

ANNALES DES CONCOURS

PC

Mathématiques · Informatique

2015

Sous la coordination de

Guillaume BATOG
Professeur en CPGE
Ancien élève de l'École Normale Supérieure (Cachan)

Julien DUMONT
Professeur en CPGE
Ancien élève de l'École Normale Supérieure (Cachan)

Vincent PUYHAUBERT
Professeur en CPGE
Ancien élève de l'École Normale Supérieure (Cachan)

Par

Juliette BRUN-LELOUP
Professeur en CPGE

Julien DUMONT
Professeur en CPGE

Jean-Julien FLECK
Professeur en CPGE

Damien GARREAU
ENS Ulm

François LÊ
ENS Lyon

Émilie LIBOZ
Professeur en CPGE

Benjamin MONMEGE
Enseignant-chercheur à l'université

Florence MONNA
Docteur en mathématiques

Mathilde PERRIN
Docteur en mathématiques

Sommaire

Énoncé

Corrigé

CONCOURS COMMUNS POLYTECHNIQUES

Mathématiques	Étude analytique et probabiliste d'une suite de fonctions. Matrices binaires. <i>loi de Poisson, séries génératrices, réduction de matrice</i>	17	24
---------------	---	----	----

CENTRALE-SUPÉLEC

Mathématiques 1	Sous-espaces stables d'endomorphismes. <i>réduction, hyperplans, systèmes différentiels, espaces euclidiens</i>	39	42
Mathématiques 2	Étude d'une fonction « bosse » et distributions. <i>représentations graphiques, dérivation, intégrales à paramètre, suite de fonctions</i>	56	59
Informatique	Autour de la dynamique gravitationnelle. <i>listes, boucles, schémas d'intégration, méthode d'Euler, bases de données</i>	82	86

MINES-PONTS

Mathématiques 1	Méthode de Stein. <i>séries numériques, probabilités finies</i>	100	106
Mathématiques 2	Étude des suites de Lucas et de Fibonacci. <i>suites réelles, calcul matriciel, déterminants, probabilités</i>	119	125
Informatique	Tests de validation d'une imprimante. <i>algorithmique, bases de données, méthode d'Euler, méthode des trapèzes</i>	138	148

POLYTECHNIQUE-ENS

Mathématiques	Valeurs propres de matrices symétriques réelles. <i>bornes sup et inf, topologie, valeurs propres</i>	159	163
Informatique	Enveloppes convexes dans le plan. <i>tableaux et listes, boucles for et while, piles, complexité</i>	182	189

FORMULAIRES

Développements limités usuels en 0	198
Développements en série entière usuels	199
Dérivées usuelles	200
Primitives usuelles	201
Trigonométrie	204

Sommaire thématique de mathématiques

2015

e3a PSI Maths A				•		•					•			•
e3a PSI Maths B			•		•									
CCP MP Maths 1								•			•			•
CCP MP Maths 2				•	•									
CCP PC Maths			•								•			•
CCP PSI Maths				•					•			•		
Centrale MP Maths 1											•			
Centrale MP Maths 2								•			•			
Centrale PC Maths 1			•	•										
Centrale PC Maths 2								•		•	•			
Centrale PSI Maths 1							•							•
Centrale PSI Maths 2		•	•								•		•	
Mines MP Maths 1					•			•				•		•
Mines MP Maths 2						•								•
Mines PC Maths 1							•							•
Mines PC Maths 2							•							
Mines PSI Maths 1							•							•
Mines PSI Maths 2			•											
X/ENS MP Maths A		•		•										
X MP Maths B								•		•	•		•	
X/ENS PC Maths				•										
	Structures algébriques et arithmétique	Polynômes	Algèbre linéaire générale	Réduction des endomorphismes	Produit scalaire et espaces euclidiens	Topologie des espaces vectoriels normés	Suites et séries numériques	Suites et séries de fonctions	Séries entières	Analyse réelle	Intégration	Équations différentielles	Fonctions de plusieurs variables	Dénombrement et probabilités

SESSION 2015

PCMA002

**CONCOURS COMMUNS
POLYTECHNIQUES****EPREUVE SPECIFIQUE - FILIERE PC****MATHEMATIQUES****Durée : 4 heures**

N.B. : le candidat attachera la plus grande importance à la clarté, à la précision et à la concision de la rédaction. Si un candidat est amené à repérer ce qui peut lui sembler être une erreur d'énoncé, il le signalera sur sa copie et devra poursuivre sa composition en expliquant les raisons des initiatives qu'il a été amené à prendre.

Les calculatrices sont interdites

L'épreuve est constituée de deux problèmes indépendants.

Lorsqu'un raisonnement utilise un résultat obtenu précédemment dans le problème, il est demandé au candidat d'indiquer précisément le numéro de la question utilisée.

PROBLEME 1 : ANALYSE ET PROBABILITE

On propose d'étudier dans ce premier problème le comportement d'une certaine suite de fonctions $(f_n)_{n \in \mathbb{N}}$ sous différents points de vue. Dans la partie 1, on étudie les aspects analytiques de $(f_n)_{n \in \mathbb{N}}$: convergence uniforme de la suite $(f_n)_{n \in \mathbb{N}}$, propriétés d'intégrales associées à f_n et modes de convergence de la série $\sum f_n$.

La partie 2 correspond à l'étude de la formule de Bernstein : $\lim_{n \rightarrow +\infty} \left(e^{-n} \sum_{k=0}^n \frac{n^k}{k!} \right) = \frac{1}{2}$.

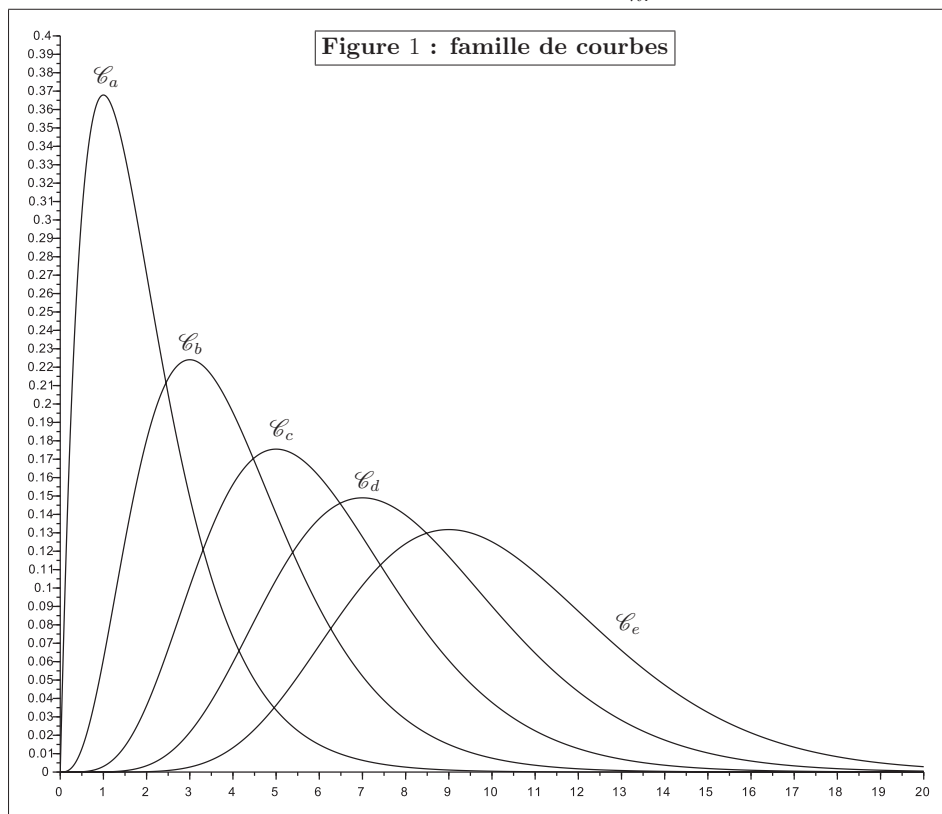
Cette formule, en lien avec la suite de fonctions introduites dans la partie 1, peut être avantageusement interprétée en terme de suite de variables aléatoires indépendantes qui suivent une loi de Poisson. Le point de vue probabiliste permet alors d'éclairer le lien avec une intégrale intervenant en partie 1 et l'existence d'une limite ℓ pour $e^{-n} \sum_{k=0}^n \frac{n^k}{k!}$. Cela permet aussi d'approcher la valeur de ℓ par différentes méthodes.

Les parties 1 et 2 peuvent être traitées, en grande partie, indépendamment l'une de l'autre.

PARTIE 1 : ANALYSE

On considère la fonction f et, pour $n \in \mathbb{N}$, la fonction f_n , définies sur $\mathbb{R}_+$ par :

$$\text{pour tout } t \in \mathbb{R}_+, f(t) = 0 \text{ et } f_n(t) = \frac{e^{-t} t^n}{n!}.$$



CCP Maths PC 2015 — Corrigé

Ce corrigé est proposé par Florence Monna (Docteur en mathématiques) ; il a été relu par Damien Garreau (ENS Ulm) et Gilbert Monna (Professeur en CPGE).

Le sujet est composé de deux problèmes indépendants, le premier traitant d'analyse et de probabilités, le second d'algèbre.

Le premier problème est consacré à une suite de fonctions $(f_n)_{n \in \mathbb{N}}$.

- La partie I étudie la convergence uniforme de la suite, les propriétés de certaines intégrales associées à f_n et les modes de convergence de la série $\sum f_n$. On trouve au début de cette partie les représentations graphiques de quelques fonctions f_n qui mettent en évidence une bosse glissante dont la hauteur tend vers 0, ce qui donne une illustration géométrique intéressante d'une situation classique de convergence uniforme.
- La partie II porte sur la formule de Bernstein, qui peut être interprétée comme une suite de variables aléatoires indépendantes qui suivent une loi de Poisson.

Le second problème s'intéresse aux matrices binaires (à coefficients dans $\{-1, 1\}$) : inversibilité des matrices, orthogonalité des colonnes, propriétés du spectre. On considère en particulier les matrices d'Hadamard (vérifiant $A^T A = n I_n$) et on donne une condition nécessaire portant sur n pour qu'il existe une matrice d'Hadamard de taille n .

Le sujet couvre une large partie du programme de PC : toute l'analyse est sollicitée hormis les équations différentielles ; en probabilités, l'utilisation d'une série génératrice d'une variable aléatoire suivant une loi de Poisson suppose que tout le chapitre ait été traité ; en algèbre, il faut être au point sur les matrices symétriques et la notion d'orthogonalité. Le découpage du sujet en parties cohérentes permet de choisir le morceau de programme que l'on veut travailler.

INDICATIONS

Analyse

- I.1.b Appliquer la formule de Stirling, rappelée dans l'énoncé.
- I.1.c Utiliser le résultat de la question I.1.b.
- I.2.c On est dans le cas d'application du théorème de dérivation d'une intégrale à paramètre. Vérifier toutes les hypothèses du théorème à l'aide des questions I.2.a et I.2.b et conclure.
- I.2.d Exploiter le résultat de la question I.2.a.
- I.2.e Intégrer par parties.
- I.3.a Utiliser la relation de Chasles dans l'expression de la fonction H_n , et conclure à l'aide de la question I.2.e.
- I.3.c Appliquer la propriété de H_n démontrée à la question I.3.b.
- I.3.d Inverser la limite et l'intégrale en appliquant le théorème correspondant, ainsi que le résultat de la question I.1.c.
- I.4.b Regarder les résultats des questions I.1.a, I.1.b et I.2.e.
- I.5.b Utiliser les questions I.1.a et I.5.a.

Probabilités

- II.2 Appliquer la formule de Taylor rappelée dans l'énoncé à la fonction $f : x \mapsto e^x$ sur $\mathbb{R}$, avec $a = 0$, $b = n$.
- II.3.a Utiliser le résultat de la question précédente.
- II.3.b Intégrer par parties.
- II.3.c Exploiter les questions I.1.a et II.3.b, ainsi que le théorème de la limite monotone.
- II.4.c Utiliser les résultats des questions II.4.a et II.4.b.

Algèbre

- 2 Développer le déterminant de A_3 par rapport à la première colonne.
- 5 Démontrer que i) implique ii), puis que ii) implique i) et terminer en montrant que i) est équivalente à iii).
- 6.b Utiliser le résultat de la question 6.a.
- 6.d Appliquer les résultats des questions 4 et 6.c.
- 8 Penser à exploiter la question 7.c.

ANALYSE

I.1.a Soit $n \in \mathbb{N}^*$. La fonction f_n est dérivable sur $\mathbb{R}_+$ par produit d'une exponentielle et d'un polynôme, et

$$\forall t \in \mathbb{R}_+ \quad f'_n(t) = \frac{t^{n-1}e^{-t}}{n!}(n-t)$$

Par ailleurs, $e^{-t}t^n$ tend vers 0 quand t tend vers 0 et quand t tend vers l'infini, par croissances comparées. On en déduit le tableau de variations ci-contre. Ainsi, f_n admet un unique maximum sur $[n; +\infty[$, atteint en n , qui vaut

t	0	n	$+\infty$
f'	+	0	-
f_n	0	$f_n(n)$	0

$$f_n(n) = \frac{e^{-n}n^n}{n!}$$

I.1.b On peut réaliser le quotient de deux équivalents tant que le dénominateur ne s'annule pas. Puisque $\sqrt{2\pi n}n^n e^{-n}$ ne s'annule pas,

$$f_n(n) \sim \frac{e^{-n}n^n}{\sqrt{2\pi n}n^n e^{-n}}$$

soit

$$f_n(n) \sim \frac{1}{\sqrt{2\pi n}} = \frac{1}{\sqrt{2\pi}} \frac{1}{n^{1/2}}$$

I.1.c Le tableau de variation de f permet d'établir que la fonction f_n est positive sur $\mathbb{R}_+$, si bien que $|f_n| = f_n$. D'après la question 1.a, on en déduit

$$\|f_n - f\|_{\infty, \mathbb{R}_+} = \|f_n\|_{\infty, \mathbb{R}_+} = \max_{t \in \mathbb{R}_+} |f_n(t)| = f_n(n)$$

Or, d'après la question 1.b, $f_n(n) \xrightarrow{n \rightarrow \infty} 0$, ce qui permet de conclure que

La suite de fonctions $(f_n)_{n \in \mathbb{N}}$ converge uniformément vers f sur $\mathbb{R}_+$.

I.2.a L'intégrale $I(x) = \int_0^{+\infty} e^{-t}t^x dt$ est impropre en $+\infty$ et en 0 seulement.

- Lorsque $t \rightarrow +\infty$, $e^{-t}t^x = o(1/t^2)$ et, par comparaison d'intégrales de fonctions positives, l'intégrale $I(x)$ est convergente quel que soit $x \in \mathbb{R}$.
- En 0, $e^{-t}t^x \sim 1/t^{-x}$, et, à nouveau par comparaison d'intégrales de fonctions positives, l'intégrale $I(x)$ est convergente si et seulement si $-x < 1$.

Finalement, $\int_0^{+\infty} e^{-t}t^x dt$ est convergente si et seulement si $x > -1$:

$$D =]-1; +\infty[\quad \text{et} \quad \mathbb{R}_+ \subset D$$

I.2.b L'intégrale $J(x) = \int_0^{+\infty} (\ln t)e^{-t}t^x dt$ est impropre en $+\infty$ et 0 seulement.

Quand t tend vers l'infini, $(\ln t)e^{-t}t^x = o(1/t^2)$ et, par comparaison d'intégrales de fonctions positives, l'intégrale $J(x)$ est convergente quel que soit $x \in \mathbb{R}$.

En 0, distinguons deux cas :

- Si $x = 0$, $\ln(t)e^{-t} \sim \ln(t)$ et on sait que l'intégrale $\int_0^1 \ln(t) dt$ est convergente.
- Si $x > 0$, $(\ln t)e^{-t}t^x$ tend vers 0 quand t tend vers 0, par croissances comparées, donc la fonction intégrée admet un prolongement par continuité en 0.

On conclut alors que

Pour tout $x \in \mathbb{R}_+$, l'intégrale $\int_0^{+\infty} (\ln t)e^{-t}t^x dt$ est convergente.

I.2.c On cherche à se placer dans le cas d'application du théorème de dérivation d'une intégrale à paramètre. On pose pour cela $I = \mathbb{R}_+$, $J =]0; +\infty[$, ainsi que

$$g(x, t) = e^{-t}t^x \text{ sur } I \times J \quad \text{et} \quad G(x) = \int_J g(x, t) dt \text{ sur } J$$

Le but est alors de montrer que la fonction G est de classe $\mathcal{C}^1$ sur I . Vérifions successivement les hypothèses du théorème.

- On peut écrire $|g(x, t)| = g(x, t)$ et le résultat de la question 2.a permet de conclure que pour tout $x \in I$, la fonction $t \mapsto g(x, t) = e^{-t}t^x$ est continue par morceaux et intégrable sur J .
- Pour tout $t \in J$, la fonction $x \mapsto g(x, t)$ est de classe $\mathcal{C}^1$ sur I et

$$\frac{\partial g}{\partial x}(x, t) = e^{-t}(\ln t)e^{x \ln t} = e^{-t}(\ln t)t^x$$

- Pour $x \in I$, $t \mapsto \frac{\partial g}{\partial x}(x, t) = e^{-t}(\ln t)t^x$ est continue par morceaux sur J .
- Reste l'hypothèse de domination sur tout segment $[0; b]$ de I . Soit $0 < b$ et notons $K = [0; b]$. On peut écrire

$$\forall x \in K \quad \forall t \in J \quad \left| \frac{\partial g}{\partial x}(x, t) \right| \leq \phi(t)$$

en posant
$$\phi(t) = \begin{cases} e^{-t} |\ln t| t^0 & \text{sur }]0; 1] \\ e^{-t}(\ln t)t^b & \text{sur } [1; +\infty[\end{cases}$$

On constate que ϕ est une fonction positive, continue par morceaux sur J .

De plus, la fonction ϕ est intégrable sur l'intervalle J d'après la question 2.b.

Les hypothèses du théorème de dérivation d'une intégrale à paramètre étant toutes vérifiées, on en déduit que G est de classe $\mathcal{C}^1$ sur I et

$$\forall x \in I \quad G'(x) = \int_0^{+\infty} e^{-t}(\ln t)t^x dt$$

d'où

La fonction $x \mapsto \int_0^{+\infty} e^{-t}t^x dt$ est de classe $\mathcal{C}^1$ sur $\mathbb{R}_+$.

I.2.d Le résultat de la question 2.a permet d'écrire $n \in D$ pour tout $n \in \mathbb{N}$, donc l'intégrale $\int_0^{+\infty} e^{-t}t^n dt$ converge, tout comme l'intégrale $\int_0^{+\infty} \frac{e^{-t}t^n}{n!} dt$. Par suite,

L'intégrale $\int_0^{+\infty} f_n(t) dt$ est convergente.

Centrale Maths 1 PC 2015 — Corrigé

Ce corrigé est proposé par Damien Garreau (ENS Ulm) ; il a été relu par Gilbert Monna (Professeur en CPGE) et Nicolas Martin (Professeur agrégé).

Ce sujet est consacré à l'étude des sous-espaces stables d'un endomorphisme. Ses cinq parties sont indépendantes.

- La première partie rappelle les liens entre sous-espaces stables et diagonalisabilité. En particulier, on montre que lorsque le corps de base est $\mathbb{C}$, la diagonalisabilité est équivalente à l'existence d'un supplémentaire stable pour tout sous-espace stable.
- Dans la deuxième partie, on montre qu'un sous-espace F de E est stable si et seulement si $F = \bigoplus_{i=1}^p (F \cap E_i)$ où les $(E_i)_{1 \leq i \leq p}$ sont les sous-espaces propres. On utilise ce résultat pour dénombrer les sous-espaces stables dans le cas où $p = n$.
- La troisième partie est consacrée à l'étude des sous-espaces stables d'un endomorphisme nilpotent. Après avoir étudié le cas particulier de l'endomorphisme de dérivation sur $\mathbb{K}[X]$, on montre un résultat général.
- Dans la quatrième partie, on montre que tout endomorphisme d'un espace vectoriel réel de dimension finie admet au moins une droite ou un plan stable. On détermine ensuite les solutions $(x(t), y(t), z(t))$ d'un système différentiel de taille 3 donné qui représentent une paramétrisation d'une droite ou d'un arc plan de $\mathbb{R}^3$.
- La dernière partie caractérise les endomorphismes diagonalisables dans un cadre euclidien : un endomorphisme est diagonalisable si, et seulement si, il admet n hyperplans stables d'intersection réduite au vecteur nul.

Les parties de ce problème sont équilibrées et chacune comporte des questions délicates. La progression est linéaire et bien guidée. On se servira utilement de ce sujet pour faire le point sur l'algèbre linéaire.

INDICATIONS

Partie I

- I.B.1 Penser aux sous-espaces vectoriels triviaux.
- I.B.2 Utiliser le noyau et l'image de l'endomorphisme et le théorème du rang pour le dernier cas.
- I.C.3 Montrer que c'est une homothétie.
- I.D.1 Utiliser le théorème de la base incomplète.
- I.D.2 Considérer $F = \bigoplus_{\lambda \in \text{sp}(\lambda)} E_\lambda$, puis montrer que $F = E$.

Partie II

- II.A.3 Reconnaître un déterminant de Vandermonde.
- II.A.5 Utiliser la question II.A.3.
- II.A.6 Utiliser la question II.A.2.
- II.B.2 Utiliser les questions II.B.1, II.A et I.A.
- II.B.3 Utiliser la question II.A.
- II.B.4 Reconnaître $(1 + 1)^n$, puis utiliser la formule du binôme.

Partie III

- III.A.2.a Considérer une base de F .
- III.A.3 Utiliser la question III.A.2.c.
- III.B.1 Montrer qu'il s'agit de $E \setminus \text{Ker}(f^{n-1})$.
- III.B.3 Multiplier les éléments de $\mathcal{B}_{f,u}$ par des constantes bien choisies.

Partie IV

- IV.B Utiliser le fait qu'un polynôme réel de degré impair admet toujours au moins une racine réelle.
- IV.C.1 Raisonner par l'absurde en montrant que $\lambda \in \mathbb{R}$.
- IV.C.2 Calculer M_Z de deux manières différentes, puis identifier les parties réelles et imaginaires.
- IV.D Utiliser les questions IV.B et IV.C.
- IV.F.1 Calculer les valeurs propres de A et les vecteurs propres associés.
- IV.F.2 Utiliser la question IV.F.1.
- IV.F.3 Calculer x'' , puis faire disparaître les termes en y .
- IV.F.4 Effectuer le changement de variable $Y = P^{-1}X$ puis utiliser les questions précédentes.

Partie V

- V.A.1 Définir le produit scalaire canonique.
- V.A.2 Montrer que $u \cdot v = {}^tU V$.
- V.C Déterminer les vecteurs propres de tA , puis utiliser la question V.B.
- V.D Utiliser l'équivalence entre A diagonalisable et tA diagonalisable.

1. PREMIÈRE PARTIE

I.A Supposons que u soit un vecteur propre de f . Il existe alors $\lambda \in \mathbb{K}$ tel que $f(u) = \lambda u$. Pour tout $x \in F$, il existe $\mu \in \mathbb{K}$ tel que $x = \mu u$. Comme

$$f(x) = f(\mu u) = \mu f(u) = \mu \lambda u \in F$$

La droite F est stable par f .

Réciproquement, supposons que F soit stable par f . Soit u un vecteur non nul de F . Par hypothèse, $f(u) \in F$ donc il existe $\lambda \in \mathbb{K}$ tel que $f(u) = \lambda u$.

Le vecteur u est propre pour f .

I.B.1 Comme f est un endomorphisme, il laisse E stable, et $f(0) = 0$. Ainsi,

Le sous-espace nul $\{0\}$ et E tout entier sont stables par f .

L'idée dans $\mathbb{R}^2$ est de trouver une application linéaire qui ne laisse stable aucune droite. Considérons donc la rotation d'angle θ . Cette application est linéaire, et elle envoie une droite sur une droite distincte pourvu que θ ne soit pas congru à 0 modulo π .

I.B.2 Si $x \in \text{Ker}(f)$, alors $f(f(x)) = f(0) = 0$, autrement dit $\text{Ker}(f)$ est stable par f . Puisque f n'est pas injective, $\text{Ker}(f) \neq \{0\}$. Comme f est non nulle, $\text{Ker}(f)$ est distinct de E . Ainsi, il existe au moins trois sous-espaces distincts stables par f .

Les sous-espaces $\{0\}$, $\text{Ker}(f)$ et E sont stables par f .

Supposons maintenant que n est impair. Soit $y = f(x) \in \text{Im}(f)$. Alors

$$f(y) = f(f(x)) \in \text{Im}(f)$$

autrement dit $\text{Im}(f)$ est stable par f . Montrons que l'image de f est distincte de $\{0\}$, $\text{Ker}(f)$ et E . Le théorème du rang appliqué à f s'écrit

$$\text{rg}(f) + \dim \text{Ker}(f) = \dim E = n$$

Comme n est impair, on ne peut avoir $\text{rg}(f) = \dim \text{Ker}(f)$. Par conséquent, l'image et le noyau de f sont distincts. Vu que f est non nulle, son image n'est pas réduite au vecteur nul. Enfin, f est non injective, d'où $\dim \text{Ker}(f) \neq 0$, ce qui implique que $\text{rg}(f) \neq n$, et donc l'image de f n'est pas E tout entier. Lorsque n est impair, il existe donc au moins quatre sous-espaces stables distincts.

Les sous-espaces $\{0\}$, $\text{Ker}(f)$, $\text{Im}(f)$ et E sont stables par f .

Considérons l'endomorphisme de $\mathbb{R}^2$ de matrice

$$M = \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix}$$

dans la base canonique. Pour trouver les droites stables de f , cherchons les vecteurs propres de f conformément à la question I.A. Il n'y a qu'une valeur propre $\lambda = 1$, et l'identité $MX = X$ conduit à $x = 0$. En conclusion, f ne possède qu'une seule droite propre, d'équation $x = 0$.

I.C.1 Soient $v_1, v_2, \dots, v_p$ des vecteurs propres de f associés aux valeurs propres $\lambda_1, \lambda_2, \dots, \lambda_p$. Montrons que $V = \text{Vect}(v_1, v_2, \dots, v_p)$ est stable par f . Soit $x \in V$. Il existe $(x_1, x_2, \dots, x_p) \in \mathbb{K}^p$ tel que $x = \sum_{i=1}^p x_i v_i$.

$$f(x) = \sum_{i=1}^p x_i f(v_i) = \sum_{i=1}^p x_i (\lambda_i v_i) = \sum_{i=1}^p (\lambda_i x_i) v_i \in V$$

Tout sous-espace engendré par une famille de vecteurs propres de f est stable par f .

Considérons le sous-espace propre E_λ associé à une valeur propre λ . Par définition, pour tout $x \in E_\lambda$, on a $f(x) = \lambda x$. Ainsi,

La restriction de f à E_λ est l'homothétie de rapport λ .

I.C.2 Soit E_λ un sous-espace propre de f de dimension au moins égale à 2. D'après la question I.C.1, la restriction de f à E_λ est une homothétie de rapport λ . En particulier, une homothétie laisse les droites stables. Comme $\dim E_\lambda \geq 2$ et $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$, E_λ contient une infinité de droites :

L'endomorphisme f laisse stable une infinité de droites.

I.C.3 Si tous les sous-espaces de E sont stables par f , alors en particulier toutes les droites de E sont stables par f . D'après la question I.A, tous les vecteurs non nuls de E sont vecteurs propres de f .

| C'est un exercice classique d'en déduire que f est une homothétie.

Pour tout $x \in E \setminus \{0\}$, notons λ_x l'élément de $\mathbb{K} \setminus \{0\}$ tel que $f(x) = \lambda_x x$. Fixons $x_0 \in E \setminus \{0\}$ et montrons que tous les λ_x coïncident avec λ_{x_0} .

- Pour tout x tel que la famille (x, x_0) soit liée, il existe $\mu \in \mathbb{K}$ tel que $x = \mu x_0$ car $x_0 \neq 0$ donc

$$f(x) = f(\mu x_0) = \mu f(x_0) = \mu \lambda_{x_0} x_0 = \lambda_{x_0} x$$

d'où $\lambda_x = \lambda_{x_0}$ par unicité de λ_x .

- Pour tout x tel que la famille (x, x_0) soit libre,

$$f(x + x_0) = \lambda_{x+x_0}(x + x_0) = \lambda_{x+x_0}x + \lambda_{x+x_0}x_0$$

$$\text{et} \quad f(x + x_0) = f(x) + f(x_0) = \lambda_x x + \lambda_{x_0} x_0$$

$$\text{En soustrayant,} \quad 0 = (\lambda_{x+x_0} - \lambda_x)x + (\lambda_{x+x_0} - \lambda_{x_0})x_0$$

Puisque (x, x_0) est libre, $\lambda_{x+x_0} = \lambda_x$ et $\lambda_{x+x_0} = \lambda_{x_0}$ d'où $\lambda_x = \lambda_{x_0}$.

Puisque $f(x) = \lambda_{x_0} x$ pour tout $x \in E \setminus \{0\}$ et pour $x = 0$,

L'endomorphisme f est une homothétie.

I.D.1 Supposons f diagonalisable. Soit $\mathcal{B}_D = (v_1, v_2, \dots, v_n)$ une base de vecteurs propres. Soit F un sous-espace stable de E . Notons p la dimension de F . Si $p = 0$ ou n , il n'y a rien à prouver. Sinon, considérons $\mathcal{B}_F = (e_1, e_2, \dots, e_p)$ une base de F . D'après le théorème de la base incomplète, on peut compléter $\mathcal{B}_F$ en une base $\mathcal{B}$ de E avec $n - p$ vecteurs de $\mathcal{B}_D$. Quitte à réordonner $\mathcal{B}_D$, on peut supposer qu'il s'agit de

Centrale Maths 2 PC 2015 — Corrigé

Ce corrigé est proposé par Mathilde Perrin (Docteur en mathématiques) ; il a été relu par Émilie Liboz (Professeur en CPGE) et Benjamin Monmege (Enseignant-chercheur à l'université).

Ce sujet d'analyse porte sur les distributions, qui généralisent la notion de fonction. Elles sont utiles aux physiciens, qui les ont manipulées bien avant leur formation, afin d'obtenir des « solutions » à des problèmes discontinus. Elles ont de nombreuses applications en ingénierie, en traitement du signal...

- La partie I.A définit l'espace $\mathcal{D}$ des fonctions de classe $\mathcal{C}^\infty$ de $\mathbb{R}$ dans $\mathbb{R}$ qui sont nulles en dehors d'un segment. On montre que $\mathcal{D}$ est un espace vectoriel stable par dérivation qui contient une fonction « bosse » donnée par l'énoncé. Puis on construit une suite de fonctions de $\mathcal{D}$ définies par l'intégrale sur $[-1; 1]$ d'une fonction bosse dilatée. Cette partie très classique représente un tiers du problème et porte sur de nombreux outils d'analyse : représentation graphique, prolongement dérivable, intégrale à paramètre, convergence uniforme... Ses résultats ne sont pas réutilisés dans la suite.
- La partie I.B est complètement indépendante des autres. Elle établit quelques résultats théoriques sur l'espace de Schwartz des fonctions de $\mathbb{R}$ dans $\mathbb{R}$ de classe $\mathcal{C}^\infty$ à décroissance rapide, c'est-à-dire dont toutes les dérivées sont négligeables en $+\infty$ et $-\infty$ par rapport à toute fonction puissance. Cet espace intervient en théorie des distributions, d'où sa présence dans le problème.
- La partie II.A définit une distribution comme une application linéaire T sur $\mathcal{D}$ à valeurs dans $\mathbb{R}$ qui est continue en un sens faisant intervenir la convergence uniforme des dérivées de fonctions de $\mathcal{D}$. On étudie des exemples comme les distributions régulières T_f définies à partir d'une fonction réelle f continue par morceaux sur $\mathbb{R}$ par

$$\forall \varphi \in \mathcal{D} \quad T_f(\varphi) = \int_{-\infty}^{+\infty} f(x)\varphi(x) \, dx$$

ou comme les distributions de Dirac $\delta_a : \varphi \mapsto \varphi(a)$ pour $a \in \mathbb{R}$ qui ne sont pas régulières. Le but est de se familiariser avec cette nouvelle notion.

- La partie II.B introduit la dérivation de distributions. Elle s'intéresse plus particulièrement aux distributions régulières T_f , dont on donne une formule générale de dérivation pour les fonctions f de classe $\mathcal{C}^1$ sauf en un nombre fini de points :

$$(T_f)' = T_{f'} + \sum_i \sigma(a_i) \delta_{a_i}$$

où $\sigma(a_i)$ représente le saut de discontinuité de f au point a_i . En ce sens, on peut définir une « dérivée » d'une fonction non dérivable !

- La partie II.C étudie des suites de distributions et de dérivées de distributions. L'épreuve se termine sur une question difficile qui demande d'étudier la convergence de trois suites de distributions régulières.

Cette épreuve est longue et difficilement faisable en 4 heures, mais elle a le mérite d'aborder un sujet très important en mathématiques. Elle constitue un bon sujet de révision ou d'entraînement sur l'analyse réelle, les suites de fonctions, l'intégration et tout particulièrement le théorème de convergence dominée.

INDICATIONS

- I.A.1.c Montrer par récurrence que φ est de classe $\mathcal{C}^k$ sur $\mathbb{R}$ en établissant une formule générale pour $\varphi^{(k)}$ et en utilisant le théorème de la limite de la dérivée.
- I.A.1.d Montrer que $\mathcal{D}$ est un sous-espace vectoriel de $\mathcal{F}(\mathbb{R}, \mathbb{R})$ contenant φ .
- I.A.3.b Effectuer le changement de variable $y = nx$.
- I.A.4 Montrer que la fonction $f * \rho_n$ est de classe $\mathcal{C}^k$ sur $\mathbb{R}$ pour tout entier naturel k en utilisant le théorème de la classe $\mathcal{C}^k$ des intégrales à paramètre.
- I.A.5.b Montrer que la fonction I_n est nulle en dehors de $[-1 - 1/n; 1 + 1/n]$ en utilisant l'expression de I_n obtenue à la question I.A.5.a.
- I.A.5.d Utiliser l'expression de I_n obtenue à la question I.A.5.a. et étudier la convergence de la suite $(I_n(x))_{n \in \mathbb{N}^*}$ dans les cas $|x| < 1$, $|x| = 1$ et $|x| > 1$.
- I.A.5.e Raisonner par l'absurde.
- I.B.3 Utiliser la formule de Leibniz.
- II.A.1 Appliquer le théorème de convergence dominée sur $[-a; a]$ à une suite de fonctions bien choisie.
- II.A.3.b Comparer $\delta_a(\varphi_n)$ et $\lim_{n \rightarrow \infty} T_f(\varphi_n)$.
- II.B.1 Montrer que si $(\varphi_n)_{n \in \mathbb{N}}$ est une suite de fonctions de $\mathcal{D}$ qui converge dans $\mathcal{D}$ vers $\varphi \in \mathcal{D}$, alors la suite $(\varphi_n')_{n \in \mathbb{N}}$ converge dans $\mathcal{D}$ vers φ' .
- II.B.2 Lorsque f est de classe $\mathcal{C}^1$, effectuer une intégration par parties pour des fonctions de classe $\mathcal{C}^1$. Lorsque f est supposée continue et de classe $\mathcal{C}^1$ par morceaux, choisir une subdivision adaptée à f et décomposer (en utilisant la relation de Chasles) l'intégrale sur les intervalles de cette subdivision, où la fonction f est de classe $\mathcal{C}^1$. Appliquer ensuite la formule d'intégration par parties pour des fonctions de classe $\mathcal{C}^1$ sur chaque intervalle de la subdivision.
- II.B.5.a Décomposer (en utilisant la relation de Chasles) l'intégrale sur les intervalles de la subdivision $(a_i)_{1 \leq i \leq p}$ donnée, sur lesquels la fonction f est prolongeable en une fonction continue et de classe $\mathcal{C}^1$ par morceaux. Appliquer la formule démontrée à la question II.B.2 sur chaque intervalle.
- II.C.1.a Effectuer le changement de variable $u = nt$ puis appliquer le théorème d'intégration d'une limite uniforme sur un segment à la suite de fonctions $f_n(u) = u\varphi(u/n)$ sur $[0; 1]$.
- II.C.1.b Calculer directement la dérivée $(T_{U_n})'$ en utilisant la formule d'intégration par parties pour les fonctions de classe $\mathcal{C}^1$ sur $[0; 1/n]$, ou utiliser le résultat de la question II.B.2 pour les fonctions continues et de classe $\mathcal{C}^1$ par morceaux.
- II.C.1.f Utiliser les résultats des questions II.C.1.a, II.C.1.e et II.B.3.
- II.C.2.c Montrer que la suite $(T_{f_n})_{n \in \mathbb{N}^*}$ tend vers la distribution $\pi\delta_0$ lorsque $n \rightarrow \infty$ en remarquant que pour tout $n \in \mathbb{N}^*$, $\int_{-\infty}^{+\infty} f_n(x) dx = \pi$. De la même manière, montrer que la suite $(T_{g_n})_{n \in \mathbb{N}^*}$ tend vers la distribution δ_1 en remarquant que $\int_{-\infty}^{+\infty} g_n(x) dx = n/(n+1) \xrightarrow[n \rightarrow \infty]{} 1$. Enfin, effectuer une intégration par parties pour calculer $T_{h_n}(\varphi)$ puis conclure en utilisant l'égalité des accroissements finis d'une part, le théorème de convergence dominée d'autre part.

I. ÉTUDE DE NOUVEAUX ESPACES FONCTIONNELS

I.A.1.a Commençons par étudier la régularité de φ . Puisqu'elle est paire, il suffit de l'étudier sur $\mathbb{R}_+$. Elle est continue sur $]0; 1[$ comme composée de fonctions qui le sont, et sur $]1; +\infty[$ où elle est constante. Étudions ses limites à gauche et à droite en $x = 1$. Par définition de φ , la limite à droite est nulle. Lorsque x tend vers 1 par valeurs inférieures, on a

$$\lim_{x \rightarrow 1^-} (1 - x^2) = 0^+ \quad \text{donc} \quad \lim_{x \rightarrow 1^-} -\frac{x^2}{1 - x^2} = -\infty \quad \text{puis} \quad \lim_{x \rightarrow 1^-} \varphi(x) = 0$$

Puisque $\lim_{x \rightarrow 1^-} \varphi(x) = \lim_{x \rightarrow 1^+} \varphi(x) = \varphi(1) = 0$, φ est continue en 1 donc sur $\mathbb{R}_+$ entier.

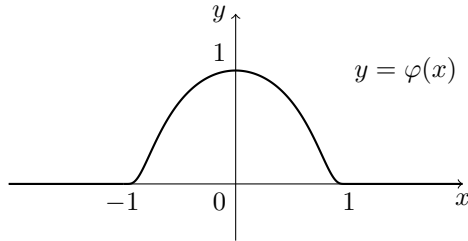
φ est dérivable sur $]0; 1[$ comme composée de fonctions qui le sont, avec

$$\forall x \in]0; 1[\quad \varphi'(x) = -\frac{2x}{(1 - x^2)^2} \exp\left(-\frac{x^2}{1 - x^2}\right) \leq 0$$

d'où le tableau de variation suivant :

x	$-\infty$	-1	0	1	$+\infty$
$\varphi'(x)$	0	$+$	0	$-$	0
$\varphi(x)$	$0 \longrightarrow 0 \nearrow 1 \searrow 0 \longrightarrow 0$				

I.A.1.b La représentation graphique de φ admet une tangente horizontale en $x = 0$.



D'après un calcul effectué à la question suivante,

$$\lim_{x \rightarrow 1^-} \varphi'(x) = \lim_{x \rightarrow 1^+} \varphi'(x) = 0$$

Cela assure que le graphe de φ admet une demi-tangente horizontale à gauche en $x = 1$ (et une demi-tangente horizontale à droite en $x = -1$, par parité).

I.A.1.c Par parité de φ , il suffit de montrer qu'elle est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}_+$. Elle est constante donc de classe $\mathcal{C}^\infty$ sur $]1; +\infty[$. Sur $]0; 1[$, la fonction $x \mapsto -x^2/(1 - x^2)$ est de classe $\mathcal{C}^\infty$ comme produit de fonctions qui le sont. L'exponentielle étant de classe $\mathcal{C}^\infty$ sur $]0; 1[$, φ l'est aussi comme composée de fonctions qui le sont.

Étudions maintenant φ en $x = 1$. Montrons par récurrence que la propriété

$\mathcal{P}(k)$: « φ est de classe $\mathcal{C}^k$ sur $\mathbb{R}_+$ et il existe un polynôme à coefficients réels P_k tel que pour tout $x \in]0; 1[$

$$\varphi^{(k)}(x) = \frac{P_k(x)}{(1 - x^2)^{2k}} \exp\left(-\frac{x^2}{1 - x^2}\right)$$

et $\varphi^{(k)}(x) = 0$ pour $x \in [1; +\infty[$. »

est vraie pour tout $k \in \mathbb{N}$.

- $\mathcal{P}(0)$ est vérifiée avec le polynôme $P_0(x) = 1$ car φ est continue au point 1 d'après la question I.A.1.a.

- $\mathcal{P}(k) \implies \mathcal{P}(k+1)$: $\varphi^{(k)}$ est dérivable sur $]1; +\infty[$, de dérivée nulle, et dérivable sur $]0; 1[$ comme produit et composée de fonctions qui le sont, avec

$$\forall x \in [0; 1[\quad \varphi^{(k+1)}(x) = \frac{P_{k+1}(x)}{(1-x^2)^{2(k+1)}} \exp\left(-\frac{x^2}{1-x^2}\right)$$

$$\text{où} \quad P_{k+1}(x) = P_k'(x)(1-x^2)^2 + 4kxP_k(x)(1-x^2) - 2xP_k(x)$$

qui est bien un polynôme réel.

Par hypothèse de récurrence, $\varphi^{(k)}$ est continue sur $\mathbb{R}_+$ et dérivable sur $\mathbb{R}_+ \setminus \{1\}$. Le théorème de la limite de la dérivée permet alors d'affirmer que φ est de classe $\mathcal{C}^{k+1}$ sur $\mathbb{R}$ avec $\varphi^{(k+1)}(1) = 0$.

Calculons la limite de $\varphi^{(k+1)}$ en $x = 1$. Puisqu'elle est nulle sur $]1; +\infty[$, sa limite à droite en 1 est également nulle. En utilisant la croissance comparée

$$\forall m \in \mathbb{N} \quad \lim_{y \rightarrow +\infty} y^m e^{-y} = 0$$

avec $y = 1/(1-x^2)$ lorsque x tend vers 1 par valeurs inférieures, l'expression de $\varphi^{(k+1)}$ sur $]0; 1[$ établie ci-dessus entraîne

$$\lim_{x \rightarrow 1^-} \varphi^{(k+1)}(x) = 0$$

Les limites à gauche et à droite de $\varphi^{(k+1)}$ en $x = 1$ coïncident, donc la limite de la fonction $\varphi^{(k+1)}$ au point $x = 1$ existe et est égale à 0. Par le théorème de la limite de la dérivée, on en déduit que $\varphi^{(k)}$ est dérivable en $x = 1$ et $\varphi^{(k+1)}(1) = 0$. Ainsi la fonction $\varphi^{(k+1)}$ est continue sur $\mathbb{R}_+$ et φ est de classe $\mathcal{C}^{k+1}$ sur $\mathbb{R}_+$, ce qui achève de démontrer que la propriété $\mathcal{P}(k+1)$ est vraie.

- Conclusion :

 φ est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$.

I.A.1.d Montrons que l'ensemble $\mathcal{D}$ est un sous-espace vectoriel du $\mathbb{R}$ -espace vectoriel $\mathcal{F}(\mathbb{R}, \mathbb{R})$. D'une part, on observe que la fonction identiquement nulle appartient à $\mathcal{D}$. D'autre part, vérifions que $\mathcal{D}$ est stable par combinaisons linéaires. Soient $\lambda \in \mathbb{R}$ et $f, g \in \mathcal{D}$ nulles en dehors des segments $[a; b]$ et $[c; d]$ respectivement. La combinaison linéaire $\lambda f + g$ est clairement de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$ par les théorèmes généraux. Elle est de plus à support compact puisque nulle en dehors du segment $[\min(a, c); \max(b, d)]$ car ce segment contient à la fois $[a; b]$ et $[c; d]$. Ainsi $\lambda f + g \in \mathcal{D}$ et $\mathcal{D}$ est stable par combinaisons linéaires.

On a démontré à la question I.A.1.c que la fonction φ est de classe $\mathcal{C}^\infty$, et par définition elle est nulle en dehors du segment $[-1; 1]$. On en déduit que $\varphi \in \mathcal{D}$, et comme φ est non nulle,

 $\mathcal{D}$ est un espace vectoriel sur $\mathbb{R}$ non réduit à $\{0\}$.

I.A.2 Soit $f \in \mathcal{D}$. La fonction f étant de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$, elle est dérivable sur $\mathbb{R}$ et sa dérivée f' est également de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$. Par ailleurs, soient a et b deux réels vérifiant $a < b$ tels que f soit nulle en dehors du segment $[a; b]$. Sa dérivée f' est par conséquent elle aussi nulle en dehors du segment $[a; b]$, elle est donc à support compact. On en déduit que

Si $f \in \mathcal{D}$ alors $f' \in \mathcal{D}$.

Centrale Informatique PC 2015 — Corrigé

Ce corrigé est proposé par Jean-Julien Fleck (Professeur en CPGE) ; il a été relu par Julien Dumont (Professeur en CPGE) et Guillaume Batog (Professeur en CPGE).

Ce sujet s'intéresse à une application en astronomie (effectivement utilisée par les astronomes) d'un algorithme d'intégration numérique, appelé schéma d'intégration de Verlet.

La première partie part du constat qu'en Python les listes ne sont pas directement adaptées pour faire du calcul vectoriel puisque leur addition à l'aide du symbole `+` conduit à la concaténation des deux listes (qui du coup peuvent être de tailles quelconques) plutôt qu'à l'addition terme à terme comme on pourrait s'y attendre si l'on représente des vecteurs par des listes. On construit donc quelques fonctions qui permettent de pallier cet inconvénient, comme l'addition ou la soustraction terme à terme ou encore la multiplication par un scalaire. Ces questions servent essentiellement à vérifier que le candidat maîtrise les boucles sur les listes.

La deuxième partie est une étude plutôt théorique des avantages comparés des schémas d'intégration d'Euler et de Verlet pour résoudre numériquement une équation différentielle¹. L'accent est mis sur l'aspect mathématique, mais on fait aussi des liens avec le cours de mécanique sur les oscillateurs à un degré de liberté puisque l'on étudie principalement l'effet du schéma d'intégration sur un oscillateur harmonique et sa représentation graphique en terme de portrait de phase. Les implémentations du schéma d'Euler (prolongement naturel et direct du cours) et du schéma de Verlet (variation sur le thème d'Euler) constituent les applications informatiques de cette partie.

La troisième partie se concentre sur l'implémentation du schéma de Verlet dans le cadre classique du problème gravitationnel à N corps. Il s'agit dans un premier temps de coder le calcul de la force gravitationnelle d'un objet sur un autre, puis de tous les autres objets sur l'objet d'étude pour finalement intégrer les équations du mouvement à l'aide du schéma de Verlet. On termine sur une estimation expérimentale et théorique de la complexité de l'algorithme proposé.

La quatrième partie aborde des questions sur les bases de données dans l'idée de collecter une série de corps issus du système solaire pour démarrer une simulation à grande échelle. Cela commence par des requêtes SQL très simples, puis la difficulté augmente progressivement. Le problème se termine par l'écriture de la fonction simulant le système solaire, en partant des conditions initiales collectées via la base de données. Elle utilise les fonctions de la partie III.

L'épreuve est progressive, même si elle démarre avec des questions légèrement hors programme. Elle construit presque complètement un intégrateur à N corps qui pourra servir de base, par exemple, à de futurs TIPE sur ce thème. La longueur est raisonnable et le sujet semble parfaitement correspondre à ce qu'on est en droit d'attendre sur le programme d'informatique commune.

¹Sans surprise, c'est Verlet qui va gagner.

INDICATIONS

Partie I

- I.A.1 Attention, les listes Python ne se comportent pas naturellement comme des vecteurs en mathématiques. Dans le doute, tester le code dans une session interactive de Python.

Partie II

- II.A.1 Par définition, on a $y' = z$, il reste donc à trouver à quoi relier z' .
- II.A.2 Il suffit de partir de la constatation que $\int_{t_i}^{t_{i+1}} y'(t) dt = y(t_{i+1}) - y(t_i)$.
- II.B.2 Utiliser la fonction **f** définissant l'équation différentielle, les conditions initiales **y0** et **z0** ainsi que le pas de temps **h** et le nombre **n** de points.
- II.B.3.a Penser au cours de physique : multiplier par la vitesse y' pour faire apparaître des dérivées de fonctions composées connues.
- II.B.3.b Ne pas oublier qu'on se restreint ici au cas de l'oscillateur harmonique.
- II.B.3.c Un schéma satisfait la conservation de l'énergie si... l'énergie se conserve.
- II.B.3.d et e Là encore, il faut s'inspirer du portrait de phase de l'oscillateur harmonique vu dans le cours de physique.
- II.C.2.a C'est la question la plus calculatoire du problème. Commencer par exprimer y_{i+1} et z_{i+1} en fonction uniquement de y_i et z_i sans oublier qu'on s'intéresse toujours uniquement à l'oscillateur harmonique. Par la suite, utiliser l'identité remarquable $a^2 - b^2 = (a - b)(a + b)$ et faire un développement en $O(h^3)$ pour ne pas s'encombrer de termes inutiles.

Partie III

- III.A.2 Définir une fonction auxiliaire **norme(vecteur)** pour simplifier l'implémentation.
- III.B.2 Utiliser la méthode de Verlet pour chaque coordonnée en position de chaque corps après avoir fait apparaître l'équation différentielle par application de la relation fondamentale de la dynamique. Attention, la fonction f dépend de la position de tous les corps et pas seulement de la coordonnée considérée.
- III.B.3 Il reste les vitesses à faire évoluer. Utiliser à nouveau la méthode de Verlet sur chaque composante de chaque vitesse à l'aide des positions présentes et à venir. Noter qu'il n'est pas nécessaire de calculer les positions suivantes pour chaque étape de la boucle sur j . Un seul calcul en dehors de la boucle est suffisant.
- III.B.5.a Il est possible de trouver une complexité cubique dans le cas où vous n'avez pas tenu compte de l'indication précédente.

Partie IV

- IV.B.1 Le mot-clé **DISTINCT** utilisé en conjonction avec **COUNT** permet d'éviter les doublons.
- IV.B.2 Pour un corps donné, la date du dernier état antérieur à **tmin()** est le **MAX** de toutes les dates des états antérieurs à **tmin()**.
- IV.B.3 Il y a des informations à prendre sur trois tables donc deux jointures à effectuer. La distance attendue pour l'ordonnancement est euclidienne.

I. QUELQUES FONCTIONS UTILITAIRES

I.A.1 Le `+` appliqué sur deux listes a pour action de les concaténer l'une à l'autre en créant une troisième liste. Il ne faut donc pas confondre avec l'addition terme à terme comme on peut l'imaginer lors de la sommation de deux vecteurs.

```
>>> [1,2,3] + [4,5,6]
[1, 2, 3, 4, 5, 6]
```

L'idée de cette partie est de montrer que les listes Python ne se comportent pas naturellement comme des vecteurs mathématiques, c'est-à-dire que la « somme » de deux listes ne donne pas une somme composante par composante des deux listes (ici on attendrait `[5, 7, 9]`) mais les concatène. Il existe bien sûr des objets Python qui permettent de faire exactement cela, ce sont les `numpy.array` qui se comportent « comme attendu » vis-à-vis de l'addition :

```
>>> import numpy
>>> a = numpy.array([1,2,3])
>>> b = numpy.array([4,5,6])
>>> a+b
array([5, 7, 9])
```

Attention, il ne faut *jamaïs* utiliser la concaténation pour construire une liste élément par élément car celle-ci, en Python, renvoie une copie des deux listes. Par conséquent, la construction suivante d'une liste contenant les n premiers entiers pairs est de complexité *quadratique* en n .

```
L= []
for i in range(n):
    L= L + [2*i] # Ne JAMAIS faire cela, préférer L.append(2*i)
```

En effet, chaque concaténation est linéaire en i , le nombre d'éléments de la liste L à l'étape i , de sorte que le nombre total d'opérations est de l'ordre de $\sum_{i=1}^N i \approx \frac{N^2}{2}$. Alors certes, cela ne sera pas trop handicapant quand on fabrique une liste d'une dizaine d'éléments, mais déjà pour mille, cela se sentira et ce sera encore bien pire pour un million !

I.A.2 Pour les entiers naturels, $n \times x$ vaut $x + \dots + x$ où x apparaît n fois. Suivant la même sémantique, la multiplication d'un entier avec une liste revient à concaténer la liste avec elle-même autant de fois que demandé.

```
>>> 2*[1,2,3]
[1, 2, 3, 1, 2, 3]
```

Le résultat « naturel » attendu (soit `[2, 4, 6]`) serait donné par l'usage d'un tableau Numpy (essayez `2*numpy.array([1,2,3])`).

Ni la concaténation de deux listes, ni la multiplication d'une liste par un entier ne sont à proprement parler au programme d'informatique pour tous, mais il est probable que vous croisieiez cette syntaxe pour définir une liste initialisée à zéro sous la forme `L = [0]*n`.

Attention, il ne faut pas utiliser cette construction pour initialiser une matrice de zéros car Python ne fait pas une « copie profonde » des objets concernés. Observez plutôt :

```
>>> a = [[0]*3]*2      # Création de la "matrice"
>>> a[0][1] = 1        # Modification de la première ligne
>>> a                  # Hein ??
[[0, 1, 0], [0, 1, 0]]
```

I.B Présentons trois versions possibles, l'une en créant une liste remplie de zéros, l'autre en itérant sur les éléments et en construisant la liste au fur et à mesure par ajouts successifs, la dernière en utilisant la Pythonnerie de construction de liste en compréhension.

```
def smul(nombre,liste):
    """ Multiplication terme à terme d'une liste par un nombre."""
    L = [0]*len(liste)      # Initialisation à une liste de 0
    for i in range(len(liste)): # Autant de fois que d'éléments
        L[i] = nombre * liste[i] # On remplit avec la valeur idoine
    return L                # On n'oublie pas de renvoyer L

def smul(nombre,liste):
    L = []                  # Initialisation à une liste vide
    for element in liste:   # On itère sur les éléments
        L.append(nombre*element) # On rajoute la valeur idoine
    return L                # On n'oublie pas de renvoyer L

def smul(nombre,liste):
    return [nombre*element for element in liste] # One-liner !
```

Bien sûr, le jour du concours, une seule version est nécessaire (et suffisante!), mais il est intéressant de comparer les diverses manières permettant d'arriver au résultat. La version la plus proche de ce que tout le monde doit pouvoir faire au vu du programme officiel est la seconde. Néanmoins, les correcteurs comprendront les différents dialectes. Le rapport de concours précise même que « De nombreux candidats résolvent cette partie à l'aide de liste en compréhension, qui produisent du code concis et lisible. »

I.C.1 L'addition terme à terme sur les listes se définit par

```
def vsom(L1,L2):
    """ Fait l'addition vectorielle L1+L2 de deux listes.
    Les deux listes doivent avoir la même taille. """
    L = [0] * len(L1)      # Inialisation à une liste de 0
    for i in range(len(L1)): # On regarde toutes les positions
        L[i] = L1[i] + L2[i] # Addition à la position i
    return L                # Renvoi du résultat
```

À noter qu'ici on ne peut pas itérer sur les éléments car on a besoin de ceux de chaque liste, ce qui impose de passer par les indices des positions, communes aux deux listes.

Mines Maths 1 PC 2015 — Corrigé

Ce corrigé est proposé par Émilie Liboz (Professeur en CPGE) ; il a été relu par Florence Monna (Docteur en mathématiques) et Benjamin Monmege (Enseignant-chercheur à l'université).

Le sujet introduit une distance entre deux lois de probabilités à valeurs dans $\mathbb{N}$, connue sous le nom de « distance en variation totale » dans les livres. L'objectif du problème est de majorer la distance entre une loi de Poisson et la loi d'une somme S de variables aléatoires indépendantes de Bernoulli. On suit pour cela la méthode dite de Stein, qui permet par ailleurs de généraliser l'interprétation de la loi de Poisson comme une loi des événements rares.

- La première partie établit des résultats préliminaires sur des séries convergentes et calcule la distance entre deux lois de Bernoulli.
- La deuxième partie caractérise l'ensemble des suites de réels $(p_n^{(\lambda)})_{n \in \mathbb{N}}$ définissant une loi de Poisson $\mathcal{P}(\lambda)$.
- La troisième partie est consacrée à la résolution de l'équation de Stein

$$\forall n \in \mathbb{N} \quad \lambda f(n+1) - n f(n) = \tilde{h}(n)$$

où le second membre fait intervenir les $p_n^{(\lambda)}$. On démontre que les suites f solutions sont bornées.

- La quatrième partie démontre la propriété de Lipschitz pour les fonctions solutions f précédentes : elle donne un majorant de l'ensemble des variations $|f(n+1) - f(n)|$ de f pour $n \in \mathbb{N}$. Les notions de borne inférieure et de borne supérieure interviennent à plusieurs reprises.
- Enfin, la cinquième partie applique les résultats précédents pour majorer la distance entre la loi de Poisson $\mathcal{P}(\lambda)$ et la loi de S . Elle fait appel à des outils de base sur les variables aléatoires finies : espérance, indépendance, théorème de transfert.

La majorité des questions portent sur les séries numériques : étude de la convergence, calculs sur les sommes de séries convergentes. Les probabilités n'apparaissent que dans la dernière partie. L'ensemble est très calculatoire et assez long. À la difficulté technique s'ajoute la forte dépendance des questions et des parties entre elles. Ce sujet permet à la fois de faire le point sur les séries numériques et de s'entraîner à suivre le fil d'un énoncé, c'est-à-dire à comprendre où il nous emmène et par quel chemin.

INDICATIONS

Partie I

- 3 Majorer f afin de comparer le terme général de cette série à celui d'une série convergente.

Partie II

- 6 Considérer $f = \mathbf{1}_{\{n_0\}}$ pour $n_0 \in \mathbb{N}$.

Partie III

- 7 Raisonner par analyse-synthèse. Utiliser la formule (3) de l'énoncé pour construire une infinité d'éléments de $\mathcal{S}_h$, en remarquant que la condition ne concerne que les $n > 0$.

- 8 Calculer $\sum_{k=0}^{+\infty} \widetilde{h}(k) \frac{\lambda^k}{k!}$.

- 9 Utiliser la formule de la question 8 en majorant $\widetilde{h}$.

Partie IV

- 10 Utiliser la formule de la question 7.
 11 Exploiter le résultat de la question 8.
 12 Utiliser les formules des questions 10 et 11.
 13 Attention, la formule donnée pour $\Delta f_0(0)$ est fausse.
 14 Distinguer les cas $n \neq m$ et $n = m$ pour utiliser les questions 12 et 13.
 15 Comparer $\widetilde{h}$ et $\widetilde{h}_+$ (les notations utilisées sont celles définies dans l'énoncé au début de la partie III).
 16 Majorer h_+ afin de comparer le terme général de cette série à celui d'une série convergente en utilisant la formule de la question 10.
 17 Utiliser le fait que $f_m \in \mathcal{S}_{\mathbf{1}_{\{m\}}}$.
 18 Remarquer que, ici, f est une fonction quelconque de $\mathcal{S}_h$ et chercher à la relier à la fonction de la question 17.

Partie V

- 19 Distinguer les cas $X_k = 0$ et $X_k = 1$ pour la première égalité. Utiliser l'indépendance des X_k pour la deuxième.
 20 Remplacer λ par $\sum_{k=1}^n r_k$ et S par $\sum_{k=1}^n X_k$ dans le membre de gauche et utiliser les formules obtenues à la question 19.
 21 Appliquer le théorème du transfert et utiliser le fait que $f_A \in \mathcal{S}_{\mathbf{1}_A}$ pour calculer $E(\lambda f_A(S+1) - S f_A(S))$ puis prendre la borne supérieure sur $A \subset \mathbb{N}$.
 22 Appliquer les résultats des questions 20 et 21 pour exprimer $\text{dist}(\text{loi}(S), P_\lambda)$ en fonction de f_A puis utiliser l'inégalité (5).

I. PRÉLIMINAIRES

1 La série numérique $\sum \frac{\lambda^n}{n!}$ est convergente de somme e^λ . La suite $\left(e^{-\lambda} \frac{\lambda^n}{n!}\right)_{n \geq 0}$ appartient donc à $\mathcal{P}$ car pour tout $n \geq 0$, $e^{-\lambda} \frac{\lambda^n}{n!} \geq 0$. Par conséquent,

Pour $c = e^{-\lambda}$, la suite $\left(c \frac{\lambda^n}{n!}\right)_{n \geq 0}$ appartient à $\mathcal{P}$.

2 Pour p et $q \in [0; 1]$, posons $p_0 = 1 - p$, $p_1 = p$, $q_0 = 1 - q$, $q_1 = q$ et, pour tout $n \geq 2$, $p_n = q_n = 0$. Alors, pour $A \subset \mathbb{N}$:

$$\sum_{n \in A} p_n - \sum_{n \in A} q_n = \begin{cases} 0 & \text{si } \{0, 1\} \subset A \\ q - p & \text{si } 0 \in A, 1 \notin A \\ p - q & \text{si } 1 \in A, 0 \notin A \\ 0 & \text{si } \{0, 1\} \subset A \end{cases}$$

Ainsi,
$$\sup_{A \subset \mathbb{N}} \left| \sum_{n \in A} p_n - \sum_{n \in A} q_n \right| = |p - q|$$

soit

$$\text{dist}\left((1 - p, p, 0, \dots), (1 - q, q, 0, \dots)\right) = |p - q|$$

3 Puisque $f \in \mathcal{F}$, il existe $M > 0$ tel que pour tout $n \in \mathbb{N}$, $|f(n)| \leq M$. De plus, si $P \in \mathcal{P}$, alors pour tout entier naturel n , on a $p_n \geq 0$ de sorte que $|f(n)p_n| \leq Mp_n$. Or Mp_n est le terme général d'une série convergente. Par comparaison de séries à termes positifs, on en déduit que la série numérique $\sum f(n)p_n$ est absolument convergente et par conséquent

La série numérique $\sum f(n)p_n$ est convergente.

II. CARACTÉRISATION

4 Soit $M > 0$ tel que pour tout $n \in \mathbb{N}$, $|f(n)| \leq M$. Ainsi, pour tout entier $n \geq 1$:

$$|nf(n)p_n^{(\lambda)}| \leq M e^{-\lambda} n \frac{\lambda^n}{n!} = M \lambda e^{-\lambda} \frac{\lambda^{n-1}}{(n-1)!}$$

Or le membre de droite est le terme d'une série convergente (on reconnaît le terme général de la série exponentielle). Par comparaison de séries à termes positifs, on en déduit que la série numérique $\sum nf(n)p_n^{(\lambda)}$ est absolument convergente, d'où

La série numérique $\sum nf(n)p_n^{(\lambda)}$ est convergente.

5 Tout d'abord, d'après la question 4, la série $\sum nf(n)p_n^{(\lambda)}$ est convergente. Ensuite, la suite P_λ est un élément de $\mathcal{P}$ d'après la question 1 et donc, la fonction f étant bornée, par le même raisonnement que celui de la question 3, on montre que la série

$\sum f(n+1)p_n^{(\lambda)}$ est convergente. Les sommes $\sum_{n=0}^{+\infty} f(n+1)p_n^{(\lambda)}$ et $\sum_{n=0}^{+\infty} nf(n)p_n^{(\lambda)}$ sont alors bien définies. De plus, à l'aide du changement d'indice $n = k+1$, on obtient

$$\sum_{n=0}^{+\infty} nf(n)p_n^{(\lambda)} = \sum_{n=1}^{+\infty} nf(n)e^{-\lambda} \frac{\lambda^n}{n!} = \sum_{k=0}^{+\infty} (k+1)f(k+1)e^{-\lambda} \frac{\lambda^{k+1}}{(k+1)!}$$

soit
$$\sum_{n=0}^{+\infty} nf(n)p_n^{(\lambda)} = \lambda \sum_{k=0}^{+\infty} f(k+1)e^{-\lambda} \frac{\lambda^k}{k!}$$

En réindexant la deuxième somme par n au lieu de k , on peut conclure que

$$\boxed{\lambda \sum_{n=0}^{+\infty} f(n+1)p_n^{(\lambda)} = \sum_{n=0}^{+\infty} nf(n)p_n^{(\lambda)}}$$

6 Soit $Q = (q_n)_{n \geq 0} \in \mathcal{P}$ tel que, pour tout $f \in \mathcal{F}$, on ait

$$\lambda \sum_{n=0}^{+\infty} f(n+1)q_n = \sum_{n=0}^{+\infty} nf(n)q_n$$

En sélectionnant $f = 1_{\{n_0\}}$ pour $n_0 \in \mathbb{N}^*$ (la fonction indicatrice définie dans l'introduction du sujet), l'identité vérifiée par Q devient :

$$\lambda \sum_{n=0}^{+\infty} 1_{\{n_0\}}(n+1)q_n = \sum_{n=0}^{+\infty} n 1_{\{n_0\}}(n)q_n$$

Or, $1_{\{n_0\}}(n+1) \neq 0$ si et seulement si $n = n_0 - 1$ donc cette égalité est équivalente à $\lambda q_{n_0-1} = n_0 q_{n_0}$. La suite $(q_n)_{n \geq 0}$ vérifie alors

$$\forall n \geq 0 \quad q_{n+1} = \frac{\lambda}{n+1} q_n$$

Montrons par récurrence sur n que la propriété

$$\mathcal{P}(n): \quad q_n = q_0 \frac{\lambda^n}{n!}$$

est vraie pour tout entier n .

- $\mathcal{P}(0)$ est évidemment vraie.
- $\mathcal{P}(n) \implies \mathcal{P}(n+1)$: si $q_n = q_0 \lambda^n / n!$ pour un n donné, alors

$$q_{n+1} = \frac{\lambda}{n+1} q_n = \frac{\lambda}{n+1} \frac{\lambda^n}{n!} q_0 = \frac{\lambda^{n+1}}{(n+1)!} q_0$$

ce qui montre que la propriété est vraie au rang $n+1$.

- Conclusion: $\forall n \in \mathbb{N} \quad q_n = q_0 \frac{\lambda^n}{n!}$

De plus, la suite Q est un élément de $\mathcal{P}$ donc, d'après la question 1, $q_0 = e^{-\lambda}$ ce qui implique que

$$\boxed{\text{Les suites } Q \text{ et } P_\lambda \text{ sont égales.}}$$

Mines Maths 2 PC 2015 — Corrigé

Ce corrigé est proposé par Juliette Brun Leloup (Professeur en CPGE) ; il a été relu par Kévin Destagnol (ENS Cachan) et Benoît Chevalier (ENS Ulm).

Le but du sujet est d'étudier les suites de Fibonacci $(F_n(\lambda))_{n \in \mathbb{Z}}$ et de Lucas $(L_n(\lambda))_{n \in \mathbb{Z}}$ généralisées où λ est un réel tel que $\tilde{\lambda} = \lambda^2 + \lambda - 1 \neq 0$. Elles vérifient sur $\mathbb{N}$ une même relation de récurrence linéaire d'ordre 2 donnée par

$$\forall n \geq 0 \quad u_{n+2} = (1 + 2\lambda) u_{n+1} - \tilde{\lambda} u_n$$

Elles diffèrent par leurs deux premiers termes dans $\mathbb{N}$ et sont définies dans $\mathbb{Z} \setminus \mathbb{N}$ par

$$\forall n \geq 1 \quad F_{-n}(\lambda) = -F_n(\lambda) \tilde{\lambda}^{-n} \quad \text{et} \quad L_{-n}(\lambda) = L_n(\lambda) \tilde{\lambda}^{-n}$$

On retrouve les suites de Fibonacci $(F_n)_{n \in \mathbb{Z}}$ et de Lucas $(L_n)_{n \in \mathbb{Z}}$ lorsque $\lambda = 0$.

- La première partie établit quelques résultats sur les matrices J réelles carrées de taille 2, non multiple de l'identité, et vérifiant $J^2 = (5/4)I$. On fixe une telle matrice J dans toute la suite.
- La deuxième partie a pour but de généraliser à $\mathbb{Z}$ la « formule de Moivre » admise par l'énoncé sur $\mathbb{N}$:

$$\forall n \in \mathbb{Z} \quad R(\lambda)^n = F_n(\lambda) J + \frac{1}{2} L_n(\lambda) I \quad \text{où} \quad R(\lambda) = J + \left(\lambda + \frac{1}{2} \right) I$$

Pour cela, on montre que $R(\lambda)$ est inversible et on calcule $R(\lambda)^{-1}$, puis on établit deux nouvelles relations de récurrence où sont imbriqués les termes des suites $(F_n(\lambda))_{n \geq 1}$ et $(L_n(\lambda))_{n \geq 1}$.

- Dans la troisième partie, on établit des identités remarquables sur les suites de Fibonacci et de Lucas généralisées. Dans un premier temps, des calculs matriciels fournissent les formules sommatoires

$$\forall n \in \mathbb{N} \quad \sum_{k=0}^n \tilde{\lambda}^k F_{n-2k}(\lambda) = 0 \quad \text{et} \quad \sum_{k=0}^n \tilde{\lambda}^k L_{n-2k}(\lambda) = 2 F_{n+1}(\lambda)$$

Dans un second temps, on fait le lien avec les suites de Fibonacci et de Lucas standard via les formules

$$\forall n \in \mathbb{Z} \quad \forall k \geq 1 \quad U_n \left(\frac{L_{k-1}}{L_k} \right) = \frac{5^{\lfloor \frac{n}{2} \rfloor}}{L_k^n} U_{nk} \quad \text{avec} \quad U = F \text{ ou } L$$

- Enfin, dans l'unique question de la dernière partie, on utilise les identités remarquables précédentes pour démontrer que la famille de réels

$$p_k = \frac{L_i L_{2i(n-k)}}{2 L_{i(2n+1)}} \quad \text{pour } k \in \{0, 1, \dots, 2n\}$$

définit une probabilité.

Ce sujet très calculatoire fait intervenir les suites mais aussi un peu d'algèbre linéaire et un soupçon de probabilités. Il fait essentiellement appel au programme de première année.

INDICATIONS

Partie I

- 2 Résoudre l'équation matricielle $J^2 = \frac{5}{4} I$ en posant $J = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ où a, b, c, d sont quatre inconnues réelles.
- 3 Revenir à la définition pour démontrer que la famille (I, J) est libre : considérer λ et μ deux réels tels que $\lambda J + \mu I = 0$ et montrer que $\lambda = \mu = 0$.

Partie II

- 4 Effectuer la division euclidienne de $\left(X + \lambda + \frac{1}{2}\right)^2$ par $X^2 - \frac{5}{4}$.
- 5 Ne pas tenir compte de l'indication de l'énoncé « en déduire » et faire le calcul directement.
- 6 Comme λ est un élément de $\mathcal{C}$, isoler la matrice I d'un côté de l'égalité à partir de la relation (8) et mettre en facteur $R(\lambda)$ de l'autre côté de l'égalité.
- 7 Calculer de deux façons différentes $R(\lambda)^{n+1}$ et conclure en utilisant l'indépendance de la famille (J, I) sur $\mathcal{M}_2(\mathbb{R})$.
- 8 Démontrer le résultat par récurrence en utilisant les formules obtenues aux questions 6 et 7 et les formules (3) et (4) données par l'énoncé.

Partie III

- 10 Utiliser le résultat de la question 9 en remarquant que J et $R(\lambda)$ commutent.
- 11 Faire apparaître $W(\lambda)$ dans la somme que l'on cherche à calculer puis montrer que $[I - W(\lambda)] \sum_{k=0}^n W(\lambda)^k = I - W(\lambda)^{n+1}$ et conclure en utilisant la question 10.
- 12 Utiliser le résultat établi à la question 11, la formule de Moivre et la liberté de la famille (J, I) pour obtenir les valeurs des deux sommes.
- 13 Suivre l'indication de l'énoncé en exprimant, à l'aide de la formule (2), $L_{k+1}(\lambda)$ et $L_{k+2}(\lambda)$ en fonction de $L_{k-1}(\lambda)$, $L_k(\lambda)$ et $L_{k+1}(\lambda)$ dans la deuxième colonne de $\Delta_{k+1}(\lambda)$. Conclure en utilisant les propriétés du déterminant pour faire apparaître $\Delta_k(\lambda)$.
- 14 Revenir à la formule de la définition du déterminant d'une matrice de $\mathcal{M}_2(\mathbb{R})$ pour faire apparaître la formule demandée dans le calcul de $\Delta_k(\lambda)$.
- 15 Transformer la partie droite de l'égalité en utilisant la formule de Moivre ainsi que les propriétés admises par l'énoncé. Obtenir alors la partie gauche de l'égalité.
- 16 Partir de l'égalité établie à la question 15 mise à la puissance $2n$. Utiliser la formule de Moivre pour transformer la partie gauche de l'égalité. Remarquer la commutativité des matrices J et $R(0)$ pour transformer celle de droite. Conclure avec la liberté de la famille (J, I) dans $\mathcal{M}_2(\mathbb{R})$.

Partie IV

- 17 Suivre l'indication de l'énoncé pour démontrer que $\sum_{k=0}^{2n} p_k = 1$. Écrire $L_{i(2n+1)}$ sous forme de somme à l'aide de la formule (14) puis du résultat de la question 12. Utiliser les propriétés démontrées à la question 14 et la relation (13) pour transformer les termes sommés et faire apparaître le résultat demandé.

I. PRÉLIMINAIRES

1 Soit $\lambda \in \mathcal{C}$. Appliquons les formules (1) et (2) données par l'énoncé à $n = 1$:

$$F_2(\lambda) = (1 + 2\lambda) F_1(\lambda) + (1 - \lambda - \lambda^2) F_0(\lambda)$$

$$L_2(\lambda) = (1 + 2\lambda) L_1(\lambda) + (1 - \lambda - \lambda^2) L_0(\lambda)$$

d'où $\forall \lambda \in \mathcal{C} \quad F_2(\lambda) = 1 + 2\lambda \quad \text{et} \quad L_2(\lambda) = 2\lambda^2 + 2\lambda + 3$

2 Posons $J = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ où a, b, c, d sont quatre inconnues réelles. On a

$$J^2 = \begin{pmatrix} a^2 + bc & ab + bd \\ ca + dc & cb + d^2 \end{pmatrix}$$

donc $J^2 = \frac{5}{4} I \iff \begin{cases} a^2 + bc = 5/4 \\ b(a + d) = 0 \\ c(a + d) = 0 \\ d^2 + cb = 5/4 \end{cases}$

Considérons le cas où $a + d = 0$. Le système est équivalent à

$$\begin{cases} a^2 + bc = \frac{5}{4} \\ d = -a \end{cases}$$

En particulier, si l'on choisit $a = d = 0$ la première équation devient $bc = 5/4$ et on trouve une infinité de solutions possibles.

Les matrices $J = \begin{pmatrix} 0 & b \\ \frac{5}{4b} & 0 \end{pmatrix}$ pour $b \in \mathbb{R}^*$ satisfont la condition c).

Il est inutile dans cette question de donner la forme générale de toutes les matrices vérifiant $J^2 = (5/4) I$. Il suffit d'en donner une infinité qui convient, ce qui est plus simple. Géométriquement, on peut penser à toutes les symétries par rapport à une droite dans le plan. Il y en a une infinité. Il suffit de rajouter un coefficient $\sqrt{5}/2$ en facteur pour avoir une matrice J .

3 Montrons que les matrices I et J sont linéairement indépendantes sur $\mathcal{M}_2(\mathbb{R})$. Soient λ et μ deux réels tels que

$$\lambda J + \mu I = 0$$

Si $\lambda \neq 0$, alors $J = -\frac{\mu}{\lambda} I$. Comme J n'est pas multiple de I d'après la condition c), nécessairement $\lambda = 0$. Il ne reste que $\mu I = 0$ puis $\mu = 0$:

Les matrices I et J sont linéairement indépendantes sur $\mathcal{M}_2(\mathbb{R})$.

II. FORMULE DE MOIVRE GÉNÉRALISÉE

4 Soit $\lambda \in \mathcal{C}$. Effectuons la division euclidienne du polynôme $(X + \lambda + 1/2)^2$ par le polynôme $X^2 - 5/4$. Le reste de la division doit ainsi être un polynôme de degré inférieur ou égal à 1. On a

$$\begin{aligned}\left(X + \lambda + \frac{1}{2}\right)^2 &= X^2 + (2\lambda + 1)X + \lambda^2 + \lambda + \frac{1}{4} \\ &= X^2 - \frac{5}{4} + (2\lambda + 1)X + \lambda^2 + \lambda + \frac{3}{2}\end{aligned}$$

En posant $Q = 1$ et $T = (2\lambda + 1)X + \lambda^2 + \lambda + \frac{3}{2}$, on a trouvé deux polynômes Q et T de $\mathbb{R}[X]$ tels que

$$\left(X + \lambda + \frac{1}{2}\right)^2 = \left(X^2 - \frac{5}{4}\right)Q(X) + T(X)$$

5 Soit $\lambda \in \mathcal{C}$. On a $R(\lambda) = J + (\lambda + 1/2)I$. Comme les matrices J et I commutent et comme $J^2 = 5/4I$,

$$\begin{aligned}R(\lambda)^2 &= J^2 + (2\lambda + 1)J + \left(\lambda^2 + \lambda + \frac{1}{4}\right)I \\ &= (2\lambda + 1)J + \left(\lambda^2 + \lambda + \frac{3}{2}\right)I \\ &= (1 + 2\lambda)R(\lambda) - (1 + 2\lambda)\left(\lambda + \frac{1}{2}\right)I + \left(\lambda^2 + \lambda + \frac{3}{2}\right)I\end{aligned}$$

$$R(\lambda)^2 = (1 + 2\lambda)R(\lambda) + (1 - \lambda - \lambda^2)I$$

L'énoncé voudrait en fait que l'on utilise ici les polynômes de matrices, mais ils sont hors programme en PC. Pour vous présenter ce que le concepteur du sujet attendait (mais pas les correcteurs!), voici d'abord un bref complément. Si $P = \sum_{k=0}^n a_k X^k$ est un polynôme de $\mathbb{R}[X]$ et M une matrice de $\mathcal{M}_2(\mathbb{R})$, alors on définit $P(M)$ par

$$P(M) = \sum_{k=0}^n a_k M^k$$

où les puissances de matrices sont définies par récurrence par

$$M^0 = I \text{ et pour tout entier } k \geq 0, M^{k+1} = M^k M$$

Pour répondre à la question 5, on applique à la matrice J la formule obtenue à la question 4 :

$$\left(J + \left(\lambda + \frac{1}{2}\right)I\right)^2 = \left(J^2 - \frac{5}{4}I\right) + (2\lambda + 1)J + \left(\lambda^2 + \lambda + \frac{3}{2}\right)I$$

Par définition de $R(\lambda)$ et d'après le choix de la matrice J , il s'ensuit que

$$\begin{aligned}R(\lambda)^2 &= (2\lambda + 1)J + \left(\lambda^2 + \lambda + \frac{3}{2}\right)I \\ &= (1 + 2\lambda)R(\lambda) + (1 - \lambda - \lambda^2)I\end{aligned}$$

en reprenant le même calcul que précédemment.

Mines Informatique PC 2015 — Corrigé

Ce corrigé est proposé par Julien Dumont (Professeur en CPGE) ; il a été relu par Guillaume Batog (Professeur en CPGE) et Jean-Julien Fleck (Professeur en CPGE).

Le sujet porte sur différents aspects de tests de validation d'une imprimante. Les parties sont indépendantes.

- La première partie porte sur la réception des données issues de la carte d'acquisition, en se concentrant plus particulièrement sur la trame contenant les informations.
- La deuxième, très courte et proche du cours, met en place des programmes de calcul approché d'intégrales, qui sont utilisés pour obtenir la moyenne et l'écart-type des données physiques reçues dans les trames.
- La troisième partie s'intéresse à la gestion d'un parc d'imprimantes en cours de validation à l'aide de bases de données.
- La quatrième aborde le thème de la compression de données, principalement en étudiant un long programme fourni en annexe.
- Enfin, la cinquième partie examine un schéma numérique de résolution d'équations différentielles afin de définir un test de validation des moteurs utilisés dans l'imprimante.

Ce sujet est intéressant sur le principe mais il est décevant : beaucoup de notions sont hors programme et ne sont pas définies ; le programme proposé dans la quatrième partie ne marche pas, ce qui conduit à des questions n'ayant pas vraiment de sens pour un candidat rigoureux ; les formulations des questions sont parfois très imprécises et on ne sait pas réellement ce qui est demandé... Bref, cet énoncé a tout pour désarçonner un candidat ayant sérieusement préparé le concours. Cependant, en vue des révisions, c'est un sujet qui balaie de nombreux domaines et, si on lit les indications du corrigé pour savoir ce qu'il est possible de faire ou pas, ce problème varié mérite que l'on s'y attarde.

INDICATIONS

La mention (HP) dans une indication signale une question hors programme, pour laquelle une explication détaillée est fournie dans le corrigé.

- Q1 (HP)
- Q3 Penser à regarder soigneusement les structures demandées en retour de fonction.
- Q4 Ne pas oublier le modulo.
- Q5 (HP)
- Q6 Cette question et la suivante se rapprochent des algorithmes de première année.
- Q8 (HP) Faire comme si l'on pouvait définir une constante en SQL comme `Imin` ou `Imax`.
- Q9 Utiliser des requêtes imbriquées.
- Q10 Écrire une requête portant sur une table dans laquelle la clause `WHERE` dépend d'une requête portant sur une autre table.
- Q11 (HP)
- Q12 On suppose que le programme proposé marche.
- Q14 (HP) On peut deviner la réponse en s'inspirant de l'énoncé, sans plus de connaissances sur les dictionnaires.
- Q16 (HP) Un invariant d'itération est l'équivalent d'un invariant de boucle pour les fonctions récursives.
- Q22 (HP)
- Q23 Il faut construire ici une fonction permettant d'évaluer si un moteur est défectueux ou non. Par exemple, on peut proposer un test comparant quelques solutions simulées et la sortie mesurée.

TESTS DE VALIDATION D'UNE IMPRIMANTE

Q1

Le complément à 2 est une méthode de représentation des entiers relatifs sur un nombre fini de bits. Un des bits est utilisé pour coder le signe, les suivants permettent de représenter la valeur absolue du nombre. La façon de représenter cette valeur absolue n'est pas nécessaire pour répondre à cette question. On donne ici une réponse générale indépendante de la convention, l'emploi de la terminologie « complément à 2 » imposant en fait l'intervalle demandé.

En complément à 2, puisqu'un bit sert à coder le signe, cela signifie qu'il en reste 9 pour coder la valeur absolue, soit $2^9 = 512$ valeurs possibles. Traditionnellement, **on considère que l'on code ici les entiers compris entre -512 et $+511$** . Mais on peut représenter de façon générale tout intervalle de nombres entiers contenant $2^{10} = 1024$ valeurs, si l'on décide arbitrairement d'une convention.

Le résultat de la conversion peut prendre toute plage de valeurs entières de 1024 valeurs.

Cependant, rien n'interdit de coder plusieurs fois une même valeur. Ainsi, une autre méthode de représentation de nombres qui consiste à coder le signe par le premier bit et la valeur absolue par les suivants en tant qu'entier naturel conduit à coder deux fois le zéro. Par exemple, sur 4 bits, les binaires 0000 et 1000 représentent « respectivement » $+0$ et -0 , codant deux fois la même chose. L'intérêt des différentes méthodes de représentations de nombres est précisément de pallier ce genre de défaut.

Q2

La résolution de la mesure se déduit de la possibilité de représenter 1024 éléments sur $N = 10$ bits. 2^N éléments permettent de définir $2^N - 1$ intervalles. La plage de tension P étant de 10 V, la résolution σ vaut

$$\sigma = \frac{P}{2^N - 1} = 1.10^{-2} \text{ V}$$

Q3

Le principe du programme est le suivant. On initialise une variable `nonstop` à la valeur booléenne `True`. On lit alors un à un les caractères reçus jusqu'à tomber sur un caractère d'en-tête. On change alors la valeur de `nonstop` pour sortir de la boucle. On lit par la suite le nombre N de données envoyées. On sait que la longueur totale de la trame est exactement $8 + 4N$: un caractère correspond à l'en-tête, trois au nombre de données envoyées, $4N$ permettent de détailler ces dernières et quatre caractères donnent le checksum. On utilise également à répétition la conversion du type chaîne de caractères vers le type entier grâce à `int`.

```
def lect_mesures():
    '''Fonction qui renvoie une trame à partir du premier en-tête'''
    nonstop=True
    resultat=[ ] #Liste que l'on renverra à la fin
    ###Recherche de l'en-tête
    while nonstop:
        test=com.read(1)
        if test in ['U','I','P']:#A-t-on trouvé l'en-tête ?
            resultat=[test]      #Si oui, on le stocke
            nonstop=False        #Et on fait en sorte de sortir du while
```

```

###Lecture du nombre de données à venir
N=int(com.read(3)) #On stocke le nombre de données à lire
###Lecture des données
donnees=[ ] #Liste vide contenant les données
for inc in range(N):
    #On lit 4 caractères que l'on convertit
    #en entier pour les stocker
    donnees.append(int(com.read(4)))
resultat.append(donnees) #On stocke donnees
##Ajout du checksum
resultat.append(int(com.read(4)))
return resultat

```

Q4 Notons une confusion de l'énoncé entre *mesure* et *mesures* qui peut déstabiliser à la lecture du sujet, surtout lors de l'emploi de la notation `mesures[]`, utilisée par exemple en Java... Ce qui est demandé est toutefois relativement compréhensible entre les lignes. Enfin, ne pas oublier le modulo 10000.

```

def check(mesure,Checksum):
    '''Vérifie la validité d'un checksum'''
    #On calcule la somme des données reçues
    somme=0
    for x in mesure:
        somme += abs(x)
    #On teste la validité du checksum
    return somme%10000==Checksum

```

Q5 On suppose ici que la bibliothèque `matplotlib.pyplot` a été importée sous l'alias `plt`.

```

def affichage(mesure):
    ###Création du vecteur des temps
    Temps=[ ]
    #On connaît exactement la plage nécessaire : on part de 0ms
    #et on va à 400ms par pas de 2ms. On met donc 401 pour atteindre
    #400 inclus mais sans dépasser cette valeur.
    for t in range(0,401,2):
        Temps.append(t)
    ###Création du vecteur des mesures
    mesures=[ ] #Initialisation des mesures
    for m in mesure: #On balaie les données fournies
        mesures+=[m*4e-3] #Conversion des données en intensités
    ###Représentation graphique proprement dite
    plt.plot(Temps,mesures)
    ###Compléments : axes et titre
    plt.xlabel('Temps (ms)')
    plt.ylabel('Intensité (A)')
    plt.title('Courant moteur')

```

Ce programme suivant est enrichi des commandes qui auraient permis d'obtenir tout le graphique proposé. Selon l'interface de développement, on peut être amené à ajouter la fonction `show` pour que le graphique créé s'affiche effectivement. On peut également le sauver à l'aide de la fonction `savefig`.

X/ENS Maths PC 2015 — Corrigé

Ce corrigé est proposé par François Lê (ENS Lyon) ; il a été relu par Tristan Poullaouec (Professeur en CPGE) et Benjamin Monmege (ENS Cachan).

Le sujet s'intéresse à des propriétés vérifiées par les valeurs propres de matrices symétriques réelles. Il est partagé en trois parties dépendantes (ce que ne signale pas l'énoncé) et de proportions inégales.

- Un des objectifs de la première partie est d'établir plusieurs formules faisant le lien entre les valeurs propres d'une matrice symétrique réelle M et des bornes supérieures ou inférieures de produits scalaires du type $\langle x, Mx \rangle$. Grâce à ces formules, on démontre également la continuité de l'application qui associe à une matrice symétrique réelle son spectre ordonné.
- La deuxième partie considère deux matrices symétriques réelles A et B , dont les valeurs propres sont notées $a_1 \geq \dots \geq a_n$ et $b_1 \geq \dots \geq b_n$ respectivement. Si $c_1 \geq \dots \geq c_n$ sont les valeurs propres de la somme $C = A + B$, le résultat final de la partie est que $c_1 + \dots + c_j \leq a_1 + \dots + a_j + b_1 + \dots + b_j$ pour tout $j \in \llbracket 1; n \rrbracket$. On démontre aussi au passage le théorème de Schur, qui énonce que $a_1 + \dots + a_j \geq a_{11} + \dots + a_{jj}$ pour tout entier $1 \leq j \leq n$, où les a_{ii} sont les coefficients diagonaux de A .
- Pour la troisième partie, on se place dans le cas $n = 2$. Quatre réels a_1, a_2, b_1, b_2 étant fixés, on cherche à déterminer les spectres possibles des sommes de deux matrices symétriques réelles données ayant pour spectres (a_1, a_2) et (b_1, b_2) . C'est un segment de droite dont on explicite les extrémités.

Ce sujet est fondé sur un thème d'algèbre linéaire, mais d'autres parties du programme y sont mobilisées. Ainsi, les questions 4d, 5a et 5b se rapportent à la topologie et au calcul différentiel, et les questions 9 et 10c nécessitent d'être au point sur la géométrie élémentaire du plan. En outre, presque l'intégralité des questions font appel à des manipulations sur des bornes inférieures ou supérieures. Il faut être extrêmement vigilant dans ces calculs et rédiger calmement les solutions.

INDICATIONS

- 1b Trouver une matrice diagonale M et un réel λ tels que $s^\perp(\lambda M) \neq \lambda s^\perp(M)$.
- 2a Utiliser le théorème spectral pour trouver P orthogonale telle que ${}^t P M P$ soit égale à la matrice diagonale $\text{diag}(m_1, \dots, m_n)$, puis considérer les colonnes de P .
- 2b Commencer par majorer $\langle x, Mx \rangle$ par une quantité indépendante de x , puis montrer que cette majoration est atteinte pour un vecteur x particulier.
- 2c Montrer séparément que les bornes inférieure et supérieure proposées sont égales à m_j . Pour cela, effectuer des calculs analogues à ceux de la question 2b.
- 3a Penser à la formule de Grassmann.
- 4a Remarquer que $\langle x, Mx \rangle = \langle x, (M - L)x \rangle + \langle x, Lx \rangle$. Montrer ensuite que pour tout x normé, $\langle x, (M - L)x \rangle \geq 0$, puis utiliser le résultat de la question 3b pour lier $\langle x, Lx \rangle$ et $\langle x, Mx \rangle$ à ℓ_j et m_j .
- 4c Appliquer la question 4b à la matrice $L - M$ puis utiliser la question 4a pour montrer que $\ell_j \leq \|L - M\| + m_j$.
- 5a D'après la question 4c, on a $|\ell_j - m_j| < r$ dès que $\|L - M\| < r$. Sachant que les m_i sont tous distincts, utiliser cette remarque pour définir un $r > 0$ tel que, si $\|L - M\| < r$, alors les ℓ_j sont aussi tous distincts.
- 5b Pour clarifier la situation, identifier $S_2(\mathbb{R})$ à $\mathbb{R}^3$ et montrer que $\mathcal{S}_2^\dagger(\mathbb{R})$ s'identifie alors à $\{(\lambda, \mu, h) \in \mathbb{R}^3 \mid \lambda \neq \mu \text{ ou } h \neq 0\}$.
- 6b Utiliser la question 2b et l'égalité $\langle x, Cx \rangle = \langle x, Ax \rangle + \langle x, Bx \rangle$.
- 6c Se servir de la question 2c ou appliquer la question 6b en remplaçant A et B par $-A$ et $-B$.
- 7a Appliquer d'abord la formule de Grassmann aux sous-espaces $\mathcal{U} \cap \mathcal{V}$ et $\mathcal{W}$.
- 7b Pour la première partie de la question, trouver, grâce à la question 2c, des sous-espaces $\mathcal{U}, \mathcal{V}, \mathcal{W}$ associés respectivement à a_j, b_k et c_{j+k-1} , et partir ensuite de l'égalité $\langle x, Cx \rangle = \langle x, Ax \rangle + \langle x, Bx \rangle$. Pour la seconde partie, s'appuyer sur le résultat de la première en remplaçant A et B par $-A$ et $-B$.
- 8a Se souvenir que $\langle x, Ax \rangle \leq a_1$ pour tout vecteur x normé.
- 8b Développer les facteurs du membre de droite de l'inégalité proposée et faire apparaître $f(t_1, \dots, t_k)$.
- 8c Grâce à une écriture $A = P D {}^t P$ avec P orthogonale, montrer qu'il existe un n -uplet $(t_1, \dots, t_n) \in \mathcal{D}_{j,n}$ tel que $\sum_{i=1}^j a_{ii} = f(t_1, \dots, t_k)$. Conclure à l'aide de la question 8b.
- 8d Étant donnée une famille $(x_1, \dots, x_j) \in \mathcal{R}_j$, la compléter en une base orthonormée de $\mathbb{R}^n$. Considérer alors la matrice de passage de la base canonique à cette base pour trouver une relation de la forme $\langle x_i, Ax_i \rangle = \langle e_i, A' e_i \rangle$.
- 9 Dans le cas où A et B sont diagonales, chercher les possibilités pour $s^\perp(A + B)$ et mettre ainsi en évidence deux points distants de $\sqrt{2}(b_1 - b_2)$. Dans le cas général, utiliser les relations données par les questions 6a, 6b, 6c et 7b.
- 10b Montrer qu'une matrice $B \in S(b_1, b_2)$ peut être diagonalisée avec une matrice de passage orthogonale directe et paramétrer alors de telles matrices par $t \in [-\pi, \pi]$.
- 10c Écrire Σ comme l'image directe de $[-\pi; \pi]$ par une application continue. En paramétrant le segment L par $[0; 1]$, trouver ensuite une application continue $[-\pi; \pi] \rightarrow [0; 1]$. Pour conclure, utiliser le fait que l'image d'un intervalle de $\mathbb{R}$ par une application continue est un intervalle de $\mathbb{R}$.

I. PREMIÈRE PARTIE

1a L'ensemble $\mathcal{S}_n(\mathbb{R})$ est un sous-espace vectoriel de $\mathcal{M}_n(\mathbb{R})$. En effet,

- $\mathcal{S}_n(\mathbb{R}) \neq \emptyset$ car la matrice nulle est symétrique.
- Soient $\lambda \in \mathbb{R}$ et $M, N \in \mathcal{S}_n(\mathbb{R})$. Alors ${}^t(\lambda M + N) = \lambda {}^tM + {}^tN = \lambda M + N$, ce qui montre que $\lambda M + N \in \mathcal{S}_n(\mathbb{R})$.

En tant que sous-espace vectoriel de $\mathcal{M}_n(\mathbb{R})$, l'ensemble $\mathcal{S}_n(\mathbb{R})$ hérite d'une structure d'espace vectoriel (sur $\mathbb{R}$). Ainsi,

$\mathcal{S}_n(\mathbb{R})$ est un espace vectoriel.

Dans la plupart des sujets de concours, pour montrer qu'un ensemble est un espace vectoriel, on montre qu'il s'agit d'un sous-espace vectoriel d'un espace vectoriel connu.

On sait par ailleurs, grâce au cours, que la dimension de $\mathcal{S}_n(\mathbb{R})$ est égale à

$$\dim \mathcal{S}_n(\mathbb{R}) = \frac{n(n+1)}{2}$$

Enfin, d'après le théorème spectral, toute matrice symétrique réelle M (de taille n) est diagonalisable donc possède en particulier n valeurs propres (éventuellement confondues). Il est ainsi possible de parler du n -uplet $s^\downarrow(M)$ de ses valeurs propres ordonnées par ordre décroissant.

L'application $s^\downarrow$ est bien définie sur $\mathcal{S}_n(\mathbb{R})$.

Le théorème spectral énonce que toute matrice symétrique réelle est diagonalisable et que la diagonalisabilité peut s'effectuer à l'aide d'une matrice de passage orthogonale. Dans cette première question, il suffit d'invoquer la diagonalisabilité, sans parler de matrices orthogonales.

1b Considérons la matrice diagonale $M \in \mathcal{S}_n(\mathbb{R})$ dont les éléments diagonaux sont $1, 2, \dots, n$. Puisque les valeurs propres de M sont ses éléments diagonaux, on en déduit que $s^\downarrow(M) = (n, n-1, \dots, 2, 1)$. Par ailleurs, la matrice $-M$ est également diagonale. Comme ses éléments diagonaux sont $-1, -2, \dots, -n$, son spectre ordonné est égal à $s^\downarrow(-M) = (-1, -2, \dots, -n)$. Cela montre que $s^\downarrow(-M) \neq -s^\downarrow(M)$ donc que

L'application $s^\downarrow$ n'est pas linéaire.

Il est aussi possible de trouver deux matrices symétriques A et B telles que $s^\downarrow(A+B) \neq s^\downarrow(A) + s^\downarrow(B)$. On vérifie par exemple que si

$$A = \begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix} \quad \text{et} \quad B = \begin{pmatrix} 0 & -1 \\ -1 & 0 \end{pmatrix}$$

alors $s^\downarrow(A) = \left(\frac{1-\sqrt{5}}{2}, \frac{1+\sqrt{5}}{2} \right)$ et $s^\downarrow(B) = (-1, 1)$

d'une part, tandis que $s^\downarrow(A+B) = (0, 1) \neq s^\downarrow(A) + s^\downarrow(B)$ d'autre part.

1c Soit $M \in \mathcal{S}_n(\mathbb{R})$. Comme rappelé dans la question 1a, M est diagonalisable : il existe une matrice $P \in GL_n(\mathbb{R})$ telle que $P^{-1}MP$ soit une matrice diagonale dont les éléments diagonaux sont les valeurs propres $m_1, \dots, m_n$ de M . De plus, il est toujours possible de supposer que P vérifie

$$P^{-1}MP = \begin{pmatrix} m_1 & 0 & \cdots & 0 \\ 0 & m_2 & & 0 \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \cdots & m_n \end{pmatrix}$$

On a alors

$$P^{-1}(-M)P = \begin{pmatrix} -m_1 & 0 & \cdots & 0 \\ 0 & -m_2 & & 0 \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \cdots & -m_n \end{pmatrix}$$

ce qui prouve que le spectre de $-M$ est $\{-m_1, -m_2, \dots, -m_n\}$. Compte tenu des relations $m_1 \geq m_2 \geq \dots \geq m_n$, on en déduit que $-m_1 \leq -m_2 \leq \dots \leq -m_n$. Ainsi,

$$s^\downarrow(-M) = (-m_n, -m_{n-1}, \dots, -m_1)$$

1d Soit $M = \begin{pmatrix} \lambda & h \\ h & \mu \end{pmatrix}$ la matrice proposée. Son polynôme caractéristique vaut

$$\begin{aligned} \chi_M(X) &= X^2 - \text{Tr}(M)X + \det(M) \\ &= X^2 - (\lambda + \mu)X + \lambda\mu - h^2 \end{aligned}$$

Le discriminant de ce polynôme caractéristique est

$$\begin{aligned} \Delta &= (\lambda + \mu)^2 - 4(\lambda\mu - h^2) \\ &= \lambda^2 + 2\lambda\mu + \mu^2 - 4\lambda\mu + 4h^2 \\ &= \lambda^2 - 2\lambda\mu + \mu^2 + 4h^2 \\ \Delta &= (\lambda - \mu)^2 + 4h^2 \end{aligned}$$

Cette dernière expression montre bien que Δ est positif donc que M possède deux valeurs propres (éventuellement confondues) : cela est cohérent avec le théorème spectral.

Les valeurs propres de M sont les racines de χ_M , c'est-à-dire

$$\xi = \frac{\lambda + \mu + \sqrt{(\lambda - \mu)^2 + 4h^2}}{2} \quad \text{et} \quad \xi' = \frac{\lambda + \mu - \sqrt{(\lambda - \mu)^2 + 4h^2}}{2}$$

Puisque $\xi \geq \xi'$, on en déduit que

$$s^\downarrow(M) = \left(\frac{\lambda + \mu + \sqrt{(\lambda - \mu)^2 + 4h^2}}{2}, \frac{\lambda + \mu - \sqrt{(\lambda - \mu)^2 + 4h^2}}{2} \right)$$

2a D'après le théorème spectral, il existe une matrice orthogonale P telle que

$$M = P \begin{pmatrix} m_1 & 0 & \cdots & 0 \\ 0 & m_2 & & 0 \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \cdots & m_n \end{pmatrix} {}^tP$$

X/ENS Informatique B (MP/PC) 2015 — Corrigé

Ce corrigé est proposé par Benjamin Monmege (Enseignant-chercheur à l'université); il a été relu par Céline Chevalier (Enseignant-chercheur à l'université) et Guillaume Batog (Professeur en CPGE).

Ce sujet propose l'étude de deux algorithmes calculant l'enveloppe convexe d'un nuage de points en position générale dans le plan affine, c'est-à-dire qui ne contient pas trois points distincts alignés. Le calcul d'enveloppe convexe est utilisé dans de nombreux domaines : robotique, traitement d'image, théorie des jeux, vérification formelle... Les parties II et III sont indépendantes.

- La première partie propose d'implémenter deux fonctions préliminaires permettant respectivement de trouver le point le plus bas du nuage et de tester l'orientation d'un triangle, opération cruciale dans la suite.
- La deuxième partie étudie l'algorithme du paquet cadeau, qui consiste à envelopper petit à petit le nuage de points. Cette construction nécessite un temps d'exécution en $O(nm)$, où n est le nombre de points et m celui de leur enveloppe convexe.
- Enfin, la troisième partie étudie l'algorithme de balayage qui résout le même problème en temps $O(n \log n)$ grâce à la construction, à l'aide de piles, des enveloppes supérieure et inférieure des n points.

Ce sujet est d'une longueur raisonnable pour une épreuve de deux heures. Il se concentre sur les parties d'algorithmique et de programmation du programme et ne présente aucune difficulté majeure sous réserve de maîtriser les rudiments du langage de programmation choisi.

INDICATIONS

Partie II

- 4 Dans la propriété d'antisymétrie, il s'agit en fait de montrer que si $p_j \preceq p_k$ et $p_k \preceq p_j$ alors $p_j = p_k$: la conclusion $j = k$ vient alors du fait que le nuage P est un ensemble qui ne contient donc pas deux occurrences du même point. Pour prouver la propriété de transitivité, utiliser le fait que p_i appartient à l'enveloppe convexe, de sorte que tous les autres points sont inclus dans un demi-plan contenant p_i .
- 7 Utiliser les fonctions des questions 1 et 5. Pour ajouter un élément j à une liste `liste`, utiliser la commande `liste.append(j)`.

Partie III

- 10 Commencer par écrire une fonction auxiliaire renvoyant la paire (j, k) des dernier et avant-dernier éléments de la pile, en supprimant au passage le sommet de la pile. Utiliser alors une boucle `while` qui applique la fonction auxiliaire précédente et teste l'orientation du triangle (p_i, p_j, p_k) , tant que celle-ci est négative.
- 11 Se convaincre que seul le test d'orientation doit être modifié par rapport à la question précédente.
- 12 Une fois obtenues les piles `es` et `ei` contenant les enveloppes supérieure et inférieure respectivement, il s'agit d'insérer à l'envers le contenu de `es` duquel on a retiré les premier et dernier éléments à la pile `ei` (pour éviter les doublons).
- 13 Borner le nombre de fois où chaque indice peut être inséré ou supprimé définitivement des deux piles `es` et `ei`, puis majorer le nombre total de tests d'orientation effectués le long de l'exécution de `convGraham`.

I. PRÉLIMINAIRES

1 Initialisons un indice j à 0, ayant vertu à contenir l'indice du point de plus petite ordonnée. À l'aide d'une boucle `for`, testons chaque point afin de mettre à jour j si nécessaire, à savoir si le point courant a une ordonnée strictement inférieure, ou bien une même ordonnée mais une abscisse strictement inférieure.

```
def plusBas(tab,n):
    j = 0
    for i in range(1,n):
        if (tab[1][i] < tab[1][j] or
            (tab[1][i] == tab[1][j] and tab[0][i] < tab[0][j])):
            j = i
    return j
```

2 Pour $i = 0$, $j = 3$, $k = 4$, le tableau de l'énoncé apporte les coordonnées suivantes : le point p_0 a pour coordonnées $(0, 0)$, le point p_3 a pour coordonnées $(4, 1)$ et le point p_4 a pour coordonnées $(4, 4)$. Ainsi, les vecteurs $\overrightarrow{p_0p_3}$ et $\overrightarrow{p_0p_4}$ ont pour coordonnées respectives $(4, 1)$ et $(4, 4)$. L'aire signée du triangle (p_0, p_3, p_4) est donc

$$\frac{1}{2} \times \begin{vmatrix} 4 & 4 \\ 1 & 4 \end{vmatrix} = \frac{16 - 4}{2} = 6 > 0$$

de sorte que le triangle (p_0, p_3, p_4) est orienté positivement.

Le résultat du test d'orientation sur (p_0, p_3, p_4) est $+1$.

De même, pour $i = 8$, $j = 9$, $k = 10$, le point p_8 a pour coordonnées $(7, 2)$, le point p_9 a pour coordonnées $(8, 5)$ et le point p_{10} a pour coordonnées $(11, 6)$. Par conséquent, les vecteurs $\overrightarrow{p_8p_9}$ et $\overrightarrow{p_8p_{10}}$ ont pour coordonnées respectives $(1, 3)$ et $(4, 4)$. L'aire signée du triangle (p_8, p_9, p_{10}) est ainsi

$$\frac{1}{2} \times \begin{vmatrix} 1 & 4 \\ 3 & 4 \end{vmatrix} = \frac{4 - 12}{2} = -4 < 0$$

ce qui implique que le triangle (p_8, p_9, p_{10}) est orienté négativement.

Le résultat du test d'orientation sur (p_8, p_9, p_{10}) est -1 .

3 La fonction `orient` calcule le déterminant utilisé dans l'aire signée du triangle (p_i, p_j, p_k) et teste son signe. En particulier, le déterminant est nul si et seulement si deux des trois points au moins sont confondus, auquel cas le résultat du test d'orientation est 0.

```
def orient(tab,i,j,k):
    pi_pj = [tab[0][j]-tab[0][i], tab[1][j]-tab[1][i]]
    pi_pk = [tab[0][k]-tab[0][i], tab[1][k]-tab[1][i]]
    det = pi_pj[0] * pi_pk[1] - pi_pj[1] * pi_pk[0]
    if det > 0:
        return 1
    elif det < 0:
        return -1
    else:
        return 0
```

II. ALGORITHME DU PAQUET CADEAU

4 Pour la réflexivité, notons que pour tout $j \neq i$, $\text{orient}(\text{tab}, i, j, j) = 0$ puisque deux des points sont égaux, de sorte que $p_j \preceq p_j$.

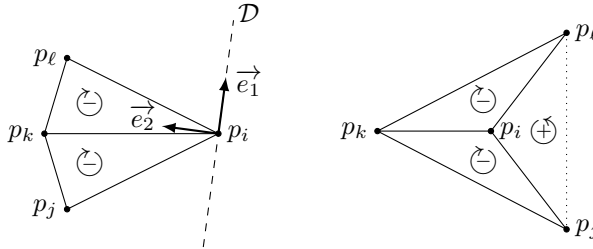
Pour l'antisymétrie, supposons fixés j et k tels que $p_j, p_k \in P \setminus \{p_i\}$, $p_j \preceq p_k$ et $p_k \preceq p_j$. Par définition,

$$\det(\overrightarrow{p_i p_j}, \overrightarrow{p_i p_k}) \leq 0 \quad \text{et} \quad \det(\overrightarrow{p_i p_k}, \overrightarrow{p_i p_j}) \leq 0$$

Puisque ces deux déterminants sont opposés l'un de l'autre, ils sont nuls, ce qui prouve que les vecteurs $\overrightarrow{p_i p_j}$ et $\overrightarrow{p_i p_k}$ sont colinéaires, c'est-à-dire que p_i , p_j et p_k sont alignés. Grâce à l'hypothèse de position générale et le fait que p_j et p_k sont distincts de p_i , cela montre que $p_j = p_k$.

Une coquille s'est glissée dans l'énoncé : la propriété d'antisymétrie consiste à montrer a priori que $p_j \preceq p_k$ et $p_k \preceq p_j$ implique $p_j = p_k$, et non $j = k$. On peut cependant conclure que $j = k$ puisque le nuage P est un ensemble qui ne contient donc pas deux occurrences du même point.

Montrons la transitivité $\preceq$. Pour cela, fixons j, k, ℓ tels que $p_j, p_k, p_\ell \in P \setminus \{p_i\}$, $p_j \preceq p_k$ et $p_k \preceq p_\ell$. Supposons également que le point p_i appartient à l'enveloppe convexe de P , de sorte que les points p_j , p_k et p_ℓ sont tous dans un demi-plan dont la frontière $\mathcal{D}$ contient p_i (et aucun des trois autres points grâce à l'hypothèse de position générale), comme représenté dans la figure de gauche ci-dessous.



La figure de droite ci-dessus montre que la propriété de transitivité est fautive si p_i n'est pas sur l'enveloppe convexe de P , puisque (p_i, p_j, p_k) et (p_i, p_k, p_ℓ) sont orientés négativement, mais que (p_i, p_j, p_ℓ) est orienté positivement.

Munissons le plan affine d'un repère orthonormé direct $(p_i, \vec{e}_1, \vec{e}_2)$ avec $\vec{e}_1$ un vecteur directeur de la droite $\mathcal{D}$ et $\vec{e}_2$ orienté vers le demi-plan contenant les points p_j , p_k et p_ℓ . Dans le plan complexe associé, les arguments principaux respectifs θ_j , θ_k et θ_ℓ des points p_j , p_k et p_ℓ appartiennent à $[0; \pi]$. L'interprétation géométrique du déterminant de deux vecteurs dans le plan euclidien assure que

$$\|\overrightarrow{p_i p_j}\| \times \|\overrightarrow{p_i p_k}\| \times \sin(\theta_k - \theta_j) \leq 0 \quad \text{et} \quad \|\overrightarrow{p_i p_k}\| \times \|\overrightarrow{p_i p_\ell}\| \times \sin(\theta_\ell - \theta_k) \leq 0$$

Puisque $\theta_k - \theta_j$ et $\theta_\ell - \theta_k$ appartiennent à $[-\pi; \pi]$, le fait que le sinus de ces angles soit négatif implique que $\theta_k - \theta_j$ et $\theta_\ell - \theta_k$ appartiennent en fait à $[-\pi; 0]$. Par somme, $\theta_\ell - \theta_j \in [-2\pi; 0]$. Comme cet angle est par ailleurs inclus dans $[-\pi; \pi]$, on en déduit que $\theta_\ell - \theta_j \in [-\pi; 0]$, de sorte que $\sin(\theta_\ell - \theta_j) \leq 0$. Finalement,

$$\det(\overrightarrow{p_i p_j}, \overrightarrow{p_i p_\ell}) = \|\overrightarrow{p_i p_j}\| \times \|\overrightarrow{p_i p_\ell}\| \times \sin(\theta_\ell - \theta_j) \leq 0$$

ce qui prouve que $p_j \preceq p_\ell$.