delta > 0$ tel que :

$$\forall y \in \mathbb{R}^d, \|y\| \leq \delta \implies \|g(y)\| \leq \|y\| \quad (12)$$

Soient $0 < C < \delta$ ^(xxviii) et $0 < \mu < \mu_0$. Soit $\|x_0\| < C$, notons x la solution issue de x_0 ^(xxix) et posons T^* le supremum de l'intervalle de définition de x . On considère alors l'ensemble :

$$\mathcal{F} := \{t \in [0, T^*[\mid x(t) \leq Ce^{-\mu t}\} \quad (13)$$

Par continuité de x et de l'exponentielle, $\mathcal{F}$ est un fermé de $[0, T^*[$ ^(xxx). Le choix de $\|x_0\|$ implique que $0 \in \mathcal{F}$. Ainsi, on peut considérer $\mathcal{E}$ la composante connexe de $\mathcal{F}$ contenant 0. C'est un intervalle (car connexe dans $\mathbb{R}$) fermé dans $\mathcal{F}$ (mais pas nécessairement fermé dans $\mathbb{R}$) car les composantes connexes d'un espace métrique sont toujours fermées dans cet espace. Comme $\mathcal{F}$ est fermé dans $[0, T^*[$, $\mathcal{E}$ est fermé dans $[0, T^*[$. Si on montre que c'est également un ouvert de $[0, T^*[$, on obtiendra :

$$\mathcal{E} = [0, T^*[\quad (14)$$

par connexité de $[0, T^*[$.

Le développement limité $(\star)$ de f montre que x est également solution du système différentiel **linéaire** avec second membre :

$$y'(t) - A \cdot y(t) = g(x(t)) \quad (15)$$

et donc on peut appliquer la formule de Duhamel :

$$\forall t \in \mathcal{E}, x(t) = \exp(tA) \cdot x_0 + \int_0^t \exp((t-s)A) \cdot g(x(s)) ds \quad (16)$$

Une petite inégalité triangulaire ainsi qu'une application du lemme 130 fournit :

$$\forall t \in \mathcal{E}, \|x(t)\| \leq C_{\mu_0} e^{-\mu_0 t} \|x_0\| + \int_0^t C_{\mu_0} e^{-\mu_0(t-s)} \|g(x(s))\| ds \quad (17)$$

Or, pour tout $t \in \mathcal{E}$, on a $\|x(t)\| \leq C \leq \delta$, donc $\|g(x(t))\| \leq \|x(t)\|$ et :

$$\forall t \in \mathcal{E}, \|x(t)\| \leq C_{\mu_0} e^{-\mu_0 t} \|x_0\| + \int_0^t C_{\mu_0} e^{-\mu_0(t-s)} \|x(s)\| ds \quad (18)$$

(xxviii). On va montrer le lemme pour toute telle constante C , la majoration de $\|x\|$ pour des constantes C plus grandes étant alors immédiate.

(xxix). Attention, x_0 est amené à changer au cours de la démonstration. Pour épurer les notations déjà très lourdes, j'ai choisi de ne pas écrire explicitement la dépendance de x et x_0 , mais il faut bien avoir conscience de ce qui en dépend, et de ce qui n'en dépend pas.

(xxx). J'insiste : ça n'est pas nécessairement un fermé de $\mathbb{R}$. On va devoir faire joujou avec de la topologie relative, donc ce genre de distinction est importante.

Notons :

$$\begin{aligned}\varphi : \mathcal{E} &\rightarrow \mathbb{R}_+ \\ t &\mapsto \|x(t)\|e^{\mu_0 t}\end{aligned}\tag{19}$$

En multipliant l'inégalité précédente par $e^{\mu_0 t}$, il vient l'inégalité intégrale :

$$\forall t \in \mathcal{E}, \varphi(t) \leq C_{\mu_0} \|x_0\| + \int_0^t C_{\mu_0} \varphi(s) ds\tag{20}$$

On peut appliquer le lemme de Grönwall à φ pour obtenir la majoration :

$$\forall t \in \mathcal{E}, \varphi(t) \leq C_{\mu_0} \|x_0\| e^{C_{\mu_0}^2 \|x_0\| t}\tag{21}$$

Soit, en multipliant par $e^{-\mu_0 t}$:

$$\forall t \in \mathcal{E}, \|x(t)\| \leq \|x_0\| C_{\mu_0} e^{(C_{\mu_0}^2 \|x_0\| - \mu_0)t}\tag{22}$$

Fixons $0 < \eta < C$ tel que :

$$\eta C_{\mu_0}^2 - \mu_0 \leq \mu \wedge \eta C_{\mu_0} \leq \frac{C}{2}\tag{23}$$

Si on choisit x_0 de sorte à avoir $\|x_0\| < \eta$, on a alors bien, pour n'importe quel $t_0 \in \mathcal{E}$:

$$\|x(t_0)\| \leq \frac{C}{2} e^{-\mu t_0} < C e^{-\mu t}\tag{24}$$

et comme x est continue sur $[0, T^*]$, cette majoration reste vraie sur un voisinage de t_0 dans $[0, T^*]$. Donc $\mathcal{E}$ est ouvert dans $[0, T^*]$ et par connexité, ces deux ensembles sont égaux ! Ce raisonnement étant vrai pour tout $\|x_0\| < \eta$, on obtient bien notre résultat. $\square$

Ce lemme permet de conclure presque immédiatement. Prenons μ , C et η comme dans le lemme et soit $\|x_0\| \leq \eta$. Alors la solution maximale issue de x_0 (encore notée x) est bornée en temps positif, donc par théorème des bouts, elle est globale en temps positif. Ainsi, la majoration $\|x(t)\| \leq C e^{-\mu t}$ est vraie sur $\mathbb{R}_+$ tout entier. Comme on peut choisir la valeur de C , on obtient la stabilité de l'équilibre 0. Par ailleurs :

$$x(t) \xrightarrow[t \rightarrow +\infty]{} 0\tag{25}$$

et donc le point 0 est asymptotiquement stable !

Exemple et contre-exemple

On considère l'équation du pendule amorti :

$$\theta''(t) + k\theta'(t) + \sin(\theta(t)) = 0\tag{26}$$

Qu'on peut vectoriser sous la forme :

$$\Theta'(t) = F(\Theta(t))\tag{27}$$

où $k > 0$ et $\Theta(t) = (\theta(t), \theta'(t))$ et $F(x, y) = (y, -\sin(x) - ky)$ Le point $(0, 0)$ est asymptotiquement stable. On calcule la jacobienne de F en $(0, 0)$:

$$J_F(0, 0) = \begin{pmatrix} 0 & 1 \\ -1 & -k \end{pmatrix}\tag{28}$$

dont les valeurs propres, si elles sont réelles, sont de même signe car le déterminant de cette matrice est 1. Comme leur somme est strictement négative (égale à $-k$, qui est la trace de la jacobienne), elles sont strictement négatives. Si elles ne sont pas réelles, elles sont conjuguées et donc leur partie réelle vaut $-k/2 < 0$. Le théorème de Liapounov permet de conclure.

Contrairement au cas linéaire, il faut avoir conscience que le cas des valeurs propres de partie réelle nulle est pathologique. Si on prend $k = 0$ dans notre exemple (c'est-à-dire si l'on considère l'équation du pendule simple) les jacobienes en $(0, 0)$ et en $(\pi, 0)$ sont égales de valeurs propres i et $-i$. Or le premier point est stable (il correspond au cas où le pendule est au repos tête en bas) et le second est instable (le pendule est alors en équilibre précaire, tête en haut). Notez que c'est un peu pénible à prouver (il faut utiliser une intégrale première pour prouver la stabilité, dont on trouvera une expression par exemple dans [Ben14]) et qu'il existe des exemples beaucoup plus élémentaires. J'aime bien celui-ci parce qu'il est tiré de la physique et donne une vision plus intuitive de la chose.

2.15 Théorème de Riesz-Fréchet-Kolmogorov ★★★★★

Recasages possibles : 201, 203, 208, 209, 234, 239

On dispose en analyse d'une certaine batterie de gros théorèmes visant à détecter les compacts dans des espaces de fonctions, au premier desquels on peut citer le théorème d'Arzelà-Ascoli. Notre objectif ici est d'établir le théorème de Riesz-Fréchet-Kolmogorov, qui en est une sorte d'analogue dans les espaces L^p . C'est un fort joli résultat d'analyse fonctionnelle, dont la preuve utilisera plusieurs résultats phares du programme. Son seul défaut pour l'agrégation, c'est qu'il n'est pas évident d'en donner des applications sans aller chercher très loin, mais je pense qu'il est possible de le présenter en soi comme une application satisfaisante des théorèmes du programme.

C'est ce développement que j'ai présenté à l'oral d'analyse !

Soit $d \in \mathbb{N}^*$. Soit $p \in [1, +\infty[$. On notera p' son exposant conjugué et L^p l'espace des fonctions mesurables de $\mathbb{R}^d$ dans $\mathbb{C}$ quotienté par la relation d'égalité presque partout. Pour $f \in L^p$ et $a \in \mathbb{R}^d$, on notera :

$$\check{f} : x \mapsto f(-x) \quad \tau_a f : x \mapsto f(x - a) \quad (1)$$

Notre objectif est de démontrer :

Théorème 132 (Riesz-Fréchet-Kolmogorov, ([HL09], chap 4, théorème 3.8)). Soit $H \subset L^p$. H est une partie relativement compacte de L^p si, et seulement si :

- (i) H est bornée dans L^p .
- (ii) H est une partie équitendue de L^p , c'est-à-dire :

$$\forall \varepsilon > 0, \exists R > 0 : \forall f \in H, \int_{\|x\| \geq R} |f(x)|^p dx \leq \varepsilon \quad (2)$$

- (iii) H satisfait :

$$\forall \varepsilon > 0, \exists \eta > 0 : \forall f \in H, \forall \|a\| \leq \eta, \|\tau_a f - f\|_p \leq \varepsilon \quad (3)$$

Remarque : Avant de se lancer dans la démonstration, je pense qu'il est important de s'arrêter un instant sur les hypothèses (ii) et (iii), afin de bien comprendre ce que l'on cherche à faire. L'hypothèse (iii) est une sorte d'analogue pour des fonctions L^p de la propriété d'équicontinuité qui figure parmi les hypothèses du théorème d'Arzelà-Ascoli : quand on décale les arguments de f de façon suffisamment petite, f et sa translatée deviennent arbitrairement proche au sens L^p (donc une propriété sur la norme infinie est désormais une propriété d'ordre intégral), et ce uniformément en les éléments de H . L'hypothèse (ii) signifie que la masse des fonctions de H se dissipe à l'infini de façon uniforme. Autrement dit, les fonctions de H ont leur masse localisée pour l'essentiel dans un compact qui ne dépend pas de l'élément de H observé. C'est un moyen d'approcher l'hypothèse qui est faite dans le théorème d'Arzelà-Ascoli sur le domaine de définition des fonctions, qui doit être compact. ♦

En avant pour la démonstration ! Rappelons en préambule qu'en vertu du théorème de Riesz-Fisher, l'espace L^p est complet. Ainsi, une partie H est relativement compacte dans L^p si, et seulement si, elle est précompacte (cf [Goub]) : c'est cet angle d'attaque que nous suivrons pour la preuve.

Sens direct

Soit $H \subset L^p$ une partie relativement compacte de L^p , c'est-à-dire précompacte. Soit $\varepsilon > 0$. Par précompacité, il existe des fonctions $f_1, \dots, f_k \in H^k$ telles que :

$$H \subset \bigcup_{i=1}^k B(f_i, \varepsilon) \quad (4)$$

Ceci prouve immédiatement (i). Par ailleurs, une application rapide du théorème de convergence dominée fournit l'existence de $R \geq 0$ tel que :

$$\forall i \in \llbracket 1, k \rrbracket, \int_{\|x\| \geq R} |f_i(x)|^p dx \leq \varepsilon \quad (5)$$

Ainsi, pour $f \in H$, il existe un indice i tel que $\|f - f_i\|_p \leq \varepsilon$ par construction, et donc :

$$\|f \mathbf{1}_{\|x\| \geq R}\|_p \leq \|(f - f_i) \mathbf{1}_{\|x\| \geq R}\|_p + \|f_i \mathbf{1}_{\|x\| \geq R}\|_p \leq 2\varepsilon \quad (6)$$

d'où H vérifie (ii). Enfin, (iii) provient de l'uniforme continuité des translations dans L^p :

$$\exists \eta > 0 : \forall a \leq \eta, \forall i \in \llbracket 1, k \rrbracket, \|\tau_a f_i - f_i\|_p \leq \varepsilon \quad (7)$$

d'où une fois de plus :

$$\|\tau_a f - f\|_p \leq \|\tau_a [f - f_i]\|_p + \|\tau_a f_i - f_i\|_p + \|f_i - f\|_p \leq 3\varepsilon \quad (8)$$

car les applications τ_a sont des isométries de L^p par invariance par translation de la mesure de Lebesgue. On a bien démontré le sens direct du théorème.

Remarque : En oral, je ne prouve que le sens indirect (qui est le plus difficile), le développement étant déjà assez long comme ça... Notez qu'à l'oral on m'a demandé d'expliquer comment faire le sens direct : il faut savoir le faire. $\blacklozenge$

Sens indirect

Réciproquement, on suppose que H a les propriétés (i), (ii) et (iii). La démonstration va consister à approcher les fonctions de H par un ensemble de fonctions continues qui satisfont les hypothèses du théorème d'Arzelà-Ascoli (ensemble qui sera donc relativement compact pour la topologie de la norme infinie) afin d'obtenir des candidats pour les centres des boules avec lesquelles on va recouvrir H . Considérons pour cela (Φ_n) une suite régularisante^(xxxi), et fixons $\varepsilon > 0$.

Il existe d'après (ii) un réel $R > 0$ tel que :

$$\forall f \in H, \|f \mathbf{1}_{\|x\| \geq R}\|_p \leq \varepsilon \quad (\star)$$

(xxxi). C'est-à-dire une suite de fonctions C^∞ à supports compacts, positives, de norme L^1 constante égale à 1 et telle que le support de Φ_n est contenu dans $B(0, \frac{1}{n})$.

Par ailleurs, étant donnés $f \in H$ et $n \in \mathbb{N}$, on a :

$$\|f - f * \Phi_n\|_p^p = \int_{\mathbb{R}^d} \left| f(x) - \int_{\mathbb{R}^d} f(x-y) \Phi_n(y) dy \right|^p dx \quad (9)$$

$$\leq \int_{\mathbb{R}^d} \left| \int_{\mathbb{R}^d} (f(x) - f(x-y)) \Phi_n(y) dy \right|^p dx \quad (10)$$

$$\leq \int_{\mathbb{R}^d} \int_{\mathbb{R}^d} |f(x) - f(x-y)|^p \Phi_n(y) dy dx \quad (11)$$

$$= \int_{\mathbb{R}^d} \int_{\mathbb{R}^d} |f(x) - f(x-y)|^p dx \Phi_n(y) dy \quad (12)$$

En utilisant successivement le fait que Φ_n est de masse 1, puis l'égalité de Jensen pour rentrer la valeur absolue à la puissance p sous l'intégrale (il faut alors considérer qu'on intègre non contre la mesure de Lebesgue mais contre la mesure $\Phi_n(y)dy$, **qui est une mesure de probabilités**) et enfin le théorème de Fubini-Tonelli. De cette dernière égalité, on peut réécrire :

$$\|f - f * \Phi_n\|_p^p \leq \int_{\mathbb{R}^d} \|f - \tau_y f\|_p^p \Phi_n(y) dy \quad (13)$$

$$= \int_{B(0, \frac{1}{n})} \|f - \tau_y f\|_p^p \Phi_n(y) dy \quad (14)$$

$$\leq \sup_{\|y\| \leq \frac{1}{n}} \|\tau_y f - f\|_p^p \quad (15)$$

d'où, d'après (iii), on peut fixer $N \in \mathbb{N}$ tel que :

$$\forall f \in H, \|f - f * \Phi_N\|_p \leq \varepsilon \quad (\star\star)$$

On introduit alors l'ensemble suivant :

$$\tilde{H} := \{(f * \Phi_N)|_{B(0,R)}, f \in H\} \subset C^0(B(0,R), \mathbb{C}) \quad (16)$$

Lemme 133. $\tilde{H}$ est relativement compact dans $(C^0(B(0,R), \mathbb{C}), \|\cdot\|_\infty)$.

Démonstration. On va sans grande surprise chercher à appliquer le théorème d'Arzelà-Ascoli à $\tilde{H}$. Notons que le domaine commun des fonctions de $\tilde{H}$ est compact, et que leur domaine d'arrivée est complet.

Ponctuelle relative compacité – Soit $x \in B(0,R)$. Alors :

$$|f * \Phi_N|(x) \leq \int_{\mathbb{R}^d} |f(x-y) \Phi_N(y)| dy \leq \|(\tau_x f)\|_p \|\Phi_N\|_{p'} = \|f\|_p \|\Phi_N\|_{p'} \quad (\text{xxxii}) \quad (17)$$

par l'inégalité de Hölder. Ainsi, d'après (i), l'ensemble $\{f * \Phi_N(x), f \in H\}$ est borné dans $\mathbb{C}$, d'où la ponctuelle relative compacité.

(xxxii). Φ_N est C_c^∞ , il s'agit donc bien d'une fonction de $L^{p'}$.

Equicontinuité – Soit $(x_1, x_2) \in B(0, R)^2$. Soit $f \in H$. Alors :

$$|f * \Phi_N(x_1) - f * \Phi_N(x_2)| \leq \int_{\mathbb{R}^d} |f(x_1 - y) - f(x_2 - y)| \Phi_N(y) dy \quad (18)$$

$$= \|\tau_{x_1} f - \tau_{x_2} f\|_p \|\Phi_N\|_{p'} \quad (19)$$

$$= \|\tau_{x_1 - x_2} f - f\|_p \|\Phi_N\|_{p'} \quad (20)$$

toujours par l'inégalité de Hölder. D'où, par (iii) :

$$\exists \eta > 0 : \forall (x_1, x_2) \in (\mathbb{R}^d)^2, \forall f \in H, \|x_1 - x_2\| \leq \eta \implies |f * \Phi_N(x_1) - f * \Phi_N(x_2)| \leq \varepsilon \quad (21)$$

et donc l'ensemble $\tilde{H}$ est équicontinu ^(xxxiii).

Ainsi, le théorème d'Arzelà-Ascoli s'applique : $\tilde{H}$ est relativement compact. $\square$

On va enfin pouvoir conclure : l'ensemble $\tilde{H}$ est précompact dans $(C^0(B(0, R), \mathbb{C}), \|\cdot\|_\infty)$. Il existe donc des fonctions $(f_1, \dots, f_k) \in H^k$ telles que :

$$\forall f \in H, \exists i \in \llbracket 1, k \rrbracket : \|(f * \Phi_N - f_i * \Phi_N) \mathbb{1}_{B(0, R)}\|_\infty \leq \frac{\varepsilon}{\lambda_d B(0, R)} \quad \text{(xxxiv)} \quad (\star \star \star)$$

Où λ_d est la mesure de Lebesgue en dimension d . On a ainsi :

$$\forall f \in H, \exists i \in \llbracket 1, k \rrbracket, \|f \mathbb{1}_{B(0, R)} - f_i \mathbb{1}_{B(0, R)}\|_p \leq \varepsilon \quad (22)$$

Recouvrent alors H par des boules dont les (f_i) seront les centres. Soit $f \in H$, et soit $i \in \llbracket 1, k \rrbracket$ comme dans $(\star \star \star)$. On a alors l'inégalité pour tout $x \in \mathbb{R}^d$:

$$|f(x) - f_i(x)| \leq |(f - f_i)(x) \mathbb{1}_{\|x\| \leq R}| + |(f - f_i)(x) \mathbb{1}_{\|x\| \geq R}| \quad (23)$$

$$\begin{aligned} &\leq |(f - f * \Phi_N)(x) \mathbb{1}_{\|x\| \leq R}| + |(f_i - f_i * \Phi_N)(x) \mathbb{1}_{\|x\| \leq R}| \\ &+ |(f * \Phi_N - f_i * \Phi_N)(x) \mathbb{1}_{\|x\| \leq R}| + |f \mathbb{1}_{\|x\| \geq R}| + |f_i \mathbb{1}_{\|x\| \geq R}| \end{aligned} \quad (24)$$

$$\begin{aligned} &\leq |(f - f * \Phi_N)(x)| + |(f_i - f_i * \Phi_N)(x)| \\ &+ |(f * \Phi_N - f_i * \Phi_N)(x) \mathbb{1}_{\|x\| \leq R}| + |f \mathbb{1}_{\|x\| \geq R}| + |f_i \mathbb{1}_{\|x\| \geq R}| \end{aligned} \quad (25)$$

D'où, en passant à la norme p , on obtient :

$$\|f - f_i\|_p \leq 5\varepsilon \quad (26)$$

grâce aux choix de R , N et i . Ainsi :

$$H \subset \bigcup_{i=1}^k B(f_i, 5\varepsilon) \quad (27)$$

ce qui prouve la précompacité de H , et donc sa relative compacité dans L^p ! $\square$

(xxxiii). Et on voit apparaître de façon clair le lien entre l'hypothèse (iii) et l'hypothèse d'équicontinuité du théorème d'Arzelà-Ascoli que j'évoquais plus haut !

(xxxiv). Il pourrait être tentant de se priver de cette constante un peu moche mais prudence. Si on majore simplement par ε , la majoration en norme p va faire apparaître la mesure de $B(0, R)$ qui dépend de ε ! La majoration finale deviendra un peu acrobatique car il ne sera pas clair du tout que le majorant tendra vers 0 lorsque ε tend vers 0...

Annexe : un exemple d'application

Les choses se gâtent ici. Vous l'aurez remarqué, presque toutes les leçons d'agreg s'appellent « Machin machin. Exemples et **applications** ». D'après ce que j'ai compris pendant mes oraux blancs, le jury peut être très pointilleux sur l'aspect « application » des résultats énoncés. Je suis tombée en oral blanc sur une leçon pour laquelle Riesz-Fréchet-Kolmogorov était l'un de mes développements, et pas manqué, le jury m'a demandé de donner une partie concrète de L^p pour laquelle on utilise le théorème de Riesz-Fréchet-Kolmogorov pour montrer qu'elle est compacte^(xxxv). Il s'agit évidemment de ne pas donner une partie trivialement compacte, sinon le théorème n'apporterait rien. Alors voilà ici l'exemple que j'ai trouvé. Je n'en ai pas trouvé de plus élémentaire, mais si vous en connaissez, n'hésitez pas à me contacter, je serais ravie de l'apprendre !

Le cadre sur lequel repose l'exemple est largement hors-programme. Si vous ne vous sentez pas capable de le mener au tableau, **je vous déconseille très fortement de présenter ce développement ou d'inclure le théorème dans votre plan**, à moins que vous n'ayez un exemple concret d'utilisation à proposer.

On considère l'espace $W^{1,p}(\mathbb{R})$ des fonctions de classe L^p sur $\mathbb{R}$ dont la dérivée au sens des distributions est représentée par une fonction L^p . Cet espace est muni de la norme Sobolev :

$$\|f\|_{W^{1,p}} := \|f\|_p + \|f'\|_p \quad (28)$$

Application 134. Toute partie $H \subset W^{1,p}$ équitendue et bornée en norme Sobolev est compacte pour la topologie L^p .

Démonstration. On va appliquer le théorème de Riesz-Fréchet-Kolmogorov. La norme Sobolev dominant la norme L^p , il est clair que H est bornée dans L^p . Comme H est supposée équitendue, il suffit de montrer que H est équicontinue au sens L^p . Fixons $\varepsilon > 0$. Pour $f \in H$ et $a \in \mathbb{R}^*$, on a :

$$\|\tau_a f - f\|_p^p = \int_{\mathbb{R}} \left| \int_t^{t+a} f'(s) ds \right|^p dt \quad (29)$$

$$= \int_{\mathbb{R}} |a|^p \left| \int_t^{t+a} f'(s) \frac{ds}{|a|} \right|^p dt \quad (30)$$

Or la mesure $\frac{ds}{|a|}$ est une mesure de probabilités sur $[t, t+a]$, ce quelque soit t . Comme $x \mapsto |x|^p$ est convexe, on applique l'inégalité de Jensen puis le théorème de Fubini-Tonelli :

$$\|\tau_a f - f\|_p^p \leq |a|^p \int_{\mathbb{R}} \int_{\mathbb{R}} |f'(s)|^p \mathbf{1}_{t \leq s \leq t+a} \frac{ds}{|a|} dt \quad (31)$$

$$= |a|^p \int_{\mathbb{R}} \left| \frac{1}{a} \right|' (s) |f'(s)|^p ds \quad (32)$$

$$= |a|^{p-1} \|f'\|_p^p \quad (33)$$

$$\leq |a|^{p-1} \sup_{h \in H} \|h\|_{W^{1,p}}^p \quad (34)$$

(xxxv). Cela dit, on ne m'a rien demandé de ce genre au vrai oral, mais bon, mieux vaut prévenir que guérir.

D'où, en prenant $\eta \leq \frac{\varepsilon^{\frac{p-1}{p}}}{\sup_{h \in H} \|h\|_{W^{1,p}}}$, on obtient :

$$\forall h \in H, \forall |a| \leq \eta, \|\tau_a h - f\|_p \leq \varepsilon \quad (35)$$

et par Riesz-Fréchet-Kolmogorov, H est compact dans L^p . □

2.16 Théorème taubérien fort ★★

Le bien connu théorème d'Abel radial affirme que si une série de nombres complexes $\sum a_n$ converge, alors la série entière $\sum a_n x^n$ converge vers la même limite lorsque x tend vers 1 par valeurs inférieures. Ceci permet de définir un mode de convergence de série plus faible que la convergence au sens usuel, dite « convergence au sens d'Abel ». On sait que la réciproque est fautive, ce que l'on voit au-travers de contre-exemples simples comme $\sum (-1)^n x^n$. Il reste pourtant naturel de se demander si l'on peut avoir des conditions suffisantes simples à vérifier pour obtenir la réciproque. C'est l'objet des théorèmes taubériens. Nous allons montrer dans ce développement l'un d'entre eux, dit « théorème taubérien fort ».

Théorème 135 (Taubérien fort ([Goub], chap IV, problème 19)). Soit (a_n) une suite de nombres complexes telle que :

$$|a_n| = O\left(\frac{1}{n}\right) \quad (1)$$

On suppose de plus que la série de terme général (a_n) converge au sens d'Abel, c'est-à-dire :

$$\sum_{n=0}^{+\infty} a_n x^n \xrightarrow[x < 1]{x \rightarrow 1} l \in \mathbb{C} \quad (2)$$

Alors la série des (a_n) converge, et sa somme est l .

On pose :

$$F : [0, 1[\rightarrow \mathbb{C} \quad (3)$$

$$x \mapsto \sum_{n=0}^{+\infty} a_n x^n$$

Dans le cadre de la démonstration, on introduit l'ensemble $\mathcal{F}$ des fonctions $f : [0, 1] \rightarrow \mathbb{C}$ telles que :

1. $\forall x \in [0, 1[, \sum a_n f(x^n)$ converge.

2. $\sum_{n=0}^{+\infty} a_n f(x^n) \xrightarrow[x < 1]{x \rightarrow 1} l$

Dans toute la suite, on supposera pour simplifier que $l = 0$, le cas général s'en déduisant aisément. Dans ce cas, on gagne notamment le fait que $\mathcal{F}$ est un espace vectoriel.

Reformulation du problème

On considère $g := \mathbb{1}_{[\frac{1}{2}, 1]}$. La démonstration du théorème va se résumer à tester l'appartenance de g à l'ensemble $\mathcal{F}$:

Proposition 136. Si $g \in \mathcal{F}$, alors $\sum_{n=0}^{+\infty} a_n = 0$.

Démonstration. On a, dans ce cas, pour $x < 1$:

$$\sum_{n=0}^{+\infty} a_n g(x^n) = \sum_{n=0}^{\lceil -\frac{\log(2)}{\log(x)} \rceil} a_n \quad (4)$$

$$\xrightarrow[x < 1]{x \rightarrow 1} \sum_{n=0}^{+\infty} a_n \quad (5)$$

$$(6)$$

Mais puisque $g \in \mathcal{F}$, cette dernière limite est égale à 0. $\square$

Tout l'enjeu de la démonstration va donc être de montrer que $g \in \mathcal{F}$. Remarquons en préambule que pour $x < 1$, la série des $a_n g(x^n)$ converge puisqu'il s'agit de ne sommer qu'un nombre fini de termes non nuls.

Les polynômes nuls en 0 sont dans $\mathcal{F}$

Faisons une première observation : les polynômes nuls en 0 appartiennent à $\mathcal{F}$. En effet, si m_k est la fonction monomiale $x \mapsto x^k$ avec $k \in \mathbb{N}^*$, on a pour $0 \leq x < 1$ que la série des $a_n m_k(x^n)$ est convergente et :

$$\sum_{n=0}^{+\infty} a_n m_k(x^n) = F(x^k) \xrightarrow[x < 1]{x \rightarrow 1} 0 \quad (7)$$

donc $m_k \in \mathcal{F}$. Par linéarité, on en déduit immédiatement que les polynômes s'annulant en 0 appartiennent à $\mathcal{F}$.

Ceci nous pousse à chercher à approximer g par de telles fonctions polynomiales.

Approximation de l'indicatrice par des polynômes

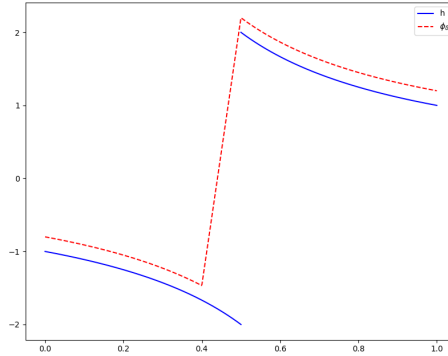
On va se ramener à travailler avec des fonctions polynomiales, dont on pourra extraire plus d'information. Fixons $\varepsilon > 0$.

Lemme 137. *Il existe p et P deux fonctions polynomiales telles que :*

1. $p \leq \frac{g-x}{x(1-x)} \leq P$
2. $\|P-p\|_1 \leq \varepsilon$

Démonstration. Commençons par remarquer que la fonction $h : x \mapsto \frac{g(x)-x}{x(1-x)}$ se prolonge par continuité en 0 et 1, avec les valeurs -1 et 1 respectivement. Soit $0 < \delta < \frac{1}{2}$. On construit une application affine a_δ par interpolation, on imposant $a_\delta(\frac{1}{2} - \delta) = h(x - \delta)$ et $a_\delta(\frac{1}{2}) = h(\frac{1}{2}) = 2$. On approche alors h par au-dessus de la façon suivante :

$$\varphi_\delta : x \mapsto \begin{cases} \max\{h(x), a_\delta(x)\} + \frac{\varepsilon}{4} & \text{si } \frac{1}{2} - \delta \leq x \leq \frac{1}{2} \\ h(x) + \frac{\varepsilon}{4} & \text{sinon} \end{cases} \quad (8)$$



La fonction φ_δ est continue (ce sont des vérifications élémentaires). Ainsi, on peut utiliser le théorème de Weierstraß : il existe P_δ une fonction polynomiale telle que $\|P_\delta - \varphi_\delta\|_\infty \leq \frac{\varepsilon}{4}$. Par construction, $P_\delta > h$, et on a :

$$\|P_\delta - h\|_1 = \int_0^{\frac{1}{2}-\delta} P(x) - h(x) dx + \int_{\frac{1}{2}}^1 P(x) - h(x) dx + \int_{\frac{1}{2}-\delta}^{\frac{1}{2}} P(x) - h(x) dx \quad (9)$$

$$\leq \frac{\varepsilon}{4}(1 - \delta) + \int_{\frac{1}{2}-\delta}^1 |P_\delta(x) - \varphi_\delta(x)| dx + \int_{\frac{1}{2}-\delta}^{\frac{1}{2}} |\varphi_\delta(x) - h(x)| dx \quad (10)$$

$$\leq \frac{\varepsilon}{4} \left(1 - \frac{\delta}{2}\right) + \|(\varphi_\delta - h)\mathbb{1}_{[\frac{1}{2}-\delta, \frac{1}{2}]}\|_1 \quad (11)$$

Ainsi, en passant à la limite supérieure lorsque $[\delta \rightarrow 0]$, et par convergence dominée, il vient :

$$\limsup_{\delta \rightarrow 0} \|P_\delta - h\|_1 \leq \frac{\varepsilon}{4} \quad (12)$$

Donc il existe un certain δ tel que $\|P_\delta - h\|_1 \leq \frac{\varepsilon}{2}$. On construit de même une approximation polynomiale p par en-dessous. On a bien alors $\|P - p\|_1 \leq \varepsilon$. $\square$

Introduisons les fonctions polynomiales suivantes :

$$T : x \mapsto x + x(1 - x)P(x) \quad t : x \mapsto x + x(1 - x)t(x) \quad (13)$$

Par construction, $t \leq g \leq T$ et $T(0) = t(0) = 0$. C'est grâce à cet encadrement qu'on va réussir à montrer que $g \in \mathcal{F}$.

g est dans $\mathcal{F}$

On est maintenant en mesure de montrer que g est dans $\mathcal{F}$, et donc d'en déduire le théorème. On a déjà vu que pour $x < 1$, la série $\sum a_n g(x^n)$ ne contient qu'un nombre fini de termes non nuls, et donc converge. Il faut, et il suffit donc de montrer que cette somme tend vers 0 lorsque x tend vers 1 par valeurs inférieures. Calculons :

$$\left| \sum_{n=0}^{+\infty} a_n g(x^n) - \sum_{n=0}^{+\infty} a_n t(x^n) \right| \leq \sum_{n=0}^{+\infty} |a_n| (g - t)(x^n) \quad (14)$$

$$\leq \sum_{n=0}^{+\infty} |a_n| (T - t)(x^n) \quad (15)$$

$$\leq \sum_{n=0}^{+\infty} |a_n| x^n (1 - x^n) (P - p)(x^n) \quad (16)$$

On va maintenant exploiter l'identité remarquable :

$$(1 - x^n) = (1 - x) \sum_{j=0}^{n-1} x^j \leq n(1 - x) \quad (17)$$

et comme a_n est un grand O de $\frac{1}{n}$, on obtient l'existence d'un réel M tel que pour tout n :

$$|a_n| (1 - x^n) \leq M(1 - x) \quad (18)$$

En injectant ceci dans l'inégalité, on trouve :

$$\left| \sum_{n=0}^{+\infty} a_n g(x^n) - \sum_{n=0}^{+\infty} a_n t(x^n) \right| \leq M(1 - x) \sum_{n=0}^{+\infty} x^n (P - p)(x^n) \quad (19)$$

Il nous reste à prouver un dernier lemme pour étudier le comportement asymptotique du majorant :

Lemme 138. *Soit q une fonction polynomiale. On a :*

$$(1 - x) \sum_{n=0}^{+\infty} x^n q(x^n) \xrightarrow[\substack{x \rightarrow 1 \\ x < 1}]{1} \int_0^1 q(s) ds \quad (20)$$

Démonstration. On le montre pour des monômes, le résultat général s'en déduisant par linéarité. Pour $k \in \mathbb{N}$:

$$(1 - x) \sum_{n=0}^{+\infty} x^n m_k(x^n) = (1 - x) \sum_{n=0}^{+\infty} x^{(k+1)n} \quad (21)$$

$$= \frac{1 - x}{1 - x^{k+1}} \quad (22)$$

$$= \frac{1}{\sum_{j=0}^k x^j} \quad (23)$$

$$\xrightarrow[\substack{x \rightarrow 1 \\ x < 1}]{1} \frac{1}{k + 1} \quad (24)$$

$$= \int_0^1 m_k(s) ds \quad (25)$$

□

On en déduit finalement que notre majorant tend, lorsque x tend vers 1 par valeurs inférieures, vers M fois l'intégrale de $P - p$, qui n'est autre que $\|P - p\|_1$! Autrement dit, en passant à la limite supérieure :

$$\limsup_{\substack{x \rightarrow 1 \\ x < 1}} \left| \sum_{n=0}^{+\infty} a_n g(x^n) - \sum_{n=0}^{+\infty} a_n t(x^n) \right| \leq M\varepsilon \quad (26)$$

Comme ceci vaut pour tout ε , et comme $t \in \mathcal{F}$, on obtient bien :

$$\sum_{n=0}^{+\infty} a_n g(x^n) \xrightarrow[\substack{x \rightarrow 1 \\ x < 1}]{} 0 \quad (27)$$

et donc $g \in \mathcal{F}$! Le théorème taubérien est ainsi démontré. □

Une application aux séries de Fourier

Recasages possibles : 223, 224, 230, 241, 243

C'est Thomas Cavalazzi qui a remarqué cette application, qui est somme toute des plus sympathiques.

Application 139. Soit f une fonction continue 2π -périodique. On suppose que la suite des coefficients de Fourier de f est un grand $O(\frac{1}{|n|})$ lorsque $|n|$ tend vers $+\infty$. Alors la série de Fourier de f converge simplement vers f .

Démonstration. On introduit le noyau de Poisson :

$$\forall r \in [0, 1[, \quad P_r : x \mapsto \sum_{n \in \mathbb{Z}} e^{inx} r^{|n|} \quad (28)$$

C'est une approximation de l'unité lorsque r tend vers 1 par valeurs inférieures (cf [Amr]). Comme f est uniformément continue et bornée, on a que $P_r * f$ converge uniformément vers f lorsque r tend vers 1 par valeurs inférieures. Or :

$$P_r * f(x) = \sum_{n \in \mathbb{Z}} c_n(f) e^{inx} r^{|n|} \quad (29)$$

En posant $a_n := c_n(f) e^{inx}$, on a par hypothèse $a_n = O(\frac{1}{n})$. Donc, comme la série de terme général $(a_n + a_{-n})$ converge au sens d'Abel, on a par le théorème taubérien qu'elle converge au sens usuel, c'est-à-dire que $f(x)$ est la somme de la série de Fourier de f en x . Pas mal, n'est-ce pas ? □

Chapitre 3

Développements mixtes



* « Rencontre du troisième type »

3.1 Déterminant circulant et suite de polygones ★

Dans ce développement, on va utiliser des outils de réduction d'endomorphismes pour étudier un algorithme permettant d'approcher l'isobarycentre d'un polygone du plan, aussi tarabiscoté soit-il ! Il s'agit du développement 62 «Suite de polygones» de [IP24] (attention à bien avoir la deuxième édition, la première étant interdite au concours car elle contient des plans de leçon).

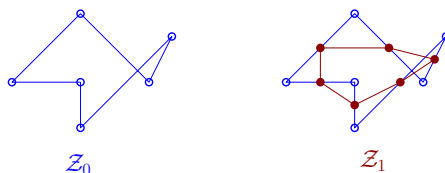
On assimilera le plan euclidien à $\mathbb{C}$ pour pouvoir faire de la géométrie plane avec des coordonnées complexes. Un polygone à n sommets sera assimilé à la liste ordonnée de ses sommets, chacun représenté par son affixe complexe. En d'autres termes, un polygone à n sommets sera identifié à un élément de $\mathbb{C}^n$. Pour plus de commodités d'écriture, on munira $\llbracket 1, n \rrbracket$ des opérations de $\mathbb{Z}/n\mathbb{Z}$, de sorte à pouvoir écrire pour un polygone $\mathcal{Z} = (z_1, \dots, z_n)$ et $i \in \llbracket 1, n \rrbracket$ z_{i+1} et z_{i-1} sans se soucier des termes de bords (on a avec cette convention $z_0 = z_n$ et $z_1 = z_{n+1}$).

Soit $\mathcal{Z} = (z_1, \dots, z_n)$ un polygone à n sommets. On construit une suite de polygones de la façon suivante : on pose $\mathcal{Z}_0 := \mathcal{Z}$ et, étant construit $\mathcal{Z}_l = (z_i^{(l)})$ au rang l , on définit $\mathcal{Z}_{l+1}$ comme le polygone dont le i -ième sommet est le milieu du segment $[z_i^{(l)}; z_{i+1}^{(l)}]$. En d'autres termes, on a :

$$\forall l \in \mathbb{N}, \mathcal{Z}_{l+1} = \frac{1}{2} \begin{pmatrix} z_1^{(l)} + z_2^{(l)} \\ \vdots \\ z_n^{(l)} + z_{n+1}^{(l)} \end{pmatrix} \quad (1)$$

On souhaite démontrer :

Théorème 140. La suite $(\mathcal{Z}_l)$ converge, et sa limite est $g \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix}$, où g est l'affixe de l'isobarycentre de $\mathcal{Z}$.



Etude de la matrice circulante

Dans cette première partie, on va étudier la forme réduite de la matrice circulante, qu'on exploitera dans un second temps pour prouver la convergence de la méthode.

Définition 141 (Matrice circulante universelle). On appelle matrice circulante universelle de taille n , à coefficients dans $\mathbb{Z}[X_0, \dots, X_{n-1}]$, la matrice :

$$C(X_0, \dots, X_{n-1}) := \begin{pmatrix} X_0 & X_1 & \dots & X_{n-1} \\ X_{n-1} & X_0 & \dots & X_1 \\ \vdots & \vdots & \ddots & \vdots \\ X_1 & \dots & X_{n-1} & X_0 \end{pmatrix} \quad (2)$$

Proposition 142 (Diagonalisation de la matrice circulante). Posons

$$P(X_0, \dots, X_{n-1}) := T^n + \sum_{i=0}^{n-1} X_i T^i \in \mathbb{Z}[T][X_1, \dots, X_n] \quad (3)$$

le polynôme unitaire générique de degré n . Il existe une matrice $F \in GL_n(\mathbb{C})$ inversible telle que $F^{-1}C(X_0, \dots, X_{n-1})F$ soit diagonale, avec pour coefficients diagonaux les $(P(X_0, \dots, X_n)(\omega^k))_{0 \leq k \leq n-1}$, avec $\omega := e^{\frac{2i\pi}{n}}$.

Démonstration. On considère la matrice $J := C(0, 1, 0, 0, \dots, 0)$. J est une matrice de permutation associée à un n -cycle. En conséquence, c'est une matrice d'ordre n , et donc son polynôme minimal est $X^n - 1$, qui est scindé à racines simples dans $\mathbb{C}$. Elle est donc diagonalisable et ses valeurs propres sont les $(\omega_k)_{0 \leq k \leq n-1}$. Par ailleurs, on a clairement $C(X_0, \dots, X_{n-1}) = P(X_0, \dots, X_{n-1})(J)$. Ainsi, si $F \in GL_n(\mathbb{C})$ est une matrice telle que :

$$F^{-1}JF = \begin{pmatrix} \omega^0 & & \\ & \ddots & \\ & & \omega^{n-1} \end{pmatrix} \quad (4)$$

On a immédiatement :

$$F^{-1}C(X_0, \dots, X_{n-1})F = \begin{pmatrix} P(X_0, \dots, X_{n-1})(\omega^0) & & \\ & \ddots & \\ & & P(X_0, \dots, X_{n-1})(\omega^{n-1}) \end{pmatrix} \quad (5)$$

et le résultat s'en suit. □

Remarque : De ce résultat, on déduit la valeur classique du déterminant circulaire. Par ailleurs, l'avantage de raisonner avec des matrices à coefficients dans des anneaux de polynôme et qu'on peut déduire sans aucun effort la valeur du polynôme caractéristique de $C(X_0, \dots, X_{n-1})$ en remplaçant X_0 par $X - X_0$ et en multipliant les autres indéterminées par -1 . L'expression diagonale obtenue est toujours valable et on peut librement y appliquer le déterminant ! ◆

Convergence de la méthode

Il nous reste à montrer le théorème. Remarquons que la suite (Z_l) construite plus haut est une suite géométrique de matrice de transition $M := C(0, 1/2, 0, 0, \dots, 0, 1/2)$:

$$\begin{cases} Z_0 = Z \\ Z_{l+1} = M Z_l \end{cases} \quad (6)$$

Ainsi, on a son expression explicite :

$$Z_l = M^l Z \quad (7)$$

et en utilisant la matrice F précédemment introduite, on peut écrire :

$$Z_l = F \begin{pmatrix} P(0, 1/2, 0, \dots, 0, 1/2)(\omega^0) & & \\ & \ddots & \\ & & P(0, 1/2, 0, \dots, 0, 1/2)(\omega^{n-1}) \end{pmatrix}^l F^{-1} Z \quad (8)$$

$$= F \begin{pmatrix} \left(\frac{1+\omega^0}{2}\right)^l & & \\ & \ddots & \\ & & \left(\frac{1+\omega^{n-1}}{2}\right)^l \end{pmatrix} F^{-1} Z \quad (9)$$

or pour $1 \leq k \leq n-1$, on a :

$$\left| \frac{1+\omega^k}{2} \right| = |\omega^{\frac{k}{2}}| \left| \frac{\omega^{-\frac{k}{2}} + \omega^{\frac{k}{2}}}{2} \right| = \cos\left(\frac{k\pi}{n}\right) < 1 \quad (10)$$

et bien évidemment, pour $k=0$, $\frac{\omega^0+1}{2} = 1$. On a donc la convergence :

$$Z_l \xrightarrow[l \rightarrow +\infty]{} F \underbrace{\begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & & & \\ \vdots & & 0 & \\ 0 & & & \end{pmatrix}}_{=:L} F^{-1} Z \quad (11)$$

On remarque alors, en notant Z_∞ la limite de la suite (Z_l) et en passant à la limite dans la relation $Z_{l+1} = M Z_l$, que Z_∞ est un vecteur propre associé à la valeur propre 1 de M . Or on sait par réduction de M que cet espace propre est de dimension 1, et un calcul immédiat montre que le vecteur $(1, \dots, 1)$ est un point fixe de M . Il existe donc $\lambda \in \mathbb{C}$ tel que la suite (Z_l) tende vers $(\lambda, \dots, \lambda)$. Tout ce qu'il nous reste à faire est de montrer que $\lambda = g$ l'abscisse de l'isobarycentre de Z . On remarque, pour $l \in \mathbb{N}$:

$$\frac{1}{n} \sum_{k=1}^n z_k^{(l+1)} = \frac{1}{n} \sum_{k=1}^n \frac{z_k^{(l)} + z_{k+1}^{(l)}}{2} = \frac{1}{n} \sum_{k=1}^n z_k^{(l)} \quad (12)$$

c'est-à-dire que l'isobarycentre de Z_{l+1} est également l'isobarycentre de Z_l . Par récurrence immédiate, la suite des isobarycentres des (Z_l) est constante d'abscisse g , et comme l'application qui à un polygone à n sommets associe son isobarycentre est continue, on a :

$$\frac{1}{n} n \lambda = \lambda = g \quad (13)$$

et la boucle est bouclée !

□

3.2 Sous-groupes petits de $GL_n(\mathbb{R})$ ★★

Recasages possibles : 106, 155, 214, 215, 241, 243

L'objectif de ce développement est de démontrer que $GL_d(\mathbb{R})$ ne contient pas de sous-groupe non-trivial arbitrairement petit. La preuve repose essentiellement sur la construction d'un logarithme matriciel au voisinage de l'identité. J'ai choisi une approche algébrique, très différente de celle de Gurwinneuse. En conséquence, je ne recase pas du tout ce développement aux mêmes endroits.

Soit $d \in \mathbb{N}^*$ un entier non-nul. Dans toute la suite, on fixe $||| \cdot |||$ une norme matricielle. On notera $\mathcal{B}(A, \eta)$ la boule ouverte de centre A et de rayon η dans $GL_d(\mathbb{R})$. Nous allons montrer le résultat topologique suivant :

Théorème 143. *Il existe $\eta > 0$ tel que si G est un sous-groupe de $GL_d(\mathbb{R})$ contenu dans $\mathcal{B}(I_d, \eta)$, alors $G = \{I_d\}$.*

Logarithme matriciel

Le cœur de la démonstration repose sur l'existence d'une inverse à l'exponentielle de matrice au voisinage de l'identité. Je propose deux méthodes différentes pour y arriver, à choisir en fonction de la leçon présentée.

Avec une série matricielle ([Rom], §24.4)

Dans $\mathbb{R}$, l'exponentielle est définie par la même série entière que l'exponentielle de matrice. L'exponentielle admet un inverse bien connu : le logarithme. Dans les réels, on sait développer le logarithme en série entière au voisinage de 1 :

$$\forall |x| < 1, \log(1+x) = \sum_{n=1}^{+\infty} (-1)^{n+1} \frac{x^n}{n} \quad (1)$$

Cette égalité va nous servir d'heuristique pour construire le logarithme matriciel.

Définition 144 (Logarithme matriciel). *Pour tout $A \in \mathcal{B}(0, 1)$, on pose :*

$$\log(I_d + A) := \sum_{n=1}^{+\infty} (-1)^{n+1} \frac{A^n}{n} \quad (2)$$

Cette série est absolument convergente.⁽ⁱ⁾

Le logarithme ainsi défini a exactement les propriétés attendues :

(i). En travaillant sur le rayon spectral, on peut en réalité montrer que cette série est absolument convergente même en l'étendant à toutes les matrices de rayon spectral strictement plus petit que 1. Cela peut servir en fonction des recasages.

Théorème 145. Pour tout $A \in \mathcal{B}(0, 1)$, on a :

$$\exp(\log(I_d + 1))^{(ii)} = I_d + A \quad (3)$$

Démonstration. Fixons $A \in \mathcal{B}(0, 1)$. On définit alors :

$$\forall t \in [0, 1], \varphi(t) := \log(I_d + tA) \quad (4)$$

La fonction φ est bien définie, continue et dérivable sur $[0, 1]$. Pour le démontrer, posons :

$$\varphi_{abs}(z) = \sum_{n=1}^{+\infty} (-1)^{n+1} \frac{|||A^n|||}{n} z^n \quad (5)$$

Cette série entière a un rayon de convergence supérieur ou égal à $1/|||A||| > 1$. Aussi, φ_{abs} et φ'_{abs} convergent normalement sur tout compact du disque ouvert de convergence. Par théorème de dérivation d'une limite, on a que φ est dérivable sur $[0, 1]$ de dérivée :

$$\varphi'(t) = \sum_{n=1}^{+\infty} (-1)^{n+1} t^{n-1} A^n = (I_d + tA)^{-1} A^{(iii)} \quad (6)$$

On peut alors considérer la fonction.

$$\psi : t \mapsto \exp(\varphi(t)) \quad (7)$$

Montrons que ψ est une fonction affine de t . On est tenté de dériver ψ et de montrer que sa dérivée est constante. ψ est en effet dérivable et sa dérivée est donnée par :

$$\psi'(t) = \varphi'(t) \exp(\varphi(t)) \quad (8)$$

Après réécriture, on obtient :

$$(I_d + tA)\psi'(t) = A \exp(\varphi(t)) \quad (9)$$

et ψ' est encore dérivable comme composée et produit de fonctions dérivables. On a alors :

$$A\psi'(t) + (I_d + tA)\psi''(t) = A\varphi'(t) \exp(\varphi(t)) = A\psi'(t) \quad (10)$$

Et donc :

$$(I_d + tA)\psi''(t) = 0 \quad (11)$$

soit $\psi''(t) = 0$ car $(I_d + tA)$ est inversible. Il suffit d'intégrer deux fois :

$$\psi(t) = t\psi'(0) + \psi(0) = tA + I_d \quad (12)$$

Ce qui prouve le théorème! □

(ii). Attention au sens de la composition!

(iii). Cette deuxième égalité est un résultat bien connu vrai dans les algèbres de Banach

Par le calcul différentiel ([Rou15], exo 65)

On va procéder par inversion local, en montrant que l'exponentielle de matrice réalise un C^1 difféomorphisme d'un voisinage de 0 dans un voisinage de l'identité. Pour cela, on va commencer par calculer la différentielle du déterminant. Posons :

$$f_n : M \mapsto M^n \quad (13)$$

de sorte à avoir :

$$\exp = \sum_{n=0}^{+\infty} \frac{f_n}{n!} = \lim_{N \rightarrow +\infty} \sum_{n=0}^N \frac{f_n}{n!} \quad (14)$$

Les fonction f_n sont faciles à différentier :

$$df_n(M) \cdot H = \sum_{k=0}^{n-1} M^k H M^{n-1-k} \quad (15)$$

Ainsi, par sous-multiplicativité de la norme d'opérateur, on trouve :

$$\|df_n(M)\| \leq (n-1)\|M\|^{n-1} \quad (16)$$

Il vient que la série des $df_n/n!$ ^(iv) est normalement convergente sur les compacts de $\mathbb{M}_d(\mathbb{R})$, donc absolument convergente par complétude. Par théorème de différentiation d'une limite, l'exponentielle est de classe C^1 et sa différentielle est obtenue en différentiant terme à terme sous le signe somme. En particulier, sa différentielle en 0 est donnée par :

$$d[\exp](0) \cdot H = \sum_{n=0}^{+\infty} df_n(0) \cdot H = H \quad (17)$$

C'est-à-dire :

$$d[\exp](0) = I_d \quad (18)$$

En particulier, celle-ci est inversible et donc le théorème d'inversion local s'applique à $\exp$: il existe des voisinages U et V respectivement de 0 dans $\mathbb{M}_d(\mathbb{R})$ et de I_d dans $GL_d(\mathbb{R})$ et une fonction $\log : V \rightarrow U$ telle que $\log = \exp|_U^{-1}$.

Sous-groupes petits de $GL_d(\mathbb{R})$, [MT]

Le logarithme matriciel va nous fournir un candidat pour η tel qu'il n'y a pas de sous-groupe non-trivial de diamètre inférieur à 2η . En effet, les résultats précédents donnent l'existence de U un voisinage ouvert de 0 dans $\mathbb{M}_d(\mathbb{R})$ et V un voisinage ouvert de I_d dans $GL_d(\mathbb{R})$ tels que l'exponentielle de matrice induise une bijection de U dans V , dont on note $\log$ la bijection réciproque. Quitte à restreindre U et V , on peut toujours les supposer bornés. Posons alors $U' := \frac{1}{2}U$ et $V' := \exp(U')$. Par l'absurde, on suppose donné G un sous-groupe de $GL_d(\mathbb{R})$ non-trivial contenu dans V' . Soit $A \in G \setminus \{I_d\}$. On dispose alors de $M := \log(A) \in U'$ tel que $\exp(M) = A$. Puisque U est borné et que la définition de U' implique qu'une somme de deux éléments de U' est dans U (mais pas nécessairement dans U' bien sûr), il existe $k_0 \in \mathbb{N}^*$ tel que $k_0 M \in U \setminus U'$. Mais comme M commute avec elle-même, on a que $\exp(k_0 M) = A^{k_0} \in V'$. L'exponentielle étant bijective de U dans V , on a alors que $k_0 M$ est l'unique logarithme de A^{k_0} et devrait donc être dans U' ... Absurde! G est donc trivial.

(iv). Attention, on parle bien de la série des fonctions $df_n : \mathbb{M}_d(\mathbb{R}) \rightarrow \mathcal{L}(\mathbb{M}_d(\mathbb{R}))$ et non pas de la série des $df_n(M) : \mathbb{M}_d(\mathbb{R}) \rightarrow \mathbb{M}_d(\mathbb{R})$!

3.3 Lemme de Morse ★★ ★

Recasages possibles : 157, 171, 206, 214, 215, 218

On sait, de par le théorème de Sylvester, qu'une forme quadratique réelle q admet toujours une base q -orthogonale de l'espace dans laquelle la matrice de q a pour forme :

$$\begin{pmatrix} I_p & 0 & 0 \\ 0 & -I_q & 0 \\ 0 & 0 & 0 \end{pmatrix} \quad (1)$$

En d'autres termes, pour un vecteur x qu'on écrirait $x = \sum x_1 e_1 + \dots x_d e_d$ dans une telle base, on obtiendrait :

$$q(x) = x_1^2 + \dots + x_p^2 - x_{p+1}^2 - \dots - x_{p+q}^2 \quad (2)$$

On cherche maintenant à s'intéresser à des objets un peu plus sauvages que des formes quadratiques : des fonctions deux fois différentiables nulles en 0. Si f est une telle fonction, qu'on suppose par ailleurs nulle et de différentielle nulle en 0, la formule de Taylor-Young donne :

$$\forall x \in \mathbb{R}^d, f(x) = d^2 f(0) \cdot (x, x) + o(\|x\|^2) \quad (3)$$

Ce qui semble indiquer qu'au voisinage de 0, f se comporte sensiblement comme sa hessienne (qui est une forme quadratique d'après le lemme de Schwarz). On est alors en droit de se demander s'il existe toujours un changement de coordonnées, qu'on ne cherchera plus linéaire (ce serait bien trop restrictif!) mais C^1 , pour lequel l'équation (2) tienne encore. On va voir ici qu'en rajoutant quelques hypothèses sur f , c'est effectivement le cas !

Remarque : Pour un recasage dans la 171, je pense qu'il est bon de prendre un peu de temps pour raconter cette petite histoire au début, quitte à aller un peu plus vite sur les parties analytiques, car sinon il est facile de passer à côté de l'aspect "formes quadratiques" qui se cache derrière ce développement et le recasage pourrait avoir l'air nébuleux... ♦

On considère un entier d .

Etant donnée une application f n -fois différentiable en a de $\mathbb{R}^d$ dans $\mathbb{R}$, on notera $d^n f(a)$ sa différentielle n -ième, qui s'identifie à une forme n -linéaire. Dans le cas particulier $n = 2$, la différentielle seconde s'identifie à une forme bilinéaire symétrique de matrice dans la base canonique de $\mathbb{R}^d$ $H_f(a)$, de sorte qu'on puisse écrire :

$$d^2 f(a) \cdot (x, x) = {}^t x \cdot H_f(a) \cdot x \quad (4)$$

où l'on assimile x avec le vecteur colonne de ses coordonnées dans la base canonique. Nous allons démontrer ici :

Théorème 146 (lemme de Morse ([Rou15], exercice 114)). Soit $f : U \rightarrow \mathbb{R}$ une application de classe C^3 , avec U un ouvert de $\mathbb{R}^d$. On fait les hypothèses suivantes :

1. $f(0) = 0$ et $df(0) = 0$.
2. La forme quadratique $x \mapsto d^2 f(0) \cdot (x, x)$ est non dégénérée. On notera $(p, d - p)$ sa signature.

Il existe V et W deux voisinages ouverts de l'origine et $u : V \rightarrow W$ un C^1 -difféomorphisme envoyant 0 sur 0 tel que, en notant (u_i) les composantes de $u^{(v)}$, on ait :

$$\forall x \in V, f(x) = u_1(x)^2 + \dots + u_p(x)^2 - u_{p+1}(x)^2 - \dots - u_d(x)^2 \quad (5)$$

Réduction des formes quadratiques et calcul différentiel

Avant de se lancer à l'assaut du lemme de Morse, on va chercher à mieux comprendre la réduction des formes quadratiques sous le prisme du calcul différentiel, car on l'aura compris, c'est bien au croisement de ces deux domaines que se trouve notre résultat.

Lemme 147 ([Rou15], exercice 66). Soit $A_0 \in GL_d(\mathbb{R}) \cap S_d(\mathbb{R})$ une matrice symétrique inversible. Il existe V un voisinage ouvert de A_0 dans $S_d(\mathbb{R})$ et $\Phi : V \rightarrow GL_d(\mathbb{R})$ un C^∞ difféomorphisme tel que :

$$\forall A \in V, A = {}^t\Phi(A) \cdot A_0 \cdot \Phi(A) \wedge \Phi(A) = I_d \quad (6)$$

Remarque : La forme matricielle de ce lemme peut faussement donner l'impression qu'il s'agit d'un résultat d'algèbre linéaire, mais celui-ci se comprend beaucoup mieux en utilisant le langage des formes quadratiques. En effet, une matrice symétrique peut s'interpréter comme la matrice d'une forme quadratique plutôt que d'une application linéaire. Ici, on encadre la matrice A_0 entre une matrice inversible et sa transposée, ce qui correspond au changement de bases pour des formes quadratiques, et non pour des applications linéaires. Ce qu'exprime ce lemme, c'est que les matrices « proches » de A_0 (qui sont dans le voisinage V) sont congruentes à A_0 , c'est-à-dire qu'elles représentent la même forme quadratique dans des bases différentes, et que par ailleurs le changement de base peut se faire de manière lisse. ♦

Démonstration. Pas de grand mystère, ce lemme ne cache pas le chemin de sa démonstration : la solution va venir du théorème d'inversion locale. On considère l'application :

$$\begin{aligned} \psi : GL_d(\mathbb{R}) &\rightarrow S_d(\mathbb{R}) \\ M &\mapsto {}^tM \cdot A_0 \cdot M \end{aligned} \quad (7)$$

Cette application est C^∞ , car toutes ses composantes sont polynomiales en les coefficients de M . Par ailleurs, on calcule sa différentielle à vue :

$$\forall M \in GL_d(\mathbb{R}), \forall H \in \mathbb{M}_d(\mathbb{R}), d\psi(M) \cdot H = {}^tH \cdot A_0 \cdot M + {}^tM \cdot A_0 \cdot H \quad (8)$$

On s'intéresse au noyau de cette différentielle en l'identité. Prenons $H \in \mathbb{M}_d(\mathbb{R})$. On a :

$$H \in \text{Ker}(d\psi(I_d)) \iff {}^tH \cdot A_0 = -A_0 \cdot H \quad (9)$$

$$\iff {}^t(A_0 \cdot H) = -A_0 \cdot H \quad \text{car } A_0 \text{ est symétrique} \quad (10)$$

$$\iff H \in A_0^{-1} \mathfrak{A}_d(\mathbb{R}) \quad (11)$$

avec $\mathfrak{A}_d(\mathbb{R})$ l'ensemble des matrices antisymétriques de taille d . Or, on a la classique décomposition $\mathbb{M}_d(\mathbb{R}) = S_d(\mathbb{R}) \oplus \mathfrak{A}_d(\mathbb{R})$. Le théorème du rang implique donc que $d\psi(I_d)$ est surjective. Par ailleurs, si W est un supplémentaire de $A_0^{-1} \mathfrak{A}_d(\mathbb{R})$ dans $\mathbb{M}_d(\mathbb{R})$, $d[\psi(I_d)]_W$ est inversible. Ainsi, par théorème d'inversion locale, $\psi|_W$ est un C^∞ -difféomorphisme local au voisinage de I_d . Puisque $\psi(I_d) = A_0$, on obtient l'application Φ en prenant la réciproque de $\psi|_W$ là où elle est définie. □

(v). Par « composantes de u », on entend les fonctions (u_i) telles que, si $(e_1, \dots, e_d)$ est la base canonique de $\mathbb{R}^d$, alors $u_i = e_i^* \circ u$. En d'autres termes, $u(x) = \sum u_i(x) e_i$.

Le lemme de Morse

Allons-y pour la réduction plus générale de f ! Comme f est C^3 , elle est C^2 et on peut utiliser la formule de Taylor avec reste intégral à l'ordre 2 :

$$\forall x \in \mathbb{R}^d, f(x) = f(0) + \underbrace{df(0)}_{=0} \cdot x + \int_0^1 (1-\theta) d^2 f(\theta x) \cdot (x, x) d\theta \quad (12)$$

Histoire de se raccrocher plus facilement aux formes quadratiques, passons en écriture matricielle, qui pour une fois s'avérera bien utile :

$$\forall x \in \mathbb{R}^d, f(x) = {}^t x \cdot \underbrace{\int_0^1 (1-\theta) H_f(\theta x) d\theta}_{=: Q(x)} \cdot x \quad (13)$$

Or l'application Q est à valeurs dans les matrices symétriques, car définie comme une intégrale de matrices symétriques. Comme f est de classe C^3 , elle est par ailleurs C^1 en vertu du théorème de différentiation sous l'intégrale^(vi) : si K est un compact de $\mathbb{R}^d$, l'application $(\theta, x) \mapsto (1-\theta)H_f(\theta x)$ est continûment différentiable par rapport à x , et l'application $(\theta, x) \mapsto (1-\theta)d_x[H_f(\theta x)]$ (à valeurs dans $\mathcal{L}_c(\mathbb{R}^d, S_d(\mathbb{R}))$) est dominée uniformément sur K , en norme d'opérateur, par une constante, qui est intégrable sur $[0, 1]$.

Considérons un voisinage $\tilde{V}$ de $Q(0) = \frac{1}{2}H_f(0)$ (matrice inversible car $H_f(0)$ est supposée non dégénérée) et une fonction Φ C^∞ de V dans $GL_d(\mathbb{R})$ fournie par le lemme précédent. Comme Q est continue, il existe un voisinage V de l'origine telle que si $x \in V$, alors $Q(x) \in \tilde{V}$, et on a alors :

$$\forall x \in V, Q(x) = {}^t \Phi(Q(x)) \cdot Q(0) \cdot \Phi(Q(x)) \quad (14)$$

On peut alors appliquer le théorème d'inertie de Sylvester à $H_f(0) = 2Q(0)$: il existe B une matrice inversible telle que :

$$Q(0) = {}^t \left(\frac{1}{\sqrt{2}} B \right) \cdot \begin{pmatrix} I_p & 0 \\ 0 & -I_q \end{pmatrix} \cdot \left(\frac{1}{\sqrt{2}} B \right) \quad (15)$$

Posons alors :

$$\begin{aligned} u : V &\rightarrow \mathbb{R}^d \\ x &\mapsto \frac{1}{\sqrt{2}} B \cdot \Phi(Q(x)) \cdot x \end{aligned} \quad (16)$$

u est de classe C^1 et $u(0) = 0$. On a de plus :

$$\forall x \in V, f(x) = {}^t x \cdot Q(x) \cdot x = {}^t u(x) \cdot \begin{pmatrix} I_p & 0 \\ 0 & -I_q \end{pmatrix} \cdot u(x) \quad (17)$$

Ce qui est exactement l'expression souhaitée. Il reste à montrer que u induit un C^1 difféomorphisme d'un voisinage de l'identité sur un autre voisinage de l'identité. On va encore une fois appliquer le théorème d'inversion locale^(vii) :

$$\forall h \in V, u(0+h) = \frac{1}{\sqrt{2}} B \cdot (\Phi(Q(0)) + d[\Phi \circ Q](0) \cdot h + o(\|h\|)) \cdot h \quad (18)$$

$$= u(0) + B \cdot I_d \cdot h + O(\|h\|^2) \quad (19)$$

(vi). J'ai l'impression que cette étape est souvent balayée par un rapide « f est C^3 donc il est trivial que Q est C^1 , mais la réalité me semble un peu plus subtile que ça. J'ai ajouté en annexe ledit théorème de différentiation qui n'est pas si connu que cela, histoire d'avoir bien en tête les outils du calcul différentiel qui sont invoqués.

(vii). Cette étape là me paraît également négligée dans la référence...

d'où u a pour jacobienne en 0 la matrice B , inversible par hypothèse. Le théorème d'inversion locale s'applique, ce qui achève la démonstration du lemme de Morse.

Remarque : Si on augmente la régularité de f , la régularité de Q augmente également. Ainsi, si f est supposée C^{2+n} avec $n \in \mathbb{N}^*$, la même preuve permet d'affirmer que u peut être choisi C^n . ♦

Remarque : Le lemme de Morse intervient en géométrie différentielle, où il est à la base d'une théorie, dite « Théorie de Morse », qui consiste à étudier des variétés différentielles à partir de leurs lignes de niveau. On pourra voir [IP24] ou encore le dernier chapitre de [Rou15] pour plus de détails.

On peut par ailleurs comprendre ce résultat d'un point de vue géométrique. si $f : \mathbb{R}^d \rightarrow \mathbb{R}$ satisfait les hypothèses du lemme de Morse, le changement de coordonnées u permet de réaliser un difféomorphisme local dans $\mathbb{R}^{d+1}$ envoyant le graphe de f sur une quadrique qui dépend de la signature de $d^2 f(0)$. En dimension 2 par exemple, si $d^2 f(0)$ est définie positive (de signature $(2, 0)$), alors on peut construire à partir de u un difféomorphisme qui envoie le graphe de f au voisinage de $(0, f(0))$ sur un voisinage du sommet de la parabolöide d'équation $z = x^2 + y^2$ via $(x, y, z) \mapsto (u_1(x, y), u_2(x, y), z)$. ♦

Annexe : différentiation sous l'intégrale

On a l'habitude de dériver sous l'intégrale lorsqu'on manipule des fonctions définies comme intégrale dépendant d'un paramètre réel, voire complexe. Mais ici, notre paramètre est un vecteur de $\mathbb{R}^d$, et donc ce n'est pas le même théorème qui s'applique (même si en dimension finie on peut s'y ramener)! En fait, le bon critère est le suivant :

Théorème 148 (Différentiation sous l'intégrale ([Ouv00], théorème 12.13)). Soit E et G deux e.v.n. de dimension finie et soit $(X, \mathfrak{F}, \mu)$ un espace mesuré. Soit $f : E \times X \rightarrow G$ telle que :

1. $\forall v \in E, f(v, -) \in L^1(\mu)$
2. pour presque tout $x \in X$, $f(-, x)$ est de classe C^1 .
3. pour tout compact K de E , il existe $g_K \in L^1(\mu)$ telle que pour tout $(v, x) \in K \times X$, $|||d_x f(v, x)||| \leq g_K(x)$.

Alors l'application :

$$F : v \mapsto \int_X f(v, x) d\mu(x) \quad (20)$$

est de classe C^1 sur E et on a :

$$\forall v \in E, \forall h \in E, dF(v) \cdot h = \int_X d_x f(v, x) \cdot h d\mu(x) \quad (21)$$

Démonstration. En dimension finie, on s'épargne l'abstraction du calcul différentiel en revenant aux dérivées partielles et en utilisant l'équivalence des normes. En effet, pour tout $x \in X$, $f(-, x)$ admet des dérivées partielles selon toutes ses coordonnées, et en exploitant l'équivalence de la norme d'opérateur avec d'autres normes sur $\mathcal{L}_c(E, G)$ où la dépendance aux coordonnées est plus explicite, on obtient que toutes les dérivées partielles sont dominées uniformément sur les compacts de E par des applications L^1 sur X . Le théorème de dérivation sous l'intégrale

s'applique alors et F admet toutes ses dérivées partielles, et celles-ci sont continues et s'obtiennent en dérivant directement sous l'intégrale. Ceci implique que F est C^1 et qu'on peut intervertir différentielle et intégrale. $\square$

Remarque : Le théorème reste vrai en dimension infinie. Il faut alors faire une démonstration directe, sans utiliser la dérivation sous l'intégrale, en exploitant l'inégalité des accroissements finis. Ce n'est pas spécialement plus compliqué, mais c'est beaucoup plus abstrait, alors j'ai préféré me limiter ici à la dimension finie, qui nous suffit amplement. $\blacklozenge$

3.4 Nombre de cycles d'une permutation aléatoire ★

Recasages possibles : 105, 190, 261, 264

On va dans ce développement établir certaines propriétés moyennes d'une permutation prise au hasard dans $\mathfrak{S}_n$. Outre que ce sujet semble beaucoup plaire au jury au vu du nombre de fois qu'ils l'évoquent dans le rapport, il donne lieu à de jolies démonstrations combinatoires et trouve des applications en informatique, permettant par exemple d'établir s'il est raisonnable de chercher à générer une permutation ayant certaines propriétés voulues par méthode de rejet.

On retrouvera ce développement éparpillé dans les questions d'un problème de probabilités dans [Goua]

On considère dans tout ce développement un entier non nul n et une variable aléatoire suivant une loi uniforme sur $\mathfrak{S}_n$. Posons C_n la variable aléatoire qui compte le nombre de cycles apparaissant dans sa décomposition en cycles à supports disjoints, avec la convention que tout point fixe compte pour un 1-cycle (par exemple, l'identité est considérée comme composée de n -cycles). L'objectif est d'obtenir la loi et l'espérance de C_n . Pour cela, on va calculer sa fonction génératrice en établissant une relation de récurrence double sur les probabilités atomiques.

Une formule de récurrence

Introduisons les ensembles suivants :

$$\mathfrak{S}_{n,k} := \{\sigma \in \mathfrak{S}_n \mid \sigma \text{ se décompose en } k \text{ cycles à supports disjoints}\} \quad (1)$$

Il est alors clair, par probabilités uniformes, que :

$$\mathbb{P}(C_n = k) = \frac{|\mathfrak{S}_{n,k}|}{n!} \quad (2)$$

On va établir une formule de récurrence sur n et k qui lie les cardinaux des $\mathfrak{S}_{n,k}$.

Tout d'abord, notons que la seule permutation de $\llbracket 1, n \rrbracket$ étant produit de n -cycles à supports disjoints est l'identité, tandis qu'aucune permutation n'est produit de zéro cycles. D'où :

$$\forall n \in \mathbb{N}^*, \quad |\mathfrak{S}_{n,n}| = 1 \wedge |\mathfrak{S}_{n,0}| = 0 \quad (3)$$

Notons, pour $l \in \llbracket 1, n \rrbracket$:

$$\mathfrak{S}_{n+1,k+1}(l) := \{\sigma \in \mathfrak{S}_{n+1,k+1} \mid \sigma(n+1) = l\} \quad (4)$$

On alors la partition :

$$\mathfrak{S}_{n+1,k+1} = \bigsqcup_{l=1}^n \mathfrak{S}_{n+1,k+1}(l) \quad (5)$$

On va dénombrer selon la valeur de l .

Pour $l = n+1$, on considère l'application :

$$\begin{aligned} \mathfrak{S}_{n+1,k+1}(n+1) &\rightarrow \mathfrak{S}_{n,k} \\ \sigma &\mapsto \sigma|_{\llbracket 1, n \rrbracket} \end{aligned}$$

Cette application est bien définie car $n + 1$ est un point fixe pour toute les permutations de $\mathfrak{S}_{n+1,k}(n+1)$. C'est une bijection dont on donne explicitement l'inverse :

$$\begin{aligned} \mathfrak{S}_{n,k} &\rightarrow \mathfrak{S}_{n+1,k+1}(n+1) \\ \sigma &\mapsto \left(i \mapsto \begin{cases} \sigma(k) & \text{si } k < n+1 \\ n+1 & \text{sinon} \end{cases} \right) \end{aligned}$$

Ainsi :

$$|\mathfrak{S}_{n+1,k+1}(n+1)| = |\mathfrak{S}_{n,k}| \quad (6)$$

Pour $l \leq n$, on considère la transposition $\tau = (l \ n+1)$. Alors :

$$\forall \sigma \in \mathfrak{S}_{n+1,k+1}(l), \tau\sigma \in \mathfrak{S}_{n+1,k+2}(n+1) \quad (\text{viii}) \quad (7)$$

Par ailleurs, la multiplication par τ est une bijection. Ainsi :

$$|\mathfrak{S}_{n+1,k+1}(l)| = |\mathfrak{S}_{n+1,k+2}(n+1)| = |\mathfrak{S}_{n,k+1}| \quad (8)$$

Finalement, il vient :

$$|\mathfrak{S}_{n+1,k+1}| = \left| \bigsqcup_{l=1}^{n+1} \mathfrak{S}_{n+1,k+1}(l) \right| \quad (9)$$

$$= \sum_{l=1}^{n+1} |\mathfrak{S}_{n+1,k+1}(l)| \quad (10)$$

$$= |\mathfrak{S}_{n,k}| + n|\mathfrak{S}_{n,k+1}| \quad (11)$$

Factorielle croissante

Grâce à cette relation de récurrence, on va pouvoir déterminer la fonction génératrice de C_n . Considérons le polynôme suivant :

Définition 149 (Factorielle croissante). Pour $n \in \mathbb{N}$, on appelle factorielle croissante le polynôme à coefficients entiers :

$$X^{\overline{n}} := \prod_{k=0}^{n-1} (X+k) \quad (12)$$

Notons sous forme développée :

$$X^{\overline{n}} = \sum_{k=0}^n s(n,k) X^k \quad (13)$$

On va montrer que $s(n,k) = |\mathfrak{S}_{n,k}|$ en montrant que ces deux suites doubles satisfont la même relation de récurrence.

(viii). On a bien $k+2 \leq n+1$ car l étant différent de $n+1$, $n+1$ n'est pas un point fixe de σ et donc σ est différente de l'identité. De plus, on a augmenté le nombre de cycles dans la décomposition de σ d'exactly 1 en multipliant par τ , car cycle contenant auparavant $n+1$ a été divisé d'un côté en $\{n+1\}$, de l'autre en un cycle identique dont on a retiré $n+1$.

Puisque $|s(m, m)|$ est le coefficient dominant de la factorielle croissante, qui est un polynôme unitaire, ce nombre est égal à 1. Par ailleurs, $s(m, 0)$ est le produit des racines de la factorielle croissante, parmi lesquelles on compte 0 : c'est donc 0. De plus :

$$X^{\overline{n+1}} = (X + n)X^{\overline{n}} \quad (14)$$

$$= (X + n) \sum_{k=0}^n s(n, k) X^k \quad (15)$$

$$= X^{n+1} + \sum_{k=0}^n \underbrace{(s(n, k) + ns(n, k+1))}_{=|s(k+1, n+1)|} X^k \quad (16)$$

Finalement, on trouve bien la même relation de récurrence.

On déduit donc :

Proposition 150. *La fonction génératrice de C_n est donnée par :*

$$G_{C_n}(t) = \frac{1}{n!} t^{\overline{n}} \quad (17)$$

Calcul de l'espérance

Terminons cette petite étude en calculant l'espérance de C_n . Puisqu'on connaît sa fonction génératrice, on dispose de la formule suivante :

$$\mathbb{E}[C_n] = G'_{C_n}(1) = \frac{1}{n!} (X^{\overline{n}})'(1) \quad (18)$$

Pour calculer ce terme, on va utiliser la petite astuce suivante :

$$\frac{(X^{\overline{n}})'}{X^{\overline{n}}} = \sum_{k=0}^{n-1} \frac{1}{X+k} \quad (19)$$

Comme par ailleurs $1^{\overline{n}} = n!$, on trouve :

$$\mathbb{E}[C_n] = \sum_{k=1}^n \frac{1}{k} = H_n \quad (20)$$

où H_n est le n -ième nombre harmonique. En particulier :

$$\mathbb{E}[C_n] \sim \log(n) \quad (21)$$

lorsque n tend vers l'infini. □

3.5 Sous-espaces de C^0 stables par translations ★★

Recasages possibles : 105, 190, 262, 264

On va s'intéresser ici à un problème à la croisée de l'analyse et de l'algèbre. Quand on manipule des espaces de fonctions régulières sur $\mathbb{R}$, on dispose de plusieurs opérateurs usuels : les translations et la dérivation. On va voir ici que si un espace de dimension finie est stable par les premiers, alors il est automatiquement composé de fonctions très régulières, qui sont solutions d'une équation différentielle linéaire. En particulier, il est alors stable par le second.

A ma connaissance, ce résultat n'a pas de grand intérêt en tant que tel. Il a plus une portée philosophique, car il caractérise les espaces de solutions d'équa diff linéaires à partir de la stabilité par certains opérateurs usuels. On doit j'imagine pouvoir exploiter sa contraposée pour montrer que certaines fonctions ne sont pas solutions d'équations différentielles linéaires à coefficients constants, ce qui est un problème loin d'être facile a priori.

On s'intéresse ici à $C^0(\mathbb{R})$, l'espace des fonctions continues de $\mathbb{R}$ dans $\mathbb{R}$. Pour $a \in \mathbb{R}$ et $f \in C^0(\mathbb{R})$, on note $\tau_a f : x \mapsto f(x+a)$ la translatée de f par a . On note par ailleurs $\nabla : C^1(\mathbb{R}) \rightarrow C^0(\mathbb{R})$ l'opérateur de dérivation.

Etant donné un espace vectoriel E de dimension finie, on notera E^* son dual. Si $(f_1, \dots, f_d)$ est une base de E , on note $(f_1^*, \dots, f_d^*)$ la base duale associée. Par ailleurs, si A est une partie de E^* , on notera $A^\circ \subset E^*$ son annulateur (ou orthogonal au sens de la dualité). Si B est une partie de V^* , on notera ${}^\circ B \subset E$ son annulateur. Notre objectif est de démontrer :

Théorème 151. *Soit F un sous-espace de dimension $d < +\infty$ de $C^0(\mathbb{R})$. Les assertions suivantes sont équivalentes :*

1. $\forall a \in \mathbb{R}, \tau_a(F) \subset F$
2. F est l'ensemble des solutions d'une équation différentielle linéaire homogène d'ordre d à coefficients constants. En particulier, $F \subset C^\infty(\mathbb{R})$ et $\nabla(F) \subset F$.

Un lemme algébrique

On va commencer par prouver un résultat d'ordre algébrique, qui sous sera bien utile dans la preuve du théorème :

Lemme 152. *Soit $f_1, \dots, f_n$ une famille de fonctions de $\mathbb{R}$ dans $\mathbb{R}$. On équivale :*

1. $f_1, \dots, f_n$ est une famille libre.
2. Il existe $x_1, \dots, x_n$ des réels tels que la matrice $(f_i(x_j))_{1 \leq i, j \leq n}$ est inversible.

Démonstration. Notons que le sens réciproque est relativement évident par contraposition : une relation de dépendance linéaire $\sum \lambda_i f_i = 0$ se reportant nécessairement point par point, on a pour tout $x \in \mathbb{R}$ $\sum \lambda_i f_i(x) = 0$, ce qui entraîne une relation de dépendance linéaire entre les colonnes de notre matrice, ce quelques soient les réels $x_1, \dots, x_n$. Attelons-nous donc à montrer le sens réciproque.

Supposons la famille (f_i) libre et posons $V := \text{Vect}(f_1, \dots, f_n) \subset \mathbb{R}^\mathbb{R}$. Il s'agit d'un espace de dimension finie : passons au dual et remarquons que V^* est engendré par les applications

$(ev_x)_{x \in \mathbb{R}}$, où ev_x est l'évaluation en x des fonctions de V . En effet en utilisant l'orthogonalité au sens de la dualité :

$$\text{Vect}(ev_x \mid x \in \mathbb{R}) = \sum_{x \in \mathbb{R}} \text{Vect}(ev_x) \quad (1)$$

$$= \left(\circ \left(\sum_{x \in \mathbb{R}} \text{Vect}(ev_x) \right) \right)^\circ \quad (2)$$

$$= \left(\bigcap_{x \in \mathbb{R}} \circ \text{Vect}(ev_x) \right)^\circ \quad (3)$$

$$= \left(\bigcap_{x \in \mathbb{R}} \{f \in V \mid f(x) = 0\} \right)^\circ \quad (4)$$

$$= \{0\}^\circ \quad (5)$$

$$= V \quad (6)$$

Ainsi, la famille $(ev_x)_{x \in \mathbb{R}}$ est génératrice de V^* et par théorème de la base extraite, il existe $x_1, \dots, x_n$ des réels tels que $(ev_{x_1}, \dots, ev_{x_n})$ est une base de V^* . Or, pour j fixé, en écrivant ev_{x_j} dans la base duale $(f_1, \dots, f_n)$ on a :

$$ev_{x_j} = \sum_{i=1}^n f_i(x_j) f_i^* \quad (7)$$

Ainsi, la matrice $(f_i(x_j))$ est la matrice de passage de la base $(ev_{x_1}, \dots, ev_{x_n})$ vers $(f_1^*, \dots, f_n^*)$. Elle est donc inversible. $\square$

Démonstration du théorème

Remarquons en premier lieu que F est l'ensemble des solutions d'une équation différentielle linéaire homogène à coefficients constants d'ordre d si, et seulement si, il existe P un polynôme P de degré d tel que :

$$F = \text{Ker}(P(\nabla)) \quad (8)$$

Dans ce cas, il est immédiat de constater que pour $f \in F$ et $a \in \mathbb{R}$:

$$[P(\nabla)] \cdot (\tau_a f) = \tau_a [P(\nabla) \cdot (f)] = 0 \quad (9)$$

et donc F est stable par τ_a .

Supposons pour montrer la réciproque que F est stable par les opérateurs de translation. On va, dans l'ordre, montrer que toutes les fonctions de F sont C^∞ , puis montrer que F est stable par dérivation, et enfin obtenir le théorème.

Comme F est de dimension finie, on peut en extraire $(f_1, \dots, f_d)$ une base. Soient $a \in \mathbb{R}$ et $i \in \llbracket 1, d \rrbracket$. Alors :

$$\exists (\lambda_{i,1}(a), \dots, \lambda_{i,d}(a)) \in \mathbb{R}^d : \tau_a f_i = \sum_{j=1}^d \lambda_{i,j}(a) f_j \quad (\star)$$

puisque $\tau_a f_i \in F$. Montrons que les fonctions $(\lambda_{i,j})$ sont de classe C^1 . On commence pour cela par intégrer les f_i afin de travailler avec des fonctions C^1 : posons

$$F_i : x \mapsto \int_0^x f_i(t) dt \quad (10)$$

On a :

$$\forall (a, x) \in \mathbb{R}^2, \forall i \in \llbracket 1, d \rrbracket, F_i(x+a) - F_i(a) = \int_a^{x+a} f_i(t) dt \quad (11)$$

$$= \int_0^x \tau_a f_i(t) dt \quad (12)$$

$$= \int_0^x \sum_{j=1}^d \lambda_{i,j}(a) f_j(t) dt \quad (13)$$

$$= \sum_{j=1}^d \lambda_{i,j}(a) F_j(x) \quad (14)$$

Notons par ailleurs que les (F_i) sont libres, car toute relation de dépendance linéaire entre elles se reporterait sur les (f_i) en dérivant l'équation, ce qui est exclu. Ainsi, le lemme précédent nous permet de prendre $(x_1, \dots, x_n)$ une famille de réels telle que la matrice $M := (F_i(x_j))_{1 \leq i, j \leq d}$ est inversible. Posons :

$$\Lambda_i(a) := \begin{pmatrix} \lambda_{i,1}(a) \\ \vdots \\ \lambda_{i,d}(a) \end{pmatrix} \quad B_i(a) := \begin{pmatrix} F_i(x_1+a) - F_i(x_1) \\ \vdots \\ F_i(x_d+a) - F_i(x_d) \end{pmatrix} \quad (15)$$

de sorte qu'on ait :

$$B_i(a) = M \Lambda_i(a) \quad (16)$$

et puisque M est inversible, on a :

$$\Lambda_i(a) = M^{-1} B_i(a) \quad (17)$$

Or B_i est une fonction C^1 de a . Donc Λ_i est C^1 , ce pour tout i , et donc les $(\lambda_{i,j})$ sont tous C^1 .

Or, si on évalue $(\star)$ en 0, on obtient :

$$\forall i \in \llbracket 1, d \rrbracket, \forall a \in \mathbb{R}, f_i(a) = \sum_{j=1}^d \lambda_{i,j}(a) f_j(0) \quad (18)$$

d'où les (f_i) sont tous C^1 et par linéarité, $F \subset C^1(\mathbb{R})$. Mais il se produit ici un effet « bootstrap » : les (F_i) sont alors C^2 , donc les $(\lambda_{i,j})$ aussi, et donc les (f_i) sont en fait C^2 . Etc. De cela, on déduit par récurrence immédiate que les (f_i) sont C^∞ .

Montrons maintenant que F est stable par dérivation. Pour cela, on exploite l'équation $(\star)$:

$$\forall i \in \llbracket 1, d \rrbracket, \forall (a, x) \in \mathbb{R}^2, f_i(a+x) = \sum_{j=1}^d \lambda_{i,j}(a) f_j(x) \quad (19)$$

équation qu'on peut dériver en a :

$$\forall i \in \llbracket 1, d \rrbracket, \forall (a, x) \in \mathbb{R}, \partial_a [(a, x) \mapsto f_i(a + x)](a, x) = f'_i(a + x) = \sum_{j=1}^d \lambda'_{i,j}(a) f_j(x) \quad (20)$$

Autrement dit :

$$\forall i \in \llbracket 1, d \rrbracket, f'_i = \sum_{j=1}^d \lambda_{i,j}(0) f_j \in F \quad (21)$$

Par linéarité, (f_i) engendrant F , on a $\nabla(F) \subset F$.

On peut donc considérer l'endomorphisme induit $\nabla|_F$. Puisque F est un espace de dimension finie, $\nabla|_F$ admet un polynôme minimal, noté P , tel que :

$$P(\nabla|_F) = 0 \quad (22)$$

Autrement dit :

$$F \subset \text{Ker}(P(\nabla)) \quad (23)$$

Mais les théorèmes généraux portant sur les équations différentielles linéaires indiquent que $\text{Ker}(P(\nabla))$ est un espace de dimension $\deg(P)$. Comme par ailleurs, dû au théorème de Cayley-Hamilton, $\deg(P) \leq d$, on obtient :

$$F = \text{Ker}(P(D)) \wedge \deg(P) = d \quad (24)$$

et donc F est l'ensemble des solutions de l'équation différentielle linéaire homogène à coefficients constants $P(\nabla) = 0$. $\square$

3.6 Théorème de Kakutani et sous-groupes compacts de $GL(E)$ ★★☆☆

Recasages possibles : 103, 106, 161, 181, 203

Dans ce développement, on se propose de montrer que tout sous-groupe compact du groupe des automorphismes linéaires d'un espace euclidien peut se traduire comme un certain groupe d'isométries en prouvant l'existence d'un produit scalaire laissé invariant par notre sous-groupe. Pour cela, on va exploiter un argument de point fixe, en démontrant un théorème de point fixe dans la première partie.

L'intégralité du développement est excellemment rédigé dans [IP24]. On peut aussi retrouver une preuve dans [Ale99] quoi que celle-ci exploite des arguments plus compliqués de compacité ; elle est un peu plus général mais le développement est assez long comme ça sans avoir besoin de se rajouter des difficultés.

Etant donné $(E, \langle \cdot, \cdot \rangle)$ un espace euclidien, on note $S(E)$ l'ensemble de ses endomorphismes auto-adjoints et $O(E)$ le groupe de ses isométries. Pour F un second espace euclidien et $u \in \mathcal{L}(E, F)$, on note u^* l'adjoint de u .

Théorème du point fixe de Kakutani

Théorème 153 (Kakutani). Soient $(E, \langle \cdot, \cdot \rangle)$ un espace euclidien, G un sous-groupe compact de $GL(E)$ et K une partie compacte et convexe de E . On suppose :

$$\forall g \in G, g(K) \subset K \quad (1)$$

Alors, K contient un point fixe sous l'action de G .

Pour démontrer ce théorème, on va commencer par définir une norme G -invariante sur E . Le point fixe sera alors réalisé comme l'élément de K qui en atteint le minimum !

Lemme 154.

$$N : E \rightarrow \mathbb{R}_+ \\ x \mapsto \max_{g \in G} \|g(x)\|$$

Alors N est une norme sur E . De plus :

1. $\forall g \in G, \forall x \in E, N(g(x)) = N(x)$
2. $\forall (x, y) \in E^2, [N(x + y) = N(x) + N(y)] \iff [\exists \lambda \in \mathbb{R}_+ : x = \lambda y \vee y = \lambda x]$
3. N admet un unique minimum sur K

Démonstration. Tout d'abord, notons que N est bien définie par compacité de G . En effet, l'application $g \mapsto \|g(x)\|$ est continue à x fixé, car on a :

$$\|g(x) - h(x)\| \leq \|g - h\| \cdot \|x\| \rightarrow [h \rightarrow g]0 \quad (2)$$

Il suit que cette application est bien maximisée sur G . N définit par ailleurs une norme. Vérifions les axiomes un à un :

Séparation : pour $x \in E$ tel que $N(x) = 0$, on a que pour $g \in G$ quelconque, $\|g(x)\| = 0$.

Comme g est inversible, on a nécessairement $x = 0$.

Homogénéité : soit $\lambda \in \mathbb{R}$. Alors :

$$N(\lambda x) = \max_{g \in G} |\lambda| \|g(x)\| \quad (3)$$

$$= |\lambda| \max_{g \in G} \|N(x)\| \quad \text{car } |\lambda| \geq 0 \quad (4)$$

$$= |\lambda| N(x) \quad (5)$$

Inégalité triangulaire : soit $(x, y) \in E^2$ et soit $g \in G$. Alors :

$$\|g(x + y)\| \leq \|g(x)\| + \|g(y)\| \leq N(x) + N(y) \quad (6)$$

Comme c'est vrai pour tout g , $N(x + y) \leq N(x) + N(y)$.

Montrons maintenant les propriétés de N .

1. Soit $g \in G$, soit $x \in E$. On a :

$$N(g(x)) = \max_{h \in G} \|(h \circ g)(x)\| = \max_{h \in Gg=G} \|h(x)\| = N(x) \quad (7)$$

Donc N est bien G -invariante.

2. Soit x et y des vecteurs réalisant le cas d'égalité dans l'inégalité triangulaire de N . Soit $g \in G$ tel que $N(x + y) = \|(g(x + y))\|$. On a alors :

$$N(x + y) = \|g(x + y)\| \leq \|g(x)\| + \|g(y)\| \leq N(x) + N(y) \quad (8)$$

Puisqu'on est dans le cas d'égalité, on a nécessairement $\|g(x) + g(y)\| = \|g(x)\| + \|g(y)\|$. Comme $\|\cdot\|$ est une norme euclidienne, $g(x)$ et $g(y)$ sont positivement liés, et par inversibilité de g , x et y sont positivement liés. Le sens réciproque est immédiat.

3. Notons que N est continue pour la topologie qu'elle engendre sur E , et que par équivalence des normes en dimension finie, elle est donc encore continue pour la topologie euclidienne de E . Ainsi, comme K est compact, N admet bien un minimum.

Soient y et z deux minimums de N sur K . Alors, par convexité de K , $x/2 + y/2$ est encore un élément de K et on a :

$$N\left(\frac{y + z}{2}\right) \leq \frac{N(y) + N(z)}{2} = N(z) \quad (9)$$

Par minimalité de $N(z)$, l'égalité est nécessairement atteinte, donc y et z sont positivement liés d'après le point précédent. Supposons $y = \lambda z$, où λ est un réel positif. Alors :

$$N(y) = \lambda N(z) = N(z) \quad (10)$$

d'où $\lambda = 1$ et $y = z$. Le minimum est bien unique!

□

De cette construction, on déduit immédiatement le théorème de Kakutani : soit z l'unique minimum de N sur K . Alors, comme N est G -invariante :

$$\forall g \in G, N(g(z)) = N(z) \quad (11)$$

Donc $g(z)$ réalise le minimum de N pour tout G . Puisque les éléments de G envoient K sur lui-même, $g(z)$ est encore un élément de K , et par unicité du minimum, $g(z) = z$, ce pour tout $g \in G$.

Sous-groupes compacts de $GL(E)$

Dans cette partie, on identifie E avec $\mathbb{R}^d$ via le choix d'une base orthonormée, permettant de travailler plutôt avec des matrices.

Théorème 155. *Tout sous groupe compact H de $GL_d(\mathbb{R})$ est conjugué à un sous groupe de $O_d(\mathbb{R})$, c'est-à-dire qu'il existe $R \in GL_d(\mathbb{R})$ tel que :*

$$R^{-1} \cdot H \cdot R \subset O_d(\mathbb{R}) \quad (12)$$

Pour démontrer ce théorème, on va utiliser le théorème de Kakutani pour exhiber un produit scalaire H -invariant sur $\mathbb{R}^d$ (ix). Cela revient à trouver une matrice symétrique définie positive B qui soit point-fixe sous l'action de H par congruence. Si une telle B existe, on peut en effet poser :

$$\forall (x, y) \in (\mathbb{R}^d)^2, \langle x | y \rangle_H := \langle B \cdot x, y \rangle \quad (13)$$

qui définit bien un produit scalaire sur $\mathbb{R}^d$ tel que pour tout $M \in H$, $\langle M \cdot x | M \cdot y \rangle_H = \langle x | y \rangle_H$.

Remarque : Une difficulté peut-être de ce développement est qu'on a ici un conflit sur la nature des objets qu'on manipule qui peut rendre la démarche un peu confusante. Ce conflit se voit un peu mieux en utilisant des matrices une matrice symétrique B peut-être vue sous deux angles très différents. Ce peut être la matrice d'une application linéaire, auquel cas les groupes d'isomorphismes agissent sur de telles matrice par conjugaison (changement de base pour un endomorphisme). Mais ce peut également être la matrice d'une forme quadratique, auquel cas l'action naturelle des groupes d'isomorphismes sur ces matrices est plutôt la congruence (changement de base pour les formes quadratiques). Ainsi, va plutôt utiliser ce point de vue « formes quadratiques » : chercher une matrice symétrique définie positive, c'est chercher la matrice d'un produit scalaire, et la condition de point fixe sous l'action de congruence de H revient exactement à dire que les éléments de H se comportent comme des isométries pour ce produit scalaire. ♦

Considérons l'action à droite par congruence de H sur l'espace des matrices symétriques :

$$\begin{aligned} \rho : H^{op} \text{ (x)} &\rightarrow GL(S_d(\mathbb{R})) \\ M &\mapsto (S \mapsto {}^t M \cdot S \cdot M) \end{aligned}$$

C'est un morphisme de groupes : étant donné M et N dans H , on a :

$$\forall S \in S_d(\mathbb{R}), \rho(M \cdot_{op} N)(S) = {}^t(NM)S(MN) = (\rho(N) \circ \rho(M))(S) \quad (15)$$

C'est de plus une application continue. Pour plus de clarté, on notera simplement $\|\cdot\|$ la norme d'opérateur sur $M_d(\mathbb{R})$ subordonnée à la norme euclidienne et $|||\cdot|||$ la norme sur $\mathcal{L}(S_d(\mathbb{R}))$

(ix). L'existence d'une tel produit scalaire est équivalente au théorème. Si cela ne vous paraît pas clair, rendez-vous en annexe !

(x). On désigne par H^{op} le groupe opposé de H , c'est-à-dire le groupe dont les éléments sont ceux de H mais où la loi de composition est renversée :

$$\forall (M, N) \in H^{op2}, M \cdot_{op} N := N \cdot M \quad (14)$$

On le verra, cela permettra à l'application ρ d'être un morphisme de groupes au sens propre du terme lorsqu'elle aurait été contrinvariante vis-à-vis de la loi de groupe sur H .

subordonnée à cette norme d'opérateur.

$$\forall S \in S(E), \|\rho(M)(S) - \rho(N)(S)\| = \|{}^tMSM - {}^tNSN\| \quad (16)$$

$$= \|{}^t(M - N) \cdot S \cdot (M - N) + {}^tMSN + {}^tNSM - 2{}^tNSN\| \quad (17)$$

$$\leq \|{}^t(M - N)S(M - N)\| + \|{}^t(M - N)SN\| + \|{}^tNS(M - N)\| \quad (18)$$

$$\leq (\|M - N\|^2 + 2\|M - N\| \|N\|) \|S\| \quad (19)$$

Donc en passant à la norme d'opérateur :

$$\|\rho(M) - \rho(N)\| \leq (\|M - N\|^2 + 2\|M - N\| \|N\|) \xrightarrow{N \rightarrow M} 0 \quad (20)$$

Donc ρ est bien continu. Ainsi, en posant $G := \text{Im}(\rho)$, on définit un sous-groupe compact de $GL(S_d(\mathbb{R}))$. S'il existe une partie convexe compacte de $S_d^{++}(\mathbb{R})$ stable par G , par le théorème de Kakutani, celle-ci contiendra une matrice symétrique définie positive B stable par G , c'est-à-dire stable sous l'action de congruence de H . On pose :

$$K := \text{Conv}(\{{}^tM \cdot M, M \in H\}) \quad (21)$$

où Conv désigne l'enveloppe convexe. On a trois choses à démontrer :

Lemme 156.

1. K est compact.
2. $K \subset S_d^{++}(\mathbb{R})$
3. K est stabilisé par l'action de G (i.e. pour tout $g \in G$, $g(K) \subset K$)

Démonstration.

1. Il s'agit d'un cas particulier d'un fait plus général : dans un espace vectoriel de dimension d finie, l'enveloppe convexe d'un compact est compact. C'est un corollaire du lemme de Caratheodory :

Théorème 157 (Caratheodory^(xi)). Soit E un espace vectoriel de dimension d . Soit $A \subset E$. Alors :

$$\text{Conv}(A) = \left\{ \sum_{i=1}^{d+1} \lambda_i x_i, (x_i) \subset A, (\lambda_i) \subset [0, 1] \mid \sum_{i=1}^{d+1} \lambda_i = 1 \right\} \quad (22)$$

(xi). Je pense qu'on ne peut pas espérer démontrer ce théorème en 15 minutes, le développement étant déjà assez long, mais il faut en connaître la preuve ! Puisqu'on l'utilise sans démonstration, je pense qu'il faut que le théorème soit présent dans le plan, clairement indiqué comme ne faisant pas partie du développement, afin de ne pas surprendre le jury.

on peut donc exprimer K comme l'image de l'application continue :

$$\{ {}^t M^* M, M \in H \}^{d+1} \times \left\{ (\lambda_i) \in [0, 1]^{d+1} \mid \sum_{i=1}^{d+1} \lambda_i = 1 \right\} \rightarrow K$$

$$(x_i), (\lambda_i) \mapsto \sum_{i=1}^{d+1} \lambda_i x_i$$

Comme l'ensemble de départ est compact, K est compact.

2. Cela vient du fait que $S_d^{++}(\mathbb{R})$ est convexe. En effet, étant donné $(S_1, S_2) \in S_d^{++}(\mathbb{R})^2$ et $t \in [0, 1]$, on a :

$$\forall x \in \mathbb{R}^d, \langle (tS_1 + (1-t)S_2) \cdot (x), x \rangle = t \langle S_1 \cdot x, x \rangle + (1-t) \langle S_2 \cdot x, x \rangle > 0 \quad (23)$$

donc $tu + (1-t)v \in S^{++}(E)$.

3. Soit $g = \rho(M) \in G$, et soit $N \in H$. Alors :

$$g({}^t N N) = {}^t M^t N N M = \underbrace{{}^t (M N)}_{\in H} \underbrace{M N}_{\in H} \quad (24)$$

Donc $\{ {}^t M M, M \in H \}$ est stable sous l'action de G et par linéarité, K est stable sous l'action de G .

□

On peut donc appliquer le théorème de Kakutani à K : il existe $B \in K \subset S_d^{++}(\mathbb{R})$ tel que B soit stable par congruence par des éléments de H , c'est-à-dire qu'il existe un produit scalaire H -invariant sur E . □

Annexe : produits scalaires invariants et conjugaison au groupe orthogonal

Dans ce développement, on a utilisé le résultat suivant, qu'il vaut mieux je pense ne pas démontrer, à moins que vous ne soyez très rapide à l'oral et très sûr.e de vous :

Proposition 158. *Soit $(E, \langle \cdot, \cdot \rangle)$ un espace euclidien et soit G un sous-groupe de $GL(E)$. On suppose qu'il existe $\langle \cdot | \cdot \rangle_G$ un produit scalaire sur E qui soit G -invariant. Alors G est conjugué à un sous-groupe de $O(E)$.*

Démonstration. La démonstration est plutôt simple à condition d'avoir bien conscience des objets qu'on manipule. Pour cela, il faut absolument considérer un espace euclidien comme **un couple** composé d'un espace vectoriel **et** d'un produit scalaire ! Pour aller dans ce sens, on notera $O(\langle \cdot, \cdot \rangle)$ plutôt que $O(E)$ le groupe orthogonal pour bien signifier la dépendance au choix du produit scalaire sur E .

La G -invariance de $\langle \cdot | \cdot \rangle_G$ revient exactement à dire que $G \subset O(\langle \cdot | \cdot \rangle_G)$. Adoptons le point de vue des formes quadratiques : nos deux produits scalaires sont les formes polaires de formes quadratiques définies positives, qui sont toutes les deux représentées par la matrice identité dans des bases orthonormées respectives. Il suit :

$$(E, \langle \cdot, \cdot \rangle) \cong (E, \langle \cdot | \cdot \rangle_G) \quad (25)$$

où l'isomorphisme est au sens des espaces quadratiques. En d'autres termes, on dispose d'une isométrie $r : (E, \langle \cdot, \cdot \rangle) \rightarrow (E, \langle \cdot | \cdot \rangle_G)$ ^(xii). Prenons alors $g \in G$. On a :

$$(r^{-1}gr)(rgr^{-1})^* = r^{-1}grr^*g^*(r^{-1})^* = Id_E \quad (26)$$

et donc $r^{-1}gr \in O(\langle \cdot, \cdot \rangle)$, ce qui prouve que G est conjugué à un sous-groupe de $O(\langle \cdot, \cdot \rangle)$. $\square$

Remarque : Pour celles et ceux qui aiment les diagrammes commutatifs, c'est cadeau :

$$\begin{array}{ccc} (E, \langle \cdot, \cdot \rangle) & \xrightarrow{r \circ g \circ r^{-1}} & (E, \langle \cdot, \cdot \rangle) \\ \downarrow r & & \downarrow r \\ (E, \langle \cdot | \cdot \rangle_G) & \xrightarrow{g} & (E, \langle \cdot | \cdot \rangle_G) \end{array}$$

Les deux flèches r et g sont des isomorphismes d'espaces quadratiques, donc nécessairement la dernière flèche l'est aussi, ce qui est l'exacte traduction du fait que c'est un élément du groupe orthogonal :) $\blacklozenge$

(xii). Pour éviter les formes quadratiques, on peut construire explicitement r en envoyant une base orthonormée sur une base orthogonale.

3.7 Théorème de min-max de Courant-Fischer et continuité des valeurs propres hermitiennes ★

Recasages possibles : 153, 157, 158, 219

On va dans ce développement démontrer le théorème du min-max de Courant-Fischer, qui donne une expression des valeurs propres d'un endomorphisme hermitien en fonction d'extrema du quotient de Rayleigh. Cette formule, qui pourrait paraître anecdotique à première vue, a en réalité des conséquences puissantes. On va montrer ici que les fonctions valeur propre sont continues sur l'espace des matrices hermitiennes. Si ce résultat est vrai de façon plus générale, la preuve en sera dans ce cas précis élémentaire, et nous épargnera au passage d'avoir à définir proprement une notion de continuité pour l'application valeurs propres, ce qui peut s'avérer quelque peu délicat.

Soit $(E, \langle \cdot, \cdot \rangle)$ un espace hermitien ou euclidien de dimension n . Notons $\mathcal{F}_k$ l'ensemble des sous-espaces vectoriels de dimension k de E .

Soit u un endomorphisme auto-adjoint de E . D'après le théorème spectral, u admet n valeurs propres réelles, qu'on peut ordonner de la façon suivante :

$$\lambda_1(u) \leq \dots \leq \lambda_n(u) \quad (1)$$

On définit le quotient de Rayleigh de u de la façon suivante :

$$\begin{aligned} R_u : E \setminus \{0\} &\rightarrow \mathbb{R} \\ x &\mapsto \frac{\langle u(x), x \rangle}{\|x\|^2} \end{aligned} \quad (2)$$

Nous allons montrer le théorème suivant :

Théorème 159 (Min-max de Courant-Fischer ([FGN08], exo 2.37)). *Les valeurs propres de u peuvent se calculer de la façon suivante :*

$$\forall k \in \llbracket 1, n \rrbracket, \lambda_k(u) = \min_{W \in \mathcal{F}_k} \max_{x \in W \setminus \{0\}} R_u(x) = \max_{W \in \mathcal{F}_{n-k+1}} \min_{x \in W \setminus \{0\}} R_u(x) \quad (3)$$

Démonstration du théorème

Commençons par remarquer que si W est un sous-espace vectoriel de E , alors R_u atteint un maximum et un minimum sur $W \setminus \{0\}$. En effet, pour tout $x \in W$ non nul, on a $R_u(x) = R_u(x/||x||)$. Comme la sphère unité est compacte, et que R_u est continu, il atteint un maximum et un minimum sur la sphère, qui sont des extrema globaux.

Soit $(e_1, \dots, e_n)$ une base orthonormale de E telle que e_i soit vecteur propre de u pour la valeur propre $\lambda_i(u)$. On considère les drapeaux complets (F_k) définis par :

$$F_k := \text{Vect}(e_i \mid i \in \llbracket 1, k \rrbracket), \quad G_k := \text{Vect}(e_i \mid i \in \llbracket k, n \rrbracket) \quad (4)$$

Soit $k \in \llbracket 1, n \rrbracket$, et soit $x \in F_k \setminus \{0\}$. Commençons par un petit calcul :

$$R_u(x) = \frac{\langle u \left(\sum_{i=1}^k e_i^*(x) e_i \right), \sum_{k=1}^n e_i^*(x) e_i \rangle}{\left\| \sum_{k=1}^n e_i^*(x) e_i \right\|^2} \quad (5)$$

$$= \frac{\sum_{i=1}^k \lambda_i(u) e_i^*(x)^2}{\sum_{i=1}^k e_i^*(x)^2} \quad (6)$$

$$\leq \lambda_k(u) \quad (7)$$

et donc $\max_{x \in F_k \setminus \{0\}} R_u(x) \leq \lambda_k(u)$. D'où il vient immédiatement : En travaillant de même sur G_k , on obtient :

$$\forall x \in G_k \setminus \{0\}, \quad R_u(x) \geq \lambda_k(u) \quad (8)$$

et donc $\min_{x \in G_k \setminus \{0\}} R_u(x) \geq \lambda_k(u)$. Ainsi, on a immédiatement les deux inégalités :

$$\inf_{W \in \mathcal{F}_k} \max_{x \in W \setminus \{0\}} R_u(x) \leq \lambda_k(u) \quad (9)$$

$$\sup_{W \in \mathcal{F}_{n-k+1}} \min_{x \in W \setminus \{0\}} R_u(x) \geq \lambda_k(u) \quad (10)$$

Pour montrer les inégalités inverses, considérons $W \in \mathcal{F}_k$. Alors, pour des raisons de dimension, $W \cap G_k$ n'est pas réduit à 0. Soit x est un vecteur non nul de l'intersection. D'après les calculs précédents :

$$R_u(x) \geq \lambda_k(u) \quad (11)$$

puisque x est dans G_k . Donc :

$$\max_{x \in W \setminus \{0\}} R_u(x) \geq \lambda_k(u) \quad (12)$$

Comme c'est vrai pour tout tel espace W , on peut passer à l'infimum^(xiii) :

$$\inf_{W \in \mathcal{F}_{n-k+1}} \max_{x \in W \setminus \{0\}} R_u(x) \geq \lambda_k(u) \quad (13)$$

Donc l'égalité est atteinte. On raisonne de même pour l'autre égalité, en intersectant un espace de $W \in \mathcal{F}_{n-k+1}$ et en l'intersectant avec F_k . Reste à montrer que les extrema sont atteints. Pour ça, il suffit d'observer que $e_k \in F_k \in \mathcal{F}_k$ et que $R_u(e_k) = \lambda_k(u)$. De même, $e_k \in G_k \in \mathcal{F}_{n-k+1}$. $\square$

(xiii). Attention, on n'a pas montré à ce stade que l'inf sur $\mathcal{F}_k$ est atteint.

Continuité des valeurs propres

Le théorème du min-max permet d'obtenir la continuité des valeurs propres dans le cas auto-adjoint. On commence par montrer :

Proposition 160 (Inégalité de perturbation de Weyl ([FGN08], exo 2.40)). Soient u et v des endomorphismes auto-adjoints de E . On a :

$$\forall k \in \llbracket 1, n \rrbracket, \lambda_k(u) + \lambda_1(v) \leq \lambda_k(u+v) \leq \lambda_k(u) + \lambda_n(v) \quad (14)$$

Démonstration. Soit $W \in \mathcal{F}_k$ et soit $x \in W \setminus \{0\}$ tel que $R_u(x) = \lambda_k(u)$. Remarquons que $R_{u+v} = R_u + R_v$, ce qui fournit :

$$R_u(x) + R_v(x) \leq R_u(x) + \lambda_n(v) = \lambda_k(u) + \lambda_n(v) \quad (15)$$

Si maintenant V est un sous-espace de E de dimension $n - k + 1$ quelconque, comme pour tout $y \in W \setminus \{0\}$, $R_u(y) \geq R_u(y)$ et que $V \cap W$ est non réduit à 0 :

$$\min_{y \in V} R_{u+v}(y) \leq \lambda_k(u) + \lambda_n(v) \quad (16)$$

et comme cette majoration vaut pour tout $V \in \mathcal{F}_{n-k+1}$, on peut passer au max :

$$\lambda_k(u+v) = \max_{V \in \mathcal{F}_{n-k+1}} \min_{y \in V} R_{u+v}(y) \leq \lambda_k(u) + \lambda_n(v) \quad (17)$$

L'autre inégalité se montre de manière similaire. $\square$

De ce résultat, on déduit la continuité des valeurs propres :

Théorème 161 (Continuité des valeurs propres, cas auto-adjoint). Pour tout $k \in \llbracket 1, n \rrbracket$, la fonction $\lambda_k : H(E) \rightarrow \mathbb{R}$ est 1-lipschitzienne (au sens de la norme d'opérateur). En particulier, elle est continue.

Démonstration. Soient u et v deux endomorphismes auto-adjoints de E . L'inégalité de perturbation de Weyl indique que :

$$\lambda_1(v) \leq \lambda_k(u+v) - \lambda_k(u) \leq \lambda_n(v) \quad (18)$$

Aussi, en passant à la valeur absolue :

$$|\lambda_k(u+v) - \lambda_k(u)| \leq \max\{|\lambda_1(v)|, |\lambda_n(v)|\} \leq \rho(v) \quad (19)$$

où ρ est le rayon spectral. Il suffit maintenant de remarquer que si e est un vecteur propre de norme 1 associé à la valeur propre de v de plus grand module, alors :

$$\rho(v) = \|v(e)\| \geq \|v\| \quad (20)$$

D'où finalement :

$$|\lambda_k(u+v) - \lambda_k(u)| \leq \|v\| \quad (21)$$

ce qui prouve bien le caractère 1-lipschitzien. $\square$

Remarque :

1. Il est bon de savoir qu'en général, on a un résultat de continuité pour les valeurs propres, même hors du cas auto-adjoint. Toutefois, celui-ci est à la fois dur à démontrer et dur à énoncer, puisqu'il n'y a pas d'ordre canonique pour le d -uplet de valeurs propres. En fait, il faut voir les valeurs propres de u comme un élément de $\mathbb{C}^d / \mathfrak{S}_d$ (les d -uplets complexes modulo l'action du groupe symétrique, c'est-à-dire pour lesquels l'ordre n'est pas important (c'est un peu comme un ensemble mais dans lequel on tolère plusieurs fois le même élément)). Il faut alors munir cet ensemble de sa topologie quotient, et donc rentrer dans des considérations qu'il vaut mieux éviter à l'agreg.
2. Une question qui a été posée à des amis qui ont présenté ce développement en oral blanc : est-ce que les fonctions λ_k sont dérivables ? La réponse est non, ce qu'on peut voir en considérant par exemple la famille de matrices :

$$t \mapsto \begin{pmatrix} 0 & t \\ t & 0 \end{pmatrix} \quad (22)$$

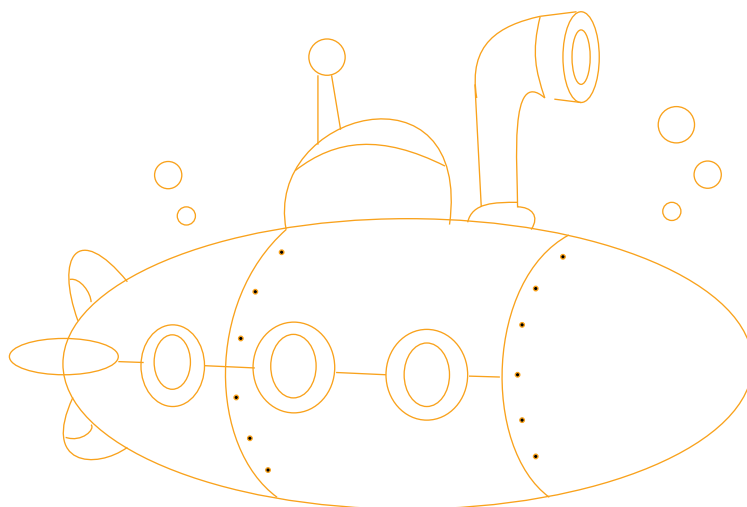
On a alors $\lambda_1(t) = -|t|$ et $\lambda_2(t) = |t|$, qui ne sont pas dérivables en 0...

◆

Annexe A

Couplage effectif

Voici enfin mon couplage sans retourche avec lequel je me suis rendu aux oraux. Peut-être vous inspirera-t-il et vous aidera-t-il à composer le vôtre.



* « *Tou.te.s dans le même bateau* »

Leçons d'algèbre

101	Probas sur $GL_2(\mathbb{F}_q)$ / Automorphismes de $\mathfrak{S}_n$
102	Cyclotomie (sur $\mathbb{Q}$) / Théorème de Kronecker
103	Probas sur $GL_2(\mathbb{F}_q)$ / Automorphismes de $\mathfrak{S}_n$
104	Probas sur $GL_2(\mathbb{F}_q)$ / Automorphismes de $\mathfrak{S}_n$
105	Automorphismes de $\mathfrak{S}_n$ / Permutations aléatoires
106	Probas sur $GL_2(\mathbb{F}_q)$ / Théorème de Kakutani
108	Automorphismes de $\mathfrak{S}_n$ / Générateurs de $O(q)$
120	Critère de Solovay-Strassen / Cyclotomie (sur $\mathbb{F}_q$)
121	Critère de Solovay-Strassen / Cyclotomie (Dirichlet faible)
122	Algorithme de Berlekamp / Forme normale de Smith
123	Cyclotomie (sur $\mathbb{F}_q$) / Norme d'un élément algébrique (carrés de $\mathbb{F}_q$)
125	Cyclotomie (sur $\mathbb{F}_q$) / Norme d'un élément algébrique (entiers algébriques)
127	Critère de Solovay-Strassen / Norme d'un élément algébrique (entiers algébriques)
141	Algorithme de Berlekamp / Cyclotomie (dans $\mathbb{F}_q$)
142	Algorithme de Berlekamp / Forme normale de Smith
144	Théorème de Kronecker / Norme d'un élément algébrique (entiers algébriques)
148	Sous-espaces stables par translations / Réduction de Frobenius
149	Par cinq points passe une conique / Norme d'un élément algébrique (entiers algébriques)
150	Décomposition de Dunford algorithmique / Réduction de Frobenius
151	Réduction de Frobenius / Théorème de Lie-Kolchin
152	Décomposition de Dunford algorithmique / Matrice circulante et suite de polygones
153	Matrice circulante et suite de polygones / Théorème de Courant-Fischer
155	IMPASSE
156	Décomposition de Dunford algorithmique / Théorème de Lie-Kolchin
157	Lemme de Morse / Théorème de Courant Fischer
158	Théorème de Kakutani / Théorème de Courant-Fischer
159	Réduction de Frobenius / Sous-espaces engendrés par les translatés
161	Théorème de Kakutani / Générateurs de $O(q)$
162	Par cinq points passe une conique / Forme normale de Smith
170	Formes quadratiques sur $\mathbb{F}_q$ / Générateurs de $O(q)$
171	Par cinq points passe une conique / Lemme de Morse
181	Par cinq points passe une conique / Théorème de Kakutani
190	Probas sur $GL_2(\mathbb{F}_q)$ / Permutations aléatoires
191	Matrice circulante et suite de polygones / Par cinq points passe une conique

Leçons d'analyse

201	Echantillonnage de Schannon / Riesz-Fréchet-Kolmogorov
203	Théorème de Kakutani / Riesz-Fréchet-Kolmogorov
204	Lemme de la grenouille / Théorème de Liapounov
205	Banach-Steinhaus et séries de Fourier / Optimisation dans un Hilbert
206	Lemme de Morse / Sous-espaces stables par translations
209	Théorème de Lévy et TCL / Riesz-Fréchet-Kolmogorov
213	Echantillonnage de Schannon / Optimisation dans un Hilbert
214	Lemme de Morse / Différentielle du flot
215	Sous-groupes petits de $GL_n(\mathbb{R})$ (différentielle) / Différentielle du flot
218	Lemme de Morse / Quadrature de Gauss-Legendre
219	Optimisation dans un Hilbert / Théorème de Courant-Fischer
220	Différentielle du flot / Théorème de Liapounov
221	Différentielle du flot (résolvante) / Théorème de Liapounov
223	Lemme de la grenouille / Théorème taubérien fort
224	IMPASSE
226	Matrice circulante et suite de polygones / Chaînes de Markov
228	Formule sommatoire de Poisson / Sous-espaces stables par translations
229	Hahn-Banach géométrique / Optimisation dans un Hilbert
230	Banach-Steinhaus et séries de Fourier / Théorème taubérien fort
234	Echantillonnage de Schannon / Riesz-Fréchet-Kolmogorov
235	Holomorphie sous l'intégrale / Banach-Steinhaus et séries de Fourier
236	Quadrature de Gauss-Legendre / Transformée de Fourier par les résidus
239	Holomorphie sous l'intégrale / Riesz-Fréchet-Kolmogorov
241	Sous-groupes petits de $GL_n(\mathbb{R})$ (différentielle) / Théorème taubérien fort
243	Sous-groupes petits de $GL_n(\mathbb{R})$ (séries matricielles) / Théorème taubérien fort
245	Holomorphie sous l'intégrale / Transformée de Fourier par les résidus
246	Formule sommatoire de Poisson / Banach-Steinhaus et séries de Fourier
250	Echantillonnage de Schannon / Formule sommatoire de Poisson
253	Hahn-Banach géométrique / optimisation dans Hilbert
261	Holomorphie sous l'intégrale / Théorème de Lévy et TCL / Chaînes de Markov
262	Théorème de Lévy et TCL / Théorème de Furstenberg-Kesten
264	Chaînes de Markov / Permutations aléatoires
266	Théorème de Lévy et TCL / Théorème de Furstenberg-Kesten

Bibliographie

- [21] *Mathématiques D (concours de l'ENS Paris)*. 2021. URL : https://www.ens-psl.eu/sites/default/files/21_mp_sujet_mathd.pdf.
- [Ale99] Michel ALESSANDRI. *Thèmes de géométrie*. Dunod, 1999.
- [AM04] Eric AMAR et Etienne MATHERON. *Analyse complexe*. Cassini, 2004. ISBN : 2-84225-052-4.
- [Amr] Mohammed El AMRANI. *Analyse de Fourier dans les espaces fonctionnels*. Ellipses.
- [BE07] Michel BENAÏM et Nicole EL KAROUI. *Promenade aléatoire*. Les éditions de l'Ecole Polytechnique, 2007. ISBN : 978-2-7302-1168-0.
- [Ben14] Sylvie BENZONI-GAVAGE. *Calcul différentiel et équations différentielles*. Dunod, 2014. ISBN : 978-2-10-070611-2.
- [BMP05] Vincent BECK, Jérôme MALICK et Gabriel PEYRÉ. *Objectif agrégation*. H&K, 2005. ISBN : 2914010923.
- [BP] Marc BRIANE et Gilles PAGÈS. *Analyse, Théorie de l'intégration*. Deboeck.
- [Bre87] Haïm BREZIS. *Analyse fonctionnelle, théorie et applications*. Masson, 1987. ISBN : 2-225-77198-7.
- [Cal23] Philippe CALDERO. *Carnet de voyage en Analystan*. Calvage & Mounet, 2023. ISBN : 978-2-49-323011-9.
- [Cha05] Antoine CHAMBERT-LOIR. *Algèbre corporelle*. Les éditions de l'Ecole Polytechnique, 2005. ISBN : 978-2-7302-1217-5.
- [CP22] Philippe CALDERO et Marie PERONNIER. *Carnet de voyage en Algèbre*. Calvage et Mounet, 2022. ISBN : 9782493230034.
- [Dem08] Michel DEMAZURE. *Cours d'algèbre*. Cassini, 2008. ISBN : 978-2-84225-127-7.
- [DLQ14] Gema-Maria DIAZ-TOCA, Henri LOMBARDI et Claude QUITTÉ. *Modules sur les anneaux commutatifs*. Calvage et Mounet, 2014.
- [Eid09] Jean-Denis EIDEN. *Géométrie analytique classique*. Calvage et Mounet, 2009. ISBN : 978-2-916352-08-4.
- [Esc97] Jean-Pierre ESCOFIER. *Théorie de Galois*. Dunod, 1997.
- [FGN08] Serge FRANCINO, Hervé GIANELLA et Serge NICOLAS. *Oraux X/ENS - Algèbre 3*. Cassini, 2008. ISBN : 978-2-84225-092-8.
- [FGN22] Serge FRANCINO, Hervé GIANELLA et Serge NICOLAS. *Oraux X/ENS - Tome 6*. Cassini, 2022. ISBN : 2842252462.
- [Goua] Xavier GOURDON. *Algèbre et probabilités*. Ellipses.

- [Goub] Xavier GOURDON. *Analyse*. Ellipses.
- [Goz97] Ivan GOZARD. *Théorie de Galois*. Ellipses, 1997.
- [GT98] Stéphane GONNORD et Nicolas TOSEL. *Thèmes d'Analyse pour l'Agrégation - Calcul différentiel*. Ellipses, 1998.
- [HL09] Francis HIRSCH et Gilles LACOMBE. *Elements d'analyse fonctionnelle*. Dunod, 2009. ISBN : 978-2-10-053594-1.
- [Hos24] Matthias HOSTEIN. *La réduction par la théorie des modules*. 2024. URL : https://perso.eleves.ens-rennes.fr/people/matthias.hostein/Fichiers_site/La_reduction_par_la_theorie_des_modules.pdf.
- [IP24] Lucas ISENMANN et Timothée PECATTE. *L'oral à l'agrégation de mathématiques, une sélection de développements*. Ellipses, 2024.
- [Kar10] Vladislav KARGIN. « Products of random matrices : dimension and growth in norm ». In : (2010). DOI : 10.1214/09-AAP658.
- [Les] Ahmed LESFARI. *Distributions, analyse de FOURIER et transformation de LAPLACE*. Ellipses.
- [Li] Daniel LI. *Cours d'analyse fonctionnelle*. Ellipses.
- [MT] Rached MNEIMÉ et Frédéric TESTARD. *Introduction à la théorie des groupes de Lie classiques*. Hermann.
- [Ouv00] Jean-Yves OUVARD. *Probabilités 2*. Cassini, 2000. ISBN : 978-2-84225-144-4.
- [Per] Daniel PERRIN. *Cours d'algèbre*. Ellipses.
- [QQ17] Hervé QUEFFÉLEC et Martine QUEFFÉLEC. *Analyse complexe*. Calvage et Mountet, 2017.
- [QZ] Hervé QUEFFÉLEC et Claude ZUILY. *Analyse pour l'agrégation*. DUNOD.
- [Rom] Jean-Etienne ROMBALDI. *Mathématiques pour l'agrégation, Algèbre et géométrie*.
- [Rou15] François ROUVIÈRE. *Petit guide du calcul différentiel*. Cassini, 2015.
- [Tau21] Paul TAUVEL. *Corps commutatifs et théorie de Galois*. Calvage et Mounet, 2021. ISBN : 2916352872.