

ANNALES DES CONCOURS

MP

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)

Par

Charles-Pierre ASTOLFI
ENS Cachan

Pierre-Yves BIENVENU
ENS Ulm

Kévin DESTAGNOL
ENS Cachan

Julien DUMONT
Professeur en CPGE

Jean-Julien FLECK
Professeur en CPGE

Alexandre LE MEUR
ENS Cachan

Thierry LIMOGES
ENS Cachan

Benjamin MONMEGE
Enseignant-chercheur à l'université

Florence MONNA
Docteur en mathématiques

Gilbert MONNA
Professeur en CPGE

Matthias MORENO
ENS Lyon

Sophie RAINERO
Professeur en CPGE

Nicolas WEISS
Professeur agrégé

Sommaire

Énoncé
Corrigé

CONCOURS COMMUNS POLYTECHNIQUES

Mathématiques 1	Autour du théorème de Weierstrass. <i>suites et séries de fonctions, convergences simple et uniforme, loi de Poisson</i>	17	21
Mathématiques 2	Surjection de l'exponentielle de matrices. <i>algorithmique, projection orthogonale, séries de matrices, éléments propres</i>	37	44

CENTRALE-SUPÉLEC

Mathématiques 1	Autour de la transformation de Radon. <i>géométrie du plan, intégrales à paramètre, changement de variable, fonctions de deux variables</i>	61	65
Mathématiques 2	Autour des sommes d'Euler. <i>intégrales à paramètre, suites et séries de fonctions</i>	87	91
Informatique commune	Autour de la dynamique gravitationnelle. <i>listes, boucles, schémas d'intégration, méthode d'Euler, bases de données</i>	119	123
Informatique optionnelle	Coloration d'un graphe d'intervalles. <i>graphes, programmation, complexité</i>	139	145

MINES-PONTS

Mathématiques 1	Opérateur de Volterra et équations différentielles. <i>endomorphismes symétriques, équation différentielle d'ordre 2, séries trigonométriques, théorème de Weierstrass</i>	163	168
Mathématiques 2	Norme d'une matrice aléatoire. <i>variable aléatoire réelle, topologie, norme matricielle</i>	183	188
Informatique commune	Tests de validation d'une imprimante. <i>algorithmique, bases de données, méthode d'Euler, méthode des trapèzes</i>	209	219
Informatique optionnelle	Automates d'arbres. <i>listes, automates, langages</i>	231	241

POLYTECHNIQUE-ENS

Mathématiques A	Enlacements de spectres de matrices symétriques. <i>réduction des matrices symétriques, compacité, groupe orthogonal, racines des polynômes, théorème des valeurs intermédiaires</i>	263	268
Mathématiques B	Quatre problèmes d'analyse. <i>intégrales à paramètre, séries numériques, développement asymptotique, système différentiel</i>	295	302
Informatique MP/PC	Enveloppes convexes dans le plan. <i>tableaux et listes, boucles for et while, piles, complexité</i>	325	332

FORMULAIRES

Développements limités usuels en 0	342
Développements en série entière usuels	343
Dérivées usuelles	344
Primitives usuelles	345
Trigonométrie	348

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

MPMA102

**CONCOURS COMMUNS
POLYTECHNIQUES****EPREUVE SPECIFIQUE - FILIERE MP****MATHEMATIQUES 1****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

Le sujet est composé de deux exercices et d'un problème tous indépendants.

EXERCICE I.

I.1. Soit X une variable aléatoire qui suit une loi de Poisson de paramètre $\lambda > 0$. Déterminer sa fonction génératrice, puis en déduire son espérance et sa variance.

EXERCICE II.

On note $I =]0, +\infty[$ et on définit pour n entier naturel non nul et pour $x \in I$, $f_n(x) = e^{-nx} - 2e^{-2nx}$.

II.1. Justifier que pour tout entier naturel non nul n , les fonctions f_n sont intégrables sur I et calculer $\int_0^{+\infty} f_n(x) dx$. Que vaut alors la somme $\sum_{n=1}^{+\infty} \left(\int_0^{+\infty} f_n(x) dx \right)$?

II.2. Démontrer que la série de fonctions $\sum_{n \geq 1} f_n$ converge simplement sur I . Déterminer sa fonction

somme S et démontrer que S est intégrable sur I . Que vaut alors $\int_0^{+\infty} \left(\sum_{n=1}^{+\infty} f_n(x) \right) dx$?

II.3. Donner, sans aucun calcul, la nature de la série $\sum_{n \geq 1} \left(\int_0^{+\infty} |f_n(x)| dx \right)$.

PROBLEME.

Toutes les fonctions étudiées dans ce problème sont à valeurs réelles. On pourra identifier un polynôme et la fonction polynomiale associée.

On rappelle le théorème d'approximation de Weierstrass pour une fonction continue sur $[a, b]$: si f est une fonction continue sur $[a, b]$, il existe une suite de fonctions polynômes (P_n) qui converge uniformément vers la fonction f sur $[a, b]$.

Le problème aborde un certain nombre de situations en lien avec ce théorème qui sera démontré dans la dernière partie.

Partie 1. Exemples et contre-exemples

III.1. Soit h la fonction définie sur l'intervalle $]0, 1]$ par : $\forall x \in]0, 1], \quad x \mapsto \frac{1}{x}$.

Expliquer pourquoi h ne peut être uniformément approchée sur l'intervalle $]0, 1]$ par une suite de fonctions polynômes. Analyser ce résultat par rapport au théorème de Weierstrass.

III.2. Soit N entier naturel non nul, on note $\mathcal{P}_N$ l'espace vectoriel des fonctions polynômiales sur $[a, b]$, de degré inférieur ou égal à N . Justifier que $\mathcal{P}_N$ est une partie fermée de l'espace des applications continues de $[a, b]$ dans $\mathbb{R}$ muni de la norme de la convergence uniforme.

Que peut-on dire d'une fonction qui est limite uniforme sur $[a, b]$ d'une suite de polynômes de degré inférieur ou égal à un entier donné ?

III.3. Cette question illustre la dépendance d'une limite vis-à-vis de la norme choisie.

Soit $\mathbb{R}[X]$ l'espace vectoriel des polynômes à coefficients réels. Soient N_1 et N_2 deux applications définies sur $\mathbb{R}[X]$ ainsi :

$$\text{pour tout polynôme } P \text{ de } \mathbb{R}[X], \quad N_1(P) = \sup_{x \in [-2, -1]} |P(x)| \quad \text{et} \quad N_2(P) = \sup_{x \in [1, 2]} |P(x)|.$$

III.3.a. Vérifier que N_1 est une norme sur $\mathbb{R}[X]$. On admettra que N_2 en est également une.

III.3.b. On note f la fonction définie sur l'intervalle $[-2, 2]$ ainsi :

pour tout $x \in [-2, -1]$, $f(x) = x^2$, pour tout $x \in [-1, 1]$, $f(x) = 1$ et pour tout $x \in [1, 2]$, $f(x) = x^3$.

Représenter graphiquement la fonction f sur l'intervalle $[-2, 2]$ et justifier l'existence d'une suite de fonctions polynômes (P_n) qui converge uniformément vers la fonction f sur $[-2, 2]$.

CCP Maths 1 MP 2015 — Corrigé

Ce corrigé est proposé par Thierry Limoges (ENS Cachan) ; il a été relu par François Lê (ENS Lyon) et Sophie Rainero (Professeur en CPGE).

Ce sujet comporte deux exercices et un problème indépendants.

Le premier exercice est une question de cours à propos de la loi de Poisson. Le deuxième traite un exemple de série de fonctions qui converge simplement mais ne s'intègre pas terme à terme.

Le problème est constitué de quatre parties indépendantes autour du théorème d'approximation de Weierstrass. De nombreux exemples et contre-exemples y sont abordés.

- La première partie illustre la dépendance de la limite vis-à-vis de la norme choisie dans un espace vectoriel normé de dimension infinie et traite un contre-exemple du théorème d'approximation de Weierstrass lorsque l'ensemble de définition n'est plus un segment.
- Dans la deuxième partie, on démontre un théorème des moments : si tous les moments d'ordre k d'une fonction f continue sur $[a; b]$ sont nuls, c'est-à-dire si

$$\forall k \in \mathbb{N} \quad \int_a^b x^k f(x) \, dx = 0$$

alors f est identiquement nulle sur $[a; b]$. On utilise ce résultat dans un cadre préhilbertien puis on établit qu'il ne s'étend pas sur un intervalle quelconque.

- La troisième partie étudie la convergence d'une suite de polynômes définis grâce à une suite récurrente.
- Enfin, la quatrième partie démontre le théorème d'approximation de Weierstrass via une méthode probabiliste, en utilisant les polynômes de Bernstein. On peut retrouver cet exercice dans la première épreuve de mathématiques du concours Mines-Ponts de la filière MP en 2015.

Ce sujet contient des exemples variés qui illustrent les théorèmes autour de la convergence uniforme, ainsi que les limites de leur cadre d'application (segments, intervalles). Il permet de s'approprier de manière solide ces notions de convergence simple ou uniforme.

INDICATIONS

- II.1 Comparer chaque terme à la fonction $x \mapsto 1/x^2$ en $+\infty$.
- II.2 Trouver des sommes géométriques, simplifier la formule de S, effectuer le changement de variable $u = e^x$, puis une décomposition en éléments simples pour calculer son intégrale.
- II.3 Raisonner par l'absurde et appliquer le théorème d'intégration terme à terme.
- III.1 Établir une contradiction à propos de la limite en 0 d'un polynôme qui approcherait la fonction h uniformément sur $]0; 1]$.
- III.2 Un sous-espace de dimension finie dans un espace vectoriel normé est fermé.
- III.3.a Un polynôme ayant une infinité de racines est le polynôme nul.
- III.3.b Utiliser le théorème de Weierstrass. Montrer que la norme de la différence tend vers 0.
- III.4.a Utiliser la linéarité de l'intégrale.
- III.4.b Montrer que l'intégrale de f^2 est nulle en l'obtenant comme une limite.
- III.5 Interpréter le résultat de la question précédente comme une orthogonalité dans l'espace préhilbertien donné.
- III.6.a Comparer à une intégrale de Riemann, puis effectuer une intégration par parties. Montrer la formule par récurrence.
- III.6.b Utiliser $\sin x = \operatorname{Im}(e^{ix})$ pour $x \in \mathbb{R}$.
- III.6.c Effectuer un changement de variable.
- III.6.d Raisonner par l'absurde.
- III.7 Étudier la fonction g_x sur l'intervalle $[0; 1]$ puis la suite récurrente associée.
- III.8 Penser à une bosse qui se déplace.
- III.9.a Utiliser la question III.7.
- III.9.b Appliquer le théorème de l'énoncé à la suite de fonctions $(P_n)_{n \in \mathbb{N}}$ sur l'intervalle $[0; 1]$.
- III.10.a Appliquer l'inégalité de Bienaymé-Tchebychev à la variable aléatoire S_n .
- III.10.b Utiliser le théorème de transfert.
- III.11.a Montrer l'uniforme continuité de f sur $[0; 1]$.
- III.11.b Montrer que les événements $\{S_n = k\}$ pour les entiers k sur lesquels la somme est indexée sont inclus dans l'événement $\left\{ \left| \frac{S_n}{n} - x \right| > \alpha \right\}$.
- III.11.c Séparer la somme en deux termes, puis les majorer grâce aux questions III.11.a et III.11.b.

EXERCICE I

I.1 Soit X une variable aléatoire qui suit une loi de Poisson de paramètre $\lambda > 0$. Cette variable aléatoire admet donc une série génératrice de rayon supérieur ou égal à 1, et étant à valeurs dans $\mathbb{N}$, c'est par définition la somme de la série entière en la variable t

$$\sum_{k \geq 0} P(X = k) t^k = \sum_{k \geq 0} e^{-\lambda} \frac{\lambda^k}{k!} t^k = \sum_{k \geq 0} e^{-\lambda} \frac{(\lambda t)^k}{k!}$$

Cette série entière est de rayon de convergence infini, de somme

$$\sum_{k=0}^{+\infty} e^{-\lambda} \frac{(\lambda t)^k}{k!} = e^{-\lambda} e^{\lambda t}$$

Par conséquent, pour tout $t \in \mathbb{R}$,

$$G_X(t) = e^{\lambda(t-1)}$$

La fonction G_X est dérivable en 1. Ainsi X admet une espérance et $E(X) = G'_X(1)$. Or, pour tout $t \in \mathbb{R}$, $G'_X(t) = \lambda e^{\lambda(t-1)}$, d'où

$$E(X) = \lambda$$

La fonction G_X est deux fois dérivable en 1 donc X admet une variance. On a

$$G''_X(1) = E(X(X-1)) = E(X^2 - X) = E(X^2) - E(X)$$

par linéarité de l'espérance. Ainsi,

$$V(X) = E(X^2) - E(X)^2 = E(X^2) - E(X) + E(X) - E(X)^2 = G''_X(1) + G'_X(1) - G'_X(1)^2$$

Il faut savoir retrouver cette formule qui ne doit pas être écrite sur la copie sans justification.

Or comme pour tout $t \in \mathbb{R}$, $G''_X(t) = \lambda^2 e^{\lambda(t-1)}$, on en déduit $V(X) = \lambda^2 + \lambda - \lambda^2 = \lambda$, c'est-à-dire

$$V(X) = \lambda$$

EXERCICE II

II.1 Soit $n \in \mathbb{N}^*$. Comme la fonction $f_n: x \mapsto e^{-nx} - 2e^{-2nx}$ est continue sur $[0; +\infty[$, la fonction f_n est intégrable au voisinage de 0. Lorsque x tend vers $+\infty$,

$$e^{-nx} = O\left(\frac{1}{x^2}\right) \quad \text{et} \quad e^{-2nx} = O\left(\frac{1}{x^2}\right) \quad \text{car } n > 0$$

donc

$$f_n(x) = O\left(\frac{1}{x^2}\right)$$

Comme la fonction $x \mapsto 1/x^2$ est intégrable au voisinage de $+\infty$, par comparaison la fonction f_n est intégrable au voisinage de $+\infty$, d'où

$$\text{Pour tout } n \in \mathbb{N}^*, f_n \text{ est intégrable sur } \mathbb{I}.$$

On a

$$\begin{aligned}
 \int_0^{+\infty} f_n(x) \, dx &= \int_0^{+\infty} (e^{-nx} - 2e^{-2nx}) \, dx \\
 &= \left[\frac{e^{-nx}}{-n} - 2 \frac{e^{-2nx}}{-2n} \right]_0^{+\infty} \\
 &= 0 + \frac{1}{n} - 0 + 2 \frac{1}{-2n} \\
 \int_0^{+\infty} f_n(x) \, dx &= 0
 \end{aligned}$$

Ainsi,

$$\boxed{\int_0^{+\infty} f_n(x) \, dx = 0 \quad \text{et donc} \quad \sum_{n=1}^{+\infty} \left(\int_0^{+\infty} f_n(x) \, dx \right) = 0}$$

II.2 Soit $x \in]0; +\infty[$. Comme $x > 0$, e^{-x} et e^{-2x} sont dans $]0; 1[$, donc les séries géométriques $\sum (e^{-x})^n$ et $\sum (e^{-2x})^n$ convergent, ainsi que $\sum f_n(x)$ par linéarité. Calculons sa somme

$$\begin{aligned}
 S(x) &= \sum_{n=1}^{+\infty} f_n(x) \\
 &= \sum_{n=1}^{+\infty} (e^{-nx} - 2e^{-2nx}) \\
 &= \sum_{n=1}^{+\infty} (e^{-x})^n - 2 \sum_{n=1}^{+\infty} (e^{-2x})^n \\
 &= e^{-x} \frac{1}{1 - e^{-x}} - 2e^{-2x} \frac{1}{1 - e^{-2x}} \\
 &= \frac{1}{e^x - 1} - 2 \frac{1}{e^{2x} - 1} \\
 &= \frac{1}{e^x - 1} - 2 \frac{1}{(e^x - 1)(e^x + 1)} \\
 &= \frac{1}{e^x - 1} \left(1 - 2 \frac{1}{e^x + 1} \right) \\
 &= \frac{1}{e^x - 1} \frac{e^x + 1 - 2}{e^x + 1} \\
 S(x) &= \frac{1}{e^x + 1}
 \end{aligned}$$

Étudions l'intégrabilité de la fonction S sur $I =]0; +\infty[$. C'est une fonction continue sur I , qui se prolonge par continuité en 0. Au voisinage de $+\infty$,

$$S(x) = O\left(\frac{1}{x^2}\right)$$

Par conséquent, la fonction S est intégrable au voisinage de $+\infty$. Ainsi,

La série de fonctions $\sum f_n$ converge simplement sur I vers la fonction S , intégrable sur I , définie par $S(x) = 1/(e^x + 1)$.

CCP Maths 2 MP 2015 — Corrigé

Ce corrigé est proposé par Alexandre Le Meur (ENS Cachan) ; il a été relu par Thierry Limoges (ENS Cachan) et Benjamin Monmege (Enseignant-chercheur à l'université).

Ce sujet comporte deux exercices et un problème indépendants les uns des autres. Le premier exercice est consacré à l'informatique, le reste à l'algèbre linéaire.

- Le premier exercice étudie d'un programme sommant les chiffres d'un entier naturel. Sans difficulté mathématique, il permettait de s'assurer que les candidats avaient des connaissances de base en algorithmique et dans le langage Python.
- Le deuxième exercice est très proche du cours. On détermine une base orthonormée (pour le produit scalaire canonique) de l'espace des matrices triangulaires supérieures, puis on utilise cette base pour calculer le projeté orthogonal d'une matrice.

Le but du problème est de déterminer, pour $A \in \mathcal{M}_3(\mathbb{C})$ donnée, des matrices M et B telles que $\exp(M) = A$ et $B^2 = A$, c'est-à-dire un logarithme et une racine carrée de A .

- Une partie de préliminaires étudie l'exponentielle de matrices en montrant la convergence de la série correspondante à l'aide d'une norme d'algèbre définie dans l'énoncé. Le principal piège de cette partie est de confondre la convergence dans $\mathbb{C}$ avec celle dans $\mathcal{M}_3(\mathbb{C})$.
- La première partie s'intéresse à un exemple où A est à coefficients réels ; on montre qu'aucune matrice à coefficients réels ne vérifie les relations définissant M et B . Ceci prouve que l'hypothèse sur le corps de base est essentielle.
- Dans la deuxième partie, la plus importante, on construit très graduellement une fonction dérivable $\gamma: \mathbb{R} \rightarrow \mathcal{M}_3(\mathbb{C})$ dont les images en -1 et $1/2$ valent A^{-1} et B respectivement, et telle que $\gamma'(0) = M$.
- La troisième partie traite un exemple. On calcule explicitement γ puis on détermine des matrices B et M .

Ce sujet est de longueur moyenne ; quelques questions, nettement plus longues que les autres, peuvent déstabiliser. Il constitue une bonne révision de l'algèbre linéaire et euclidienne.

INDICATIONS

Exercice I

- I.1 Penser à la division euclidienne.
- I.3.a Montrer que la suite d'entiers naturels $(n_k)_{k \in \mathbb{N}}$ décroît strictement.
- I.4 Identifier ce que fait la procédure `mystere` en s'aidant de la question II.2.
- I.5 Un algorithme récursif est une procédure faisant appel à elle-même.

Exercice II

- II.2 Penser à la base canonique de $\mathcal{M}_2(\mathbb{R})$. Peut-on en extraire une base de $\mathcal{J}$? Cette base est-elle orthogonale pour le produit scalaire indiqué?
- II.3 Utiliser l'expression de la projection en fonction du produit scalaire avec les éléments d'une base orthonormée de $\mathcal{J}$. Autre possibilité: trouver la décomposition de A suivant la somme directe $\mathcal{M}_2(\mathbb{R}) = \mathcal{J} \oplus \mathcal{J}^\perp$.

Problème III

- III.1 Expliciter les coefficients du produit AB en fonction de ceux de A et B et les majorer.
- III.2 Utiliser le fait que dans $\mathbb{C}$, la convergence absolue entraîne la convergence.
- III.3 Utiliser les questions III.2 et III.1.
- III.4 Si M est semblable à une matrice triangulaire T, écrire une expression de T^k pour $k \in \mathbb{N}$ et en déduire un lien entre $\det(\exp(M))$ et $\det(\exp(T))$.
- III.5 Raisonner par l'absurde et montrer que cela apporte une contradiction avec le calcul du déterminant de A.
- III.6.b Étudier le cas où f est de la forme $x^k \rho^x e^{ix\theta}$ puis conclure dans le cas général.
- III.7.b Devant une égalité de fonctions, penser à évaluer en des valeurs remarquables ou à calculer des limites.
- III.8 Penser à distinguer tous les cas suivant le nombre de valeurs propres distinctes de A. En s'inspirant de la page d'exemple de l'énoncé, montrer que les coefficients du reste de la division euclidienne de X^n par χ_A sont solutions d'un système linéaire et montrer qu'ils sont dans F.
- III.9.c Utiliser la question III.6.b pour montrer que les fonctions f et g sont dans F, puis conclure avec la question III.7.c.
- III.11 Utiliser la dérivabilité des fonctions $\omega_{i,j}$ sur $\mathbb{R}$. Expliciter la solution de l'équation différentielle (dans $\mathcal{M}_3(\mathbb{C})$) $u' = \gamma'(0)u$ vérifiant $u(0) = I_3$. Comparer cette solution avec la fonction γ .
- III.12 Déterminer les valeurs propres de A et la dimension des sous-espaces propres.
- III.13 Cette série de questions est calculatoire. Veiller à poser les calculs proprement et penser à vérifier les résultats intermédiaires. S'inspirer de la page d'exemple de l'énoncé pour calculer la fonction γ puis utiliser les questions III.10 et III.11 pour calculer A^{-1} , B et M.

I. INFORMATIQUE

I.1 Pour décomposer 21 en base 2, appliquons l'algorithme d'Euclide à 21 et 2

$$\begin{aligned} 21 &= 2 \times 10 + 1 \\ 10 &= 2 \times 5 + 0 \\ 5 &= 2 \times 2 + 1 \\ 2 &= 2 \times 1 + 0 \\ 1 &= 2 \times 0 + 1 \end{aligned}$$

La liste des restes en remontant l'algorithme donne la décomposition en base 2 de 21.

$$21 = \overline{10101}$$

On peut également faire cette question « à tâtons » en remarquant que $21 = 16 + 4 + 1$. Cependant cette technique admet rapidement ses limites si le nombre à écrire devient grand.

I.2 Avant le premier passage dans la boucle while, les variables n , b et t ont pour valeurs $n = 256$, $b = 10$ et $t = []$. La variable b ne sera pas modifiée durant l'exécution de la fonction. On obtient donc le tableau suivant

k	1	2	3
c_k	6	5	2
t_k	[6]	[6, 5]	[6, 5, 2]
n_k	25	2	0

La fonction `mystere` renvoie $t = [6, 5, 2]$.

I.3.a Montrons que pour tout entier k , $n_k \neq 0$ entraîne que $n_{k+1} < n_k$. Soit $k \in \mathbb{N}$. L'entier n_{k+1} est le quotient de la division euclidienne de n_k par 10. Par définition

$$n_k = 10n_{k+1} + c_k$$

où $0 \leq c_k < 10$. Ainsi $n_{k+1} \leq \frac{n_k}{10} < n_k$ car $n_k \neq 0$. La suite $(n_k)_{k \in \mathbb{N}}$ est une suite strictement décroissante d'entiers naturels. Donc il existe $k_1 \in \mathbb{N}$ tel que $n_{k_1} = 0$. En conclusion,

La boucle `while` termine.

I.3.b Montrons par récurrence que la propriété

$$\mathcal{P}(k) : n_k \leq \frac{n}{10^k}$$

est vraie pour tout $k \in \llbracket 0; p \rrbracket$.

- $\mathcal{P}(0)$ est trivialement vraie car $n = n_0$.
- $\mathcal{P}(k) \implies \mathcal{P}(k+1)$: soit un entier $k \in \llbracket 0; p-1 \rrbracket$. On a

$$n_{k+1} \leq \frac{n_k}{10}$$

Or d'après $\mathcal{P}(k)$, $n_k \leq \frac{n}{10^k}$. En combinant ces deux inégalités, il vient

$$n_{k+1} \leq \frac{n}{10^{k+1}}$$

- Conclusion :

$$\forall k \in \llbracket 0; p \rrbracket \quad n_k \leq \frac{n}{10^k}$$

En particulier, pour $k = p - 1$, on obtient

$$n_{p-1} \leq \frac{n}{10^{p-1}}$$

Or $n_{p-1} > 0$ par définition du nombre p . On a alors

$$10^{p-1} \leq \frac{n}{n_{p-1}}$$

puis en composant par le logarithme en base 10

$$p - 1 \leq \log_{10}(n) - \log_{10}(n_{p-1})$$

puisque $\log_{10}(n_{p-1}) \geq 0$, n_{p-1} étant strictement positif. Ainsi,

$$\boxed{p \leq \log_{10}(n) + 1}$$

I.4 En s'inspirant de la procédure `mystere(n)`, on réécrit un programme qui somme les restes dans l'algorithme d'Euclide au lieu de les stocker dans une liste.

```
def somme_chiffres(n):
    """Données: n>0 un entier.
       Résultat: la somme des chiffres de n"""
    s=0
    while n>0:
        c=n%10
        s=c+s
        n=n//10
    return s
```

I.5 La somme des chiffres de n est égale au chiffre des unités ajouté à la somme des chiffres du quotient de la division euclidienne de n par 10. De plus, si $n = 0$ la procédure doit retourner 0.

```
def somme_rec(n):
    """Données: n>0 un entier
       Résultat: la somme des chiffres de n,
       calculée récursivement"""
    if n==0:
        return 0
    else:
        return somme_rec(n//10)+n%10
```

Centrale Maths 1 MP 2015 — Corrigé

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

Ce sujet étudie la *transformation de Radon*, qui associe à une fonction f de $\mathbb{R}^2$ une fonction $\widehat{f}$ définie par

$$\forall (q, \theta) \in \mathbb{R}^2 \quad \widehat{f}(q, \theta) = \int_{-\infty}^{+\infty} f(q \cos \theta - t \sin \theta, q \sin \theta + t \cos \theta) dt$$

L'objectif du problème est d'établir une formule d'inversion qui permet de retrouver f à partir de $\widehat{f}$.

Les trois premières parties sont complètement indépendantes. La quatrième reprend les résultats montrés dans la partie III, et la cinquième dépend de la partie I.

- La première partie ne comporte que des questions très élémentaires sur la notion de groupe et sur la géométrie plane. On y introduit le groupe G des isométries affines directes du plan euclidien (c'est-à-dire les composés d'une rotation de centre O et d'une translation du plan) ainsi que l'ensemble des droites affines du plan.
- La deuxième partie, également facile, permet de définir la transformée de Radon dans le cas particulier des fonctions du plan qui sont invariantes par rotations vectorielles. On y définit une moyenne radiale $\overline{f}$, et on établit une équation faisant le lien entre $\widehat{f}$ et $\overline{f}$. Les principaux outils utilisés sont les intégrales généralisées et le théorème de changement de variable.
- La troisième partie s'attaque au cas général de fonctions qui satisfont certaines hypothèses de décroissance à l'infini. On y étudie la régularité de la fonction $\overline{f}$, ce qui fait intervenir le théorème de dérivation sous le signe intégral.
- La quatrième partie est la plus longue et la plus difficile du problème. On y démontre la formule d'inversion de Radon. Tous les théorèmes d'intégration sont utilisés, ainsi que les résultats des parties précédentes.
- La cinquième et dernière partie est élémentaire et fait la synthèse du problème. On y étudie une application de l'opérateur $f \mapsto \widehat{f}$ dans le domaine de l'imagerie médicale. Des points pouvaient aisément y être glanés.

Ce problème permet de travailler de manière approfondie les intégrales à paramètre. Sa teneur géométrique le rend particulièrement intéressant et permet de comprendre toutes les formules abordées de façon intuitive.

INDICATIONS

Partie I

- I.A.3 Chercher un inverse de la forme $M(A, \vec{b}')$.
- I.B.3 Écrire une relation du type $\overrightarrow{PM} = t\vec{u}$, où $\vec{u}$ dirige $\Delta(q, \vec{u}_\theta)$ et $P \in \Delta(q, \vec{u}_\theta)$.
- I.B.4 Injecter les formules de la question I.B.3 pour $\Delta(q, \vec{u})$ dans l'équation cartésienne de $\Delta(r, \vec{v})$ donnée par la question I.B.2.
- I.C.3 Montrer que toute droite affine peut s'écrire sous la forme $\Delta(q, \vec{u}_\theta)$.
- I.C.4.a Utiliser la question I.B.4.

Partie II

- II.A.2 Faire apparaître la dérivée de Arctan dans le calcul de $\hat{f}$.
- II.A.3 Calculer explicitement la fonction R.
- II.B.2 Au voisinage de q , décomposer $\sqrt{r^2 - q^2} = \sqrt{r+q}\sqrt{r-q}$ et minorer. Au voisinage de $+\infty$, majorer $\frac{r}{\sqrt{r^2 - q^2}}$ pour $r > 2q > 0$.
- II.B.3 Effectuer le changement de variable $r = \sqrt{q^2 + t^2}$.

Partie III

- III.B Exploiter les symétries de cos et sin.
- III.C.1 Utiliser le théorème de dérivation des intégrales à paramètre.
- III.C.2 Utiliser le fait que $f \in \mathcal{B}_1$.
- III.C.3 La question III.C.1 permet de dériver sous l'intégrale.

Partie IV

- IV.A.1 Effectuer le changement de variable $u = \sqrt{t^2 - 1}$.
- IV.A.2 On a $\tan(\operatorname{Arccos} \alpha) = \frac{\sqrt{1 - \alpha^2}}{\alpha}$ pour $\alpha \in]0; 1[$.
- IV.B.1 Appliquer le théorème de continuité des intégrales à paramètre. Utiliser la question IV.A.1 pour la domination.
- IV.B.2 Après majoration de l'intégrande, utiliser la question IV.A.1 pour simplifier.
- IV.B.3 Utiliser le théorème de dérivation des intégrales à paramètre.
- IV.C.1 Effectuer le changement de variables $r = qt$. Utiliser les parties III.C et IV.B.
- IV.C.2 Faire une intégration par parties sur $]\varepsilon; +\infty[$.
- IV.D.1 Établir que $-\frac{1}{\pi} \int_0^{+\infty} \frac{F'(q)}{q} dq = \lim_{\varepsilon \rightarrow 0} \frac{2}{\pi} \int_1^{+\infty} \frac{\bar{f}(\varepsilon u)}{u\sqrt{u^2 - 1}} du$, puis calculer cette limite à l'aide du théorème de convergence dominée.
- IV.D.2 Regarder l'exemple de la sous-partie II.A.
- IV.D.3 Utiliser la formule d'inversion de Radon au point $(0, 0)$ pour la fonction

$$(x, y) \longmapsto f(x + x_0, y + y_0)$$

I. PRÉLIMINAIRES GÉOMÉTRIQUES

I.A.1 Prenons $A = I_2$ et $\vec{b} = \vec{0} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$. Alors $A \in SO(2)$ et $\vec{b} \in \mathbb{R}^2$. Il vient

$$\begin{aligned} M(A, \vec{b}) &= \left(\begin{array}{c|c} I_2 & \begin{pmatrix} 0 \\ 0 \end{pmatrix} \\ \hline 0 & 1 \end{array} \right) \\ &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \end{aligned}$$

c'est-à-dire $M(A, \vec{b}) = I_3$

Ainsi, $\boxed{(I_2, \vec{0}) \in SO(2) \times \mathbb{R}^2 \quad \text{et} \quad M(I_2, \vec{0}) = I_3}$

I.A.2 Soient $(A, \vec{b})$ et $(A', \vec{b}')$ dans $SO(2) \times \mathbb{R}^2$. Notons $\vec{b} = \begin{pmatrix} b_1 \\ b_2 \end{pmatrix}$ et $\vec{b}' = \begin{pmatrix} b'_1 \\ b'_2 \end{pmatrix}$.

Alors

$$\begin{aligned} M(A, \vec{b}) \cdot M(A', \vec{b}') &= \left(\begin{array}{c|c} A & \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} \\ \hline 0 & 1 \end{array} \right) \cdot \left(\begin{array}{c|c} A' & \begin{pmatrix} b'_1 \\ b'_2 \end{pmatrix} \\ \hline 0 & 1 \end{array} \right) \\ &= \left(\begin{array}{c|c} AA' & A \begin{pmatrix} b'_1 \\ b'_2 \end{pmatrix} + \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} \\ \hline 0 & 1 \end{array} \right) \\ &= \left(\begin{array}{c|c} AA' & A\vec{b}' + \vec{b} \\ \hline 0 & 1 \end{array} \right) \end{aligned}$$

$$\boxed{M(A, \vec{b}) \cdot M(A', \vec{b}') = M(AA', A\vec{b}' + \vec{b})}$$

I.A.3 Soit $(A, \vec{b}) \in SO(2) \times \mathbb{R}^2$. Cherchons un inverse de $M(A, \vec{b})$ de la forme $M(A', \vec{b}')$. D'après la question I.A.2, on a $M(A, \vec{b}) \cdot M(A', \vec{b}') = M(AA', A\vec{b}' + \vec{b})$. Pour que $M(AA', A\vec{b}' + \vec{b}) = I_3$, d'après le résultat trouvé à la question I.A.1, il suffit d'avoir

$$\begin{cases} AA' = I_2 \\ A\vec{b}' + \vec{b} = \vec{0} \end{cases}$$

Puisque $A \in SO(2)$ est inversible, prenons

$$\begin{cases} A' = A^{-1} \\ \vec{b}' = -A^{-1}\vec{b} \end{cases}$$

Vérifions que $M(A^{-1}, -A^{-1}\vec{b})$ est bien l'inverse de $M(A, \vec{b})$. En utilisant la question I.A.2, on a

$$\begin{aligned} M(A, \vec{b}) \cdot M(A^{-1}, -A^{-1}\vec{b}) &= M(AA^{-1}, A(-A^{-1}\vec{b}) + \vec{b}) \\ &= M(I_2, -\vec{b} + \vec{b}) \\ &= M(I_2, \vec{0}) \end{aligned}$$

donc $M(A, \vec{b}) \cdot M(A^{-1}, -A^{-1}\vec{b}) = I_3$

avec le résultat de la question I.A.1. On a montré que

Tout $M(A, \vec{b}) \in G$ est inversible avec $M(A, \vec{b})^{-1} = M(A^{-1}, -A^{-1}\vec{b})$.

Pour une matrice carrée, être inversible, avoir un inverse à gauche et avoir un inverse à droite sont équivalents. Les inverses à droite et à gauche sont toujours égaux lorsqu'ils existent.

I.A.4 Vérifions les quatre conditions pour être un sous-groupe de $GL_3(\mathbb{R})$.

- Non vide : d'après la question I.A.1, $I_3 \in G$, donc G est non vide.
- Inclusion dans $GL_3(\mathbb{R})$: d'après la question I.A.3, les éléments de G sont des matrices carrées inversibles, donc G est un sous-ensemble de $GL_3(\mathbb{R})$.
- Stabilité par passage à l'inverse : pour tout $M(A, \vec{b}) \in G$, on a

$$M(A, \vec{b})^{-1} = M(A^{-1}, -A^{-1}\vec{b})$$

et si $A \in SO(2)$, alors $A^{-1} \in SO(2)$. Comme $-A^{-1}\vec{b} \in \mathbb{R}^2$, il vient

$$M(A, \vec{b})^{-1} \in G$$

- Stabilité par produit : soit $(A, \vec{b})$ et $(A', \vec{b}')$ dans $SO(2) \times \mathbb{R}^2$. D'après la question I.A.2 on a

$$M(A, \vec{b}) \cdot M(A', \vec{b}') = M(AA', A\vec{b}' + \vec{b})$$

Comme $SO(2)$ est un groupe, $AA' \in SO(2)$. On a également $A\vec{b}' + \vec{b} \in \mathbb{R}^2$ d'où

$$M(A, \vec{b}) \cdot M(A', \vec{b}') \in G$$

Conclusion :

G est un sous-groupe de $GL_3(\mathbb{R})$.

I.A.5 Soit $\vec{b} \in \mathbb{R}^2$. Alors $M(I_2, \vec{b}) \in G$ et $\Phi(M(I_2, \vec{b})) = \vec{b}$, donc

Φ est surjective.

Soient deux matrices A et B dans $SO(2)$ distinctes. Alors $M(A, \vec{b}) \in G$ et $M(B, \vec{b}) \in G$ avec

$$M(A, \vec{b}) \neq M(B, \vec{b})$$

et

$$\Phi(M(A, \vec{b})) = \vec{b} = \Phi(M(B, \vec{b}))$$

donc

Φ n'est pas injective.

Centrale Maths 2 MP 2015 — Corrigé

Ce corrigé est proposé par Nicolas Weiss (Docteur en mathématiques); il a été relu par Pauline Tan (ENS Cachan) et Benjamin Monmege (Enseignant-chercheur à l'université).

Le but annoncé de ce sujet est la démonstration d'une relation due à Euler :

$$2 \sum_{n=1}^{+\infty} \frac{H_n}{(n+1)^r} = r\zeta(r+1) - \sum_{k=1}^{r-2} \zeta(k+1)\zeta(r-k)$$

pour tout entier $r \geq 2$, où $\zeta: x \mapsto \sum_{n=1}^{+\infty} \frac{1}{n^x}$ est la fonction ζ de Riemann et $H_n = \sum_{k=1}^n \frac{1}{k}$ pour tout $n \geq 1$ est une somme partielle de la série harmonique. C'est l'occasion de s'intéresser aux fonctions spéciales Γ , β et digamma. Les principaux outils employés sont les intégrales généralisées et les intégrales à paramètre.

- La première partie prépare le terrain avec des calculs de séries et d'intégrales qui resserviront dans la suite. Elle aboutit à un premier calcul de

$$S_2 = \sum_{n=1}^{+\infty} \frac{H_n}{(n+1)^2}$$

- La deuxième partie étudie des propriétés des fonctions spéciales Γ et β et démontre la relation $\beta(x, y) = \Gamma(x)\Gamma(y)/\Gamma(x+y)$.
- La troisième partie concerne la fonction digamma ψ et l'exprime comme une série de fonctions liée à la fonction ζ :

$$\forall x \in]-1; 1[\quad \psi(1+x) = \psi(1) + \sum_{n=1}^{+\infty} (-1)^{n+1} \zeta(n+1) x^n$$

- La quatrième et dernière partie aboutit à la relation d'Euler en reliant la fonction digamma à certaines dérivées partielles de la fonction β .

L'ensemble du sujet demande de la précision dans la vérification des hypothèses des théorèmes du cours. Si la majorité des questions peuvent se traiter par un raisonnement direct, un bon nombre demandent de s'appuyer sur les résultats obtenus précédemment. Notons également que la longueur du sujet paraît disproportionnée par rapport aux 4 heures dont disposaient les candidats pour le traiter.

INDICATIONS

Partie I

I.A.2 Le reste d'une série convergente tend vers zéro.

I.C.2 Effectuer le produit de Cauchy des développements en série entière donnés en I.C.1.

I.D.3 Calculer $\lim_{\varepsilon \rightarrow 0^+} I_{p,q}^\varepsilon$ en utilisant la question I.D.2.

I.E Intégrer terme à terme la série de fonctions de terme général

$$t \longmapsto a_n t^n (\ln t)^{r-1}$$

I.F.1 Appliquer la question I.E à $t \longmapsto -\ln(1-t)/(1-t)$.

I.F.2 Intégrer par parties dans l'égalité de la question I.F.1.

I.F.3 Pour calculer S_2 , effectuer le changement de variable $u = 1 - t$. Pour relier S_2 à $\zeta(3)$, intégrer terme à terme à l'aide du développement en série entière de $1/(1-t)$.

Partie II

II.B.3 Intégrer par parties $\beta(x+1, y)$ et calculer $\beta(x, y+1)$ en fonction de $\beta(x, y)$ et $\beta(x+1, y)$.

II.C.1 Utiliser l'identité $\Gamma(x+1) = x\Gamma(x)$ et la question II.B.4 pour évaluer les deux membres de la relation $(\mathcal{R})$.

II.C.4 Appliquer le théorème de continuité d'une intégrale à paramètre.

II.C.5 Se servir du théorème de convergence dominée.

II.C.6 Appliquer le théorème de dérivation d'une intégrale à paramètre en utilisant la question II.A.2 pour l'hypothèse de domination sur tout segment.

II.C.8 Montrer que $G(a) = \int_0^a G'(t) dt$ puis $\lim_{a \rightarrow +\infty} G(a) = \Gamma(x)\Gamma(y)$ et conclure avec les questions II.C.5 et II.C.1.

Partie III

III.A Dériver l'identité $\Gamma(x+1) = x\Gamma(x)$.

III.B.1 Dériver la relation $(\mathcal{R})$ par rapport à y .

III.B.2 Utiliser la croissance de l'intégrale définissant β .

III.B.3 Combiner les questions III.B.1 et III.B.2.

III.C.1 Raisonner par récurrence sur n en utilisant la question III.A.

III.C.2 Utiliser les questions III.A, III.C.1 et III.B.3 et établir $\psi(1+p) - \psi(1) = H_p$.

III.C.3 Faire tendre n vers l'infini dans le résultat de la question III.C.2.

III.D.1 Appliquer le théorème de la classe $\mathcal{C}^\infty$ d'une série de fonctions.

III.D.2 Réécrire $g(x)$ et $\sum_{k=0}^n \frac{g^{(k)}(0)}{k!} x^k$ en terme de $\frac{x^k}{\ell^{k+1}}$ pour $\ell \geq 2$.

III.D.3 Utiliser le développement en série entière de $1/(1+x)$.

Partie IV

- IV.B.1 Appliquer le théorème de la classe $\mathcal{C}^2$ d'une intégrale à paramètre avec $x > 0$ et $y \geq 1$ pour pouvoir appliquer la question I.F.2, puis évaluer dans le cas $y = 1$.
- IV.B.3 Utiliser les questions I.F.2 et IV.B.2 en appliquant le théorème de continuité d'une intégrale à paramètre.
- IV.B.4 Calculer la limite du résultat de la question IV.B.3 à l'aide de la formule donnée en question IV.A.
- IV.C.1 Dériver φ à partir de sa définition, grâce à la formule de Leibniz.
- IV.C.2 Combiner les questions IV.A, IV.B.3, IV.C.1 et III.D.3.

I. REPRÉSENTATION INTÉGRALE DE SOMMES DE SÉRIES

I.A.1 Les nombres a_n sont bien définis dès $n \geq 2$ par continuité de la fonction inverse sur le segment $[n-1; n]$. Par le théorème de comparaison série/intégrale, on sait que la positivité sur $\mathbb{R}_+$, la continuité (donc continuité par morceaux) et la décroissance de la fonction inverse assurent ensemble la convergence de la série de terme général

$$\int_{n-1}^n \frac{dt}{t} - \frac{1}{n}$$

On reconnaît ci-dessus l'opposé du terme général a_n .

La série de terme général a_n est convergente.

On peut aussi prouver la convergence de la série de terme général a_n en majorant $|a_n|$ par $1/(n(n-1))$ à l'aide de la décroissance de la fonction inverse sur le segment $[n-1; n]$.

I.A.2 Commençons par évaluer la somme partielle $\sum_{k=2}^n a_k$ pour $n \geq 2$.

$$\begin{aligned} \sum_{k=2}^n a_k &= \sum_{k=2}^n \left(\frac{1}{k} - \int_{k-1}^k \frac{dt}{t} \right) \\ &= \sum_{k=2}^n \frac{1}{k} - \sum_{k=2}^n \int_{k-1}^k \frac{dt}{t} \\ &= H_n - 1 - \int_1^n \frac{dt}{t} \\ \sum_{k=2}^n a_k &= H_n - 1 - \ln n \end{aligned}$$

D'après la question I.A.1, la série $\sum a_n$ converge. Posons donc

$$A = \sum_{k=2}^{+\infty} a_k + 1$$

Ainsi, pour $n \geq 2$,

$$\sum_{k=2}^n a_k = A - 1 - \sum_{k=n+1}^{+\infty} a_k \underset{n \rightarrow +\infty}{=} A - 1 + o(1)$$

puisque le reste d'une série convergente tend vers zéro. En réunissant les deux calculs de $\sum_{k=2}^n a_k$ qui précèdent, on obtient

$$H_n - 1 - \ln n \underset{n \rightarrow +\infty}{=} A - 1 + o(1)$$

soit

Il existe une constante réelle A telle que $H_n \underset{n \rightarrow +\infty}{=} \ln n + A + o(1)$.

Factorisons l'expression ci-dessus par le terme $\ln n$, non nul pour $n > 1$:

$$H_n \underset{n \rightarrow +\infty}{=} \ln n \times \left(1 + \frac{A}{\ln n} + o\left(\frac{1}{\ln n}\right) \right) = \ln n \times (1 + o(1))$$

Cela se réécrit

$$H_n \sim \ln n$$

Centrale Informatique commune MP 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 croisiez 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.

Centrale Informatique optionnelle MP 2015

Corrigé

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

Ce sujet consiste en l'étude d'un problème d'allocation de ressources. Un ensemble de tâches (cours, avion en phase d'atterrissage, requête d'un client informatique, etc.) est modélisé par un ensemble d'intervalles correspondant au temps nécessaire à la réalisation de chacune d'elles. Il s'agit alors d'allouer une ressource (salle, piste d'atterrissage, serveur informatique, etc.) à chaque tâche de sorte qu'une ressource ne soit jamais utilisée au même instant par deux tâches différentes. On cherche par ailleurs à minimiser le nombre de ressources nécessaires. Au long de quatre parties fortement dépendantes, le sujet propose une solution grâce à l'étude de colorations de graphes d'intervalles.

- La partie I introduit les graphes d'intervalles, qui permettent de représenter à l'aide de graphes non orientés les contraintes d'intervalles d'un problème d'allocation de ressources. Une allocation de ressources consiste alors en une coloration du graphe, c'est-à-dire à l'étiquetage de chaque sommet par une couleur de telle sorte que les sommets voisins ne soient jamais étiquetés avec la même couleur. Un lien est établi entre le nombre chromatique d'un graphe, c'est-à-dire le nombre minimal de couleurs nécessaire pour le colorier, et la taille de la plus grande clique (sous-graphe complet) du graphe. Certaines fonctions utiles pour la suite du sujet sont également implémentées.
- La partie II propose un algorithme glouton pour colorier un graphe d'intervalles. Son implémentation, la preuve de sa correction et sa complexité sont étudiées.
- La partie III démontre que l'algorithme glouton précédent reste correct dans le cas plus général de graphes munis d'un ordre d'élimination parfait, c'est-à-dire une énumération de leurs sommets telle que, pour chaque sommet, ses voisins énumérés avant lui forment une clique.
- La partie IV étudie finalement une condition suffisante pour qu'un graphe possède un ordre d'élimination parfait, à savoir que le graphe soit cordal : un graphe est cordal si tout cycle de longueur supérieure à quatre possède une corde. Les graphes cordaux sont couramment appelés graphes triangulés. Après avoir montré que tout graphe d'intervalles est cordal, le sujet demande de résoudre une énigme policière à l'aide de cet outil. Le reste du sujet implémente la recherche d'ordres d'élimination parfaits lorsqu'ils existent, puis prouve que tout graphe cordal possède un ordre d'élimination parfait.

Les parties I à III du sujet sont des applications directes du cours, où l'on manipule des graphes représentés par des listes d'adjacence. L'énigme policière, ainsi que les questions d'implémentation et de complexité de la partie IV, sont nettement plus difficiles et discriminantes : elles permettent d'évaluer la créativité et la compréhension globale des parties précédentes. Enfin, les trois dernières sections de la partie IV proposent une preuve élégante de théorie des graphes.

INDICATIONS

- I.C.2 Utiliser la fonction `conflit` de la question I.A.1 afin de construire itérativement le tableau de listes de voisins de $G(I_0, \dots, I_{n-1})$.
- I.E.2 Montrer que $\omega(G) \leq \chi(G)$.
- I.E.3 Remarquer que tester si un ensemble $\{x_0, \dots, x_{p-1}\} \subset S$ de sommets est une clique de G consiste à vérifier que, pour tout $0 \leq q < p$, tous les sommets de $\{x_{q+1}, \dots, x_{p-1}\}$ sont voisins avec x_q . Utiliser la fonction `appartient` de la question I.D.2.a.
- II.B Employer la fonction `couleur_disponible` de la question I.D.2.d.
- II.C.1 Les intervalles sont énumérés dans l'ordre croissant de leurs extrémités gauches.
- II.C.3 Utiliser l'inégalité trouvée à la question I.E.2.
- II.D Déterminer dans un premier temps la complexité des fonctions des questions de la partie I.D.2.
- III.B.2 Utiliser la fonction `est_clique` de la question I.E.3 ainsi que la fonction `voisins_inferieurs` de la question III.B.1.
- III.C Appliquer le résultat de la question II.C.2.
- III.D.3 Remarquer la ressemblance de cette partie avec la partie II.C.
- IV.A.2 Grâce à la question I.A.1, expliciter le fait que I_1 et I_2 ont une intersection non vide et que I_0 et I_2 ont une intersection vide.
- IV.A.3 Utiliser la caractérisation de la question I.A.1 pour contredire le fait que $I_0 \cap I_3 \neq \emptyset$.
- IV.B Procéder par récurrence, en remarquant qu'un cycle de longueur $n + 1 \geq 5$ dans un graphe d'intervalles peut se réduire en un cycle de longueur n dans un graphe d'intervalles obtenu en fusionnant deux sommets.
- IV.C Construire le graphe d'intervalles G résumant les informations données et noter qu'il n'est pas cordal. En sachant que seul le coupable a menti, la suppression d'un seul sommet de G le rend cordal.
- IV.D.1 Utiliser la fonction `est_clique` de la question I.E.3, puis montrer qu'elle teste si l'ensemble représenté par la liste `xs` de longueur k est une clique d'un graphe à m arêtes avec une complexité $O(1 + km)$.
- IV.D.3 Employer une approche gloutonne qui consiste à éliminer itérativement des sommets simpliciaux en utilisant la fonction `trouver_simplicial` de la question IV.D.2. Justifier que l'algorithme est correct en démontrant que si un graphe $G = (S, A)$ possède un ordre d'élimination parfait, alors pour tout sommet simplicial x dans G , le graphe induit par $S \setminus \{x\}$ possède un ordre d'élimination parfait.
- IV.E.1 En notant C_1 l'ensemble des sommets de C qui n'ont pas de voisin dans S_1 , montrer que $C \setminus C_1$ est une coupure. Dédire de la minimalité de C que $C_1 = \emptyset$.
- IV.E.2 Profiter du fait que G_1 et G_2 sont des composantes connexes de G .
- IV.E.3 Dans le cycle formé à partir des chemins P_1 et P_2 , où se trouvent les cordes ?
- IV.F.3.c Utiliser le résultat de la question IV.F.3.a pour appliquer l'hypothèse $\mathcal{P}(H_1)$. Conclure grâce à la question IV.E.4.
- IV.G S'inspirer du raisonnement de la question IV.D.3 et utiliser le résultat de la partie IV.F.

I. GRAPHE D'INTERVALLES

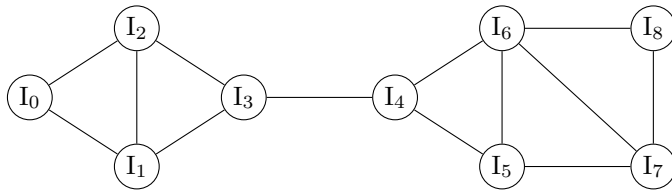
I.A.1 Soient $I = [a; b]$ et $J = [c; d]$ deux intervalles. S'ils sont en conflit, il existe un réel x appartenant à $I \cap J$. En particulier, $a \leq x \leq b$ et $c \leq x \leq d$, ce qui implique que $a \leq d$ et $c \leq b$.

Réciproquement, supposons que $a \leq d$ et $c \leq b$. Deux cas sont alors possibles : soit $a \leq c$, auquel cas $c \in I \cap J$, soit $a > c$, auquel cas $a \in I \cap J$. Dans tous les cas, les intervalles I et J sont en conflit.

Les intervalles $[a; b]$ et $[c; d]$ sont donc en conflit si et seulement si $a \leq d$ et $c \leq b$. Ainsi, la fonction `conflit` se contente d'exécuter ce test.

```
let conflit (a,b) (c,d) = (a <= d) && (c <= b);;
```

I.C.1 Le graphe d'intervalles associé au problème b de la figure 1 peut se représenter de la manière suivante :

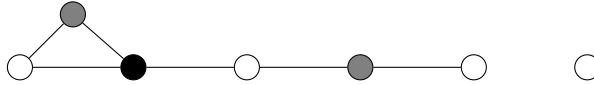


I.C.2 Afin de construire la représentation des arêtes du graphe $G(I_0, \dots, I_{n-1})$ associé au tableau de segments $(I_0, \dots, I_{n-1})$, parcourons les paires d'intervalles itérativement afin d'insérer une arête (i, j) si I_i et I_j sont en conflit. Un tableau `aretes` de n listes vides est initialement créé pour recevoir les listes de voisins du graphe. Une boucle sur les indices i parcourt ensuite les intervalles I_i et une boucle interne sur les indices $j \in \{0, \dots, i-1\}$ parcourt les intervalles I_j pour concaténer j en tête de la liste courante des voisins de i , et i en tête de la liste courante des voisins de j , si I_i et I_j sont en conflit : le test de conflit est réalisé à l'aide de la fonction `conflit` décrite à la question I.A.1.

```
let construit_graphe segments =
  let n = vect_length segments in
  let aretes = make_vect n [] in
  for i=n-1 downto 0 do
    for j=i-1 downto 0 do
      if (conflit segments.(i) segments.(j))
      then begin
        aretes.(i) <- j::aretes.(i);
        aretes.(j) <- i::aretes.(j)
      end
    done;
  done;
  aretes;;
```

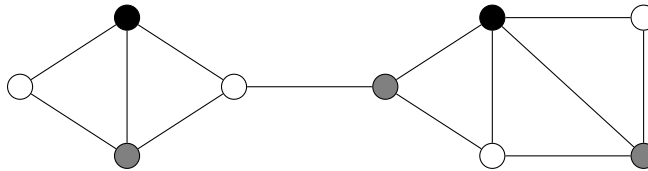
Par souci esthétique, les indices i et j sont parcourus dans l'ordre décroissant, afin de renvoyer une liste des voisins de chaque sommet dans l'ordre croissant.

I.D.1 Une coloration du graphe d'intervalles de la figure 3, associé au problème a de la figure 1, est donnée par la suite $(0, 1, 2, 0, 1, 0, 0)$:



Toute coloration de ce graphe nécessite au moins trois couleurs puisque les sommets 0, 1 et 2 sont reliés deux à deux et ne peuvent donc pas avoir une couleur en commun.

Pour les mêmes raisons, la suite $(0, 1, 2, 0, 1, 0, 2, 1, 0)$ est une coloration optimale du graphe d'intervalles trouvé en question I.C.1, associé au problème b de la figure 1, comme reproduite ci-dessous :



I.D.2.a Le test d'appartenance d'un élément x dans une liste est réalisé à l'aide d'une fonction récursive qui parcourt la liste jusqu'à trouver x et qui conclut négativement si la fin de la liste est atteinte.

```
let rec appartient l x =
  match l with
  | [] -> false
  | a::q -> (a=x) || (appartient q x);;
```

Il est sans doute maladroit dans cette question de répondre en utilisant la fonction de signature `mem : 'a -> 'a list -> bool` de la librairie standard de Caml telle que `mem x l` teste l'appartenance de l'élément x dans la liste l .

I.D.2.b Notons que le plus petit entier non présent dans une liste l de taille n est nécessairement inférieur à $n + 1$. Créons donc un vecteur v de taille n , initialisé à `false`, dont la case d'indice c a pour vocation de contenir `true` si et seulement si c appartient à l . À l'aide d'un tel vecteur, trouver le plus petit entier non présent revient à trouver l'indice de la première case du vecteur ne contenant pas `true`, donc $n + 1$ si toutes les cases contiennent `true`, ce que réalise la boucle `while` finale du programme. Construisons ainsi ce vecteur à l'aide d'une fonction auxiliaire récursive de signature `parcours : int list -> unit` telle que `parcours l` met à jour le vecteur v en fonction des entiers contenus dans l .

```
let plus_petit_absent l =
  let n = list_length l in
  let v = make_vect n false in
  let rec parcours = function
    | [] -> ()
    | c::q -> if (c < n) then v.(c) <- true; parcours q
  in parcours l;
  let c = ref 0 in
  while (!c < n && v.(!c)) do incr c done;
  !c;;
```

Mines Maths 1 MP 2015 — Corrigé

Ce corrigé est proposé par Gilbert et Florence Monna (Professeur en CPGE et Docteur en mathématiques) ; il a été relu par Pierre-Yves Bienvenu (ENS Ulm) et Nicolas Martin (Professeur agrégé).

Le problème est consacré à l'étude d'une équation différentielle du second ordre avec des conditions initiales, mais il s'agit d'un problème de Sturm-Liouville et non d'un problème usuel de Cauchy, ce qui donne une certaine originalité au sujet.

La méthode utilisée n'est pas classique non plus en classes préparatoires puisque l'on utilise des opérateurs de Volterra, qui sont introduits dans la partie A. On étudie alors un opérateur symétrique défini positif dans un espace vectoriel euclidien de dimension infinie en déterminant son spectre, et on relie les vecteurs propres de l'opérateur aux solutions d'une équation différentielle.

La partie B consiste en la démonstration du théorème de Weierstrass : toute fonction réelle définie et continue sur un segment est limite uniforme d'une suite de polynômes, en utilisant les très classiques polynômes de Bernstein. L'originalité vient de l'utilisation de l'inégalité de Bienaymé-Tchebychev pour établir la majoration technique sur les polynômes de Bernstein, qui permet d'achever la démonstration de la convergence uniforme par la découpe usuelle. Cette intervention des probabilités, nouvellement introduites au programme cette année, à un endroit où l'on ne les attendait pas, est très intéressante.

La partie suivante a pour but de déterminer un développement en série trigonométrique, appelée aussi série de Fourier. Les séries de Fourier ont disparu du programme, mais le sujet utilise une suite orthonormée totale dont l'existence est montrée par le théorème de Weierstrass trigonométrique qui se déduit du résultat de la partie précédente. La détermination effective de la série trigonométrique (question 11) nécessite tout de même quelques calculs...

Dans la dernière partie, dont certaines questions sont difficiles, on utilise les outils construits jusque-là pour étudier l'équation différentielle avec les conditions de Sturm-Liouville, en établissant une condition nécessaire et suffisante pour qu'une fonction soit solution. Ceci permet de construire une solution dans un cas (mais on n'aborde pas la question de son unicité), de trouver une infinité de solutions dans un autre cas, et on termine par un dernier cas où il n'y a pas de solution, ce qui démarque ce problème de Sturm-Liouville d'un problème de Cauchy.

En résumé, c'est un problème très intéressant qui utilise deux nouveautés du programme, l'inégalité de Bienaymé-Tchebychev et les suites orthonormées totales ainsi que de nombreux chapitres. Il ne peut être traité dans son ensemble qu'en fin de seconde année.

INDICATIONS

Partie A

- 1 Utilisez une intégration par parties.
- 2 Pour la symétrie, cherchez à utiliser la question précédente.
- 3 Pour montrer que la fonction f_λ est de classe $\mathcal{C}^2$, pensez qu'une primitive d'une fonction de classe $\mathcal{C}^k$ est de classe $\mathcal{C}^{k+1}$. Il suffit ensuite de dériver des fonctions qui sont définies comme des primitives.
- 4 Commencez par résoudre l'équation différentielle linéaire homogène du second ordre obtenue à la question précédente, puis utilisez les conditions initiales sur la solution générale.

Partie B

- 6 Appliquez l'inégalité de Bienaymé-Tchebychev à la variable aléatoire Z_n puis interprétez l'évènement $\{|Z_n - E(Z_n)| \geq \alpha\}$ comme une réunion disjointe d'évènements dont on exprime les probabilités.
- 7 Utilisez le théorème de transfert, puis observez que

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

soit par un argument probabiliste, soit par un argument algébrique. Utilisez ensuite le théorème de Heine pour déterminer un réel positif α et faites une découpe de la somme en séparant les indices tels que $\left| \frac{k}{n} - x \right| \geq \alpha$ et les autres.

Partie C

- 8 On peut procéder par récurrence.
- 9 Appliquez le théorème de Weierstrass à la fonction définie sur $[-1, 1]$ par

$$g(x) = f(\operatorname{Arccos}(x))$$
- 10 Utilisez une propriété des suites totales que vous avez vue en cours, puis la relation de comparaison entre les normes $\|\cdot\|_G$ et $\|\cdot\|_\infty$, pour conclure avec l'unicité de la limite pour la norme $\|\cdot\|_G$.
- 11 Les calculs sont un peu longs... Pour conclure, démontrez la convergence normale pour permuter la somme et l'intégrale, puis ramenez-vous à la question précédente.
- 12 On peut partir du membre de droite en séparant l'intégrale en trois, puis en faisant une intégration par parties en introduisant la primitive $V(f)$ de f .

Partie D

- 13 Utilisez la question 4 au lieu de refaire les calculs.
- 14 Pour une implication, dérivez. Pour l'autre, intégrez, mais en utilisant la bonne primitive et en tenant compte des conditions initiales. Pour la dernière égalité, utilisez la question 12.
- 15 Posez la fonction $g = \sum_{n=0}^{+\infty} \frac{1}{(2n+1)^2 - \lambda} \langle h, \phi_n \rangle \phi_n$ et montrez qu'elle vérifie la condition suffisante de la question 14.
- 16 Pour trouver une solution, modifiez celle utilisée à la question précédente et, pour montrer qu'il n'y a pas de solution, montrez que la condition nécessaire de la question 14 n'est vérifiée par aucune fonction.

A. OPÉRATEURS DE VOLTERRA

1 Par définition du produit scalaire sur E , on a, pour tout couple $(f, g) \in E^2$

$$\langle V(f), g \rangle = \int_0^{\pi/2} V(f)(x)g(x) \, dx$$

Intégrons par parties en posant

$$\begin{cases} u = V(f), & u' = f \\ v' = g, & v = -V^*(g) \end{cases}$$

puisque, ainsi que l'énoncé le faisait remarquer, $V(f)$ et $-V^*(g)$ sont des primitives de f et g , donc sont de classe $\mathcal{C}^1$. On obtient

$$\int_0^{\pi/2} V(f)(x)g(x) \, dx = [-V(f)(x)V^*(g)(x)]_0^{\pi/2} + \int_0^{\pi/2} f(x)V^*(g)(x) \, dx$$

La partie intégrée est nulle puisque $V(f)(0) = V^*(g)(\pi/2) = 0$. Il en résulte que

$$\boxed{\forall (f, g) \in E^2 \quad \langle V(f), g \rangle = \langle f, V^*(g) \rangle}$$

2 D'après la propriété de symétrie du produit scalaire, on a, pour tout $(f, g) \in E^2$,

$$\langle V^*(V(f)), g \rangle = \langle g, V^*(V(f)) \rangle$$

et d'après la question précédente,

$$\langle g, V^*(V(f)) \rangle = \langle V(g), V(f) \rangle = \langle V(f), V(g) \rangle = \langle f, V^*(V(g)) \rangle$$

Finalement, $\forall (f, g) \in E^2 \quad \langle V^* \circ V(f), g \rangle = \langle f, V^* \circ V(g) \rangle$

d'où

$$\boxed{V^* \circ V \text{ est un opérateur symétrique.}}$$

Soit f un élément de E ; d'après la question 1,

$$\langle V^* \circ V(f), f \rangle = \langle V(f), V(f) \rangle$$

qui est positif, par propriété du produit scalaire. De plus, $\langle V(f), V(f) \rangle = 0$ entraîne que $V(f) = 0$, toujours par propriété du produit scalaire, donc sa dérivée, qui est la fonction f , est nulle. On a démontré que $\langle V^* \circ V(f), f \rangle \geq 0$ et $\langle V^* \circ V(f), f \rangle = 0 \implies f = 0$. Ainsi,

$$f \neq 0 \implies \langle V^* \circ V(f), f \rangle > 0$$

On conclut que $\boxed{\text{L'opérateur symétrique } V^* \circ V \text{ est défini positif.}}$

Soit f un vecteur propre de l'opérateur $V^* \circ V$ et λ la valeur propre associée. Un vecteur propre n'étant pas nul, on a $f \neq 0$, ce qui entraîne

$$\langle V^* \circ V(f), f \rangle > 0$$

ainsi que

$$\langle V^* \circ V(f), f \rangle = \langle \lambda f, f \rangle = \lambda \langle f, f \rangle$$

On a donc $\lambda \langle f, f \rangle > 0$, ce qui implique $\lambda > 0$ puisque $\langle f, f \rangle > 0$. On en conclut que

$$\boxed{\text{Les valeurs propres de l'opérateur } V^* \circ V \text{ sont strictement positives.}}$$

| C'est toujours le cas pour un opérateur défini positif.

3 Par définition d'une valeur propre

$$V^* \circ V(f_\lambda) = \lambda f_\lambda$$

La fonction f_λ est de classe $\mathcal{C}^0$, la fonction $V(f_\lambda)$, qui en est une primitive, est de classe $\mathcal{C}^1$ et la fonction $-V^*(V(f_\lambda))$, qui est une primitive de la fonction $V(f_\lambda)$ de classe $\mathcal{C}^1$, est de classe $\mathcal{C}^2$. On en déduit que

$$\boxed{\text{La fonction } f_\lambda \text{ est de classe } \mathcal{C}^2.}$$

En dérivant l'égalité $\lambda f_\lambda = V^*(V(f_\lambda))$, on obtient

$$\lambda f_\lambda' = -V(f_\lambda)$$

puis, en dérivant une nouvelle fois, $\lambda f_\lambda'' = -f_\lambda$

D'après la question précédente, λ n'est pas nul, donc la fonction f_λ vérifie

$$\boxed{f_\lambda'' + \frac{1}{\lambda} f_\lambda = 0}$$

On a les égalités, pour tout x élément de $[0; \pi/2]$,

$$V^*(V(f_\lambda))(x) = \lambda f_\lambda(x) \quad \text{et} \quad V(f_\lambda)(x) = -\lambda f_\lambda'(x)$$

En donnant à x la valeur $\pi/2$ dans la première égalité, on arrive à

$$\lambda f_\lambda\left(\frac{\pi}{2}\right) = V^*(V(f_\lambda))\left(\frac{\pi}{2}\right) = 0$$

d'où $f_\lambda\left(\frac{\pi}{2}\right) = 0$ puisque λ n'est pas nul. En donnant à x la valeur 0 dans la deuxième égalité, on obtient

$$\lambda f_\lambda'(0) = -V(f_\lambda)(0) = 0$$

d'où

$$\boxed{f_\lambda\left(\frac{\pi}{2}\right) = 0 \quad \text{et} \quad f_\lambda'(0) = 0}$$

4 L'équation différentielle $y'' + (1/\lambda)y = 0$ du second ordre à coefficients constants a pour équation caractéristique $r^2 + 1/\lambda = 0$, de racines $r = \pm i/\sqrt{\lambda}$ ($\lambda > 0$). Une base de solutions de l'équation différentielle est formée des fonctions

$$x \mapsto \cos\left(\frac{x}{\sqrt{\lambda}}\right) \quad \text{et} \quad x \mapsto \sin\left(\frac{x}{\sqrt{\lambda}}\right)$$

La solution générale de l'équation est définie par

$$y(x) = \alpha \cos\left(\frac{x}{\sqrt{\lambda}}\right) + \mu \sin\left(\frac{x}{\sqrt{\lambda}}\right) \quad (\alpha, \mu) \in \mathbb{R}^2$$

En dérivant,

$$y'(x) = -\frac{\alpha}{\sqrt{\lambda}} \sin\left(\frac{x}{\sqrt{\lambda}}\right) + \frac{\mu}{\sqrt{\lambda}} \cos\left(\frac{x}{\sqrt{\lambda}}\right)$$

La condition initiale $y'(0) = 0$ donne donc $\mu = 0$, d'où

$$y(x) = \alpha \cos\left(\frac{x}{\sqrt{\lambda}}\right)$$

La condition initiale $y(\pi/2) = 0$ s'écrit alors

$$\cos\left(\frac{\pi}{2\sqrt{\lambda}}\right) = 0$$

Mines Maths 2 MP 2015 — Corrigé

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

Le sujet a pour objectif d'établir une « inégalité de concentration » de la forme

$$\forall a > 0 \quad \exists b > 0 \quad \mathbb{P}(\|M\|_{\text{op}} \geq a) \leq b$$

où M est une matrice aléatoire dont les coefficients sont mutuellement indépendants et « uniformément sous-gaussiens » et $\|\cdot\|_{\text{op}}$ est une norme sur $\mathcal{M}_n(\mathbb{R})$ qui donne la valeur absolue de la plus grande valeur propre pour les matrices symétriques. Ce cadre d'étude intervient, en le généralisant, dans la technique d'analyse en composantes principales : on manipule des matrices de corrélation d'un échantillon de données, dont on cherche les éléments propres pour décrire, décorréler, débruiter... les données. Une inégalité de concentration apporte par exemple un outil statistique pour quantifier la validité des résultats obtenus.

Le problème se compose de quatre parties ; les trois premières sont indépendantes et la dernière utilise les précédentes pour démontrer l'inégalité recherchée.

- La première partie, très classique, mélange topologie et algèbre linéaire. Pour $M \in \mathcal{M}_n(\mathbb{R})$, la *norme opérationnelle* $\|M\|_{\text{op}} = \max \{\|Mx\| \mid x \in S^{n-1}\}$ est introduite. On en établit quelques propriétés, on la compare aux normes usuelles sur $\mathcal{M}_n(\mathbb{R})$ et on calcule sa valeur sur les matrices symétriques.
- La deuxième partie introduit la notion de variable aléatoire sous-gaussienne. Elle utilise des notions élémentaires d'analyse et de probabilités discrètes pour en étudier quelques propriétés. Une inégalité de concentration est établie pour ce type de variables aléatoires.
- La troisième partie porte exclusivement sur la topologie de $\mathbb{R}^n$. Il s'agit de démontrer que l'on peut recouvrir toute partie compacte par une union finie de boules fermées de rayon prescrit. Il s'agit de la partie la plus délicate du sujet.
- La dernière partie montre l'inégalité de concentration recherchée. Elle nécessite d'avoir bien intégré le reste du sujet et d'avoir pris suffisamment de recul sur les résultats démontrés et admis.

Globalement, le sujet est plutôt bien guidé et abordable, hormis la troisième partie qui demande des initiatives – mais le sujet énonce précisément tous les résultats utiles dans la suite. Par ailleurs, cette épreuve nécessite de maîtriser beaucoup de notions pour espérer en venir à bout. Les thèmes abordés vont en effet des probabilités discrètes à la topologie en passant par l'algèbre linéaire. Il constitue un bon problème de révision.

INDICATIONS

Partie A

- 1 Penser au fait qu'on est en dimension finie puis remarquer que, pour une matrice $M \in \mathcal{M}_n(\mathbb{R})$, l'application $x \mapsto \|Mx\|$ est continue sur S^{n-1} .
- 2 Tirer parti du fait que pour tout x dans $\mathbb{R}^n$ non nul, le vecteur $x/\|x\|$ est unitaire.
- 3 Pour passer du cas diagonal au cas général, se souvenir qu'une matrice symétrique réelle est diagonalisable en base orthonormée, de sorte que la matrice de passage est une matrice orthogonale (qui conserve par conséquent la norme).
- 4 Remarquer que J_n est symétrique puis déterminer son rang. En déduire une première valeur propre et sa multiplicité. Conclure en calculant ensuite l'image par J_n du vecteur ${}^t(1, \dots, 1)$.
- 5 Écrire, pour tout $(i, j) \in \llbracket 1; n \rrbracket^2$, le coefficient $M_{i,j}$ comme un produit scalaire puis majorer $|M_{i,j}|$ grâce à l'inégalité de Cauchy-Schwarz.
- 6 Utiliser l'inégalité de Cauchy-Schwarz.
- 7 Se servir de la question 6 et se souvenir qu'une inégalité est une égalité si, et seulement si, toutes les inégalités écrites pour la démontrer sont des égalités. Pour dénombrer les matrices $M \in \Sigma_n$ vérifiant $\|M\|_{\text{op}} = n$, voir qu'une telle matrice est caractérisée par les coefficients de sa première colonne et les coefficients de proportionnalité entre la i^{e} colonne et la première colonne pour tout $i \in \llbracket 2; n \rrbracket$.

Partie B

- 8 Comparer pour $t \in \mathbb{R}$ les développements en séries entières de $\text{ch}(t)$ et $\exp(t^2/2)$.
- 9 Remarquer que, pour $t \in \mathbb{R}$ et $x \in [-1; 1]$, on a

$$tx = \frac{1+x}{2}t - \frac{1-x}{2}t \quad \text{et} \quad \frac{1+x}{2} + \frac{1-x}{2} = 1$$

- 10 Se servir des questions 9 et 8 pour la première partie de la question. Puis, voir que X/α vérifie les hypothèses de la première partie de la question.
- 11 Utiliser la relation fonctionnelle de l'exponentielle :

$$\forall (x, y) \in \mathbb{R}^2 \quad \exp(x+y) = \exp(x)\exp(y)$$

Se souvenir ensuite que pour deux variables aléatoires réelles indépendantes, on a $E(XY) = E(X)E(Y)$.

- 12 Utiliser l'inégalité de Markov. Pour la deuxième partie de la question, remarquer que si X est α -sous-gaussienne, il en est de même pour $-X$.
- 13 Revenir à la double inégalité définissant la partie entière et comparer les deux événements $(X \geq k)$ et $(\lfloor X \rfloor \geq k)$ pour tout $k \in \mathbb{N}^*$.
- 14 Écrire, pour $k \in \mathbb{N}^*$, l'événement $(\exp(\beta^2 X^2/2) \geq k)$ sous la forme d'un événement du type $(|X| \geq K)$ pour un certain $K > 0$ puis utiliser les questions 12 et 13.

Partie C

- 15 Remarquer qu'on peut traiter le cas où K est fini facilement. Dans le cas où K est infini, construire, en raisonnant par l'absurde comme le suggère l'énoncé, une suite $(x_n)_{n \in \mathbb{N}}$ de K vérifiant

$$\forall i \neq j \quad \|x_i - x_j\| > \frac{\varepsilon}{2}$$

Aboutir à une contradiction en utilisant le théorème de Bolzano-Weierstrass.

- 16 Considérer un sous-ensemble A fini de K tel que $K \subset \bigcup_{a \in A} B_{a, \frac{\varepsilon}{2}}$. Établir alors par l'absurde que toute boule $B_{a, \frac{\varepsilon}{2}}$ avec $a \in A$ contient au plus un point de Λ . Pour la seconde moitié de la question, raisonner à nouveau par l'absurde. Pour contredire la maximalité, exhiber alors un sous-ensemble Λ' vérifiant la même propriété mais de cardinal strictement plus grand.
- 17 Pour la deuxième partie de la question, appliquer la fonction volume μ à l'inclusion

$$\bigcup_{a \in \Lambda} B_{a, \frac{\varepsilon}{2}} \subset B_{0, 1 + \frac{\varepsilon}{2}}$$

- 18 Considérer la partie de $\mathbb{N}$ constituée de tous les $\text{Card}(\Lambda)$ pour Λ décrivant les sous-ensembles de K vérifiant

$$\forall i \neq j \in \Lambda^2 \quad \|x_i - x_j\| > \varepsilon$$

pour un ε judicieusement choisi. Utiliser alors les questions 16 et 17.

Partie D

- 19 Pour la deuxième partie de la question, établir que les variables aléatoires y_i indexées par $i \in \llbracket 1 ; n \rrbracket$ sont mutuellement indépendantes. Pour la dernière partie de la question, utiliser l'inégalité de Markov.
- 20 Concernant le début de la question, raisonner par l'absurde et obtenir une contradiction en tirant parti du fait qu'il existe $x \in S^{n-1}$ tel que l'on ait l'égalité

$$\|M^{(n)}x\| = \|M^{(n)}\|_{\text{op}}$$

Remarquer alors que la première partie de la question se traduit par l'inclusion

$$\left(\|M^{(n)}\|_{\text{op}} \geq 2r\sqrt{n} \right) \subset \bigcup_{a \in \Lambda_n} \left(\|M^{(n)}a\| \geq r\sqrt{n} \right)$$

A. NORME D'OPÉRATEUR D'UNE MATRICE

1 Soit $M \in \mathcal{M}_n(\mathbb{R})$. On note N l'application de $\mathbb{R}^n$ dans $\mathbb{R}$ définie par $N : x \mapsto \|x\|$. Par définition, on a alors $S^{n-1} = N^{-1}(\{1\})$. Or, le singleton $\{1\}$ est un fermé de $\mathbb{R}$ et N est continue (car 1-lipschitzienne). On en déduit que S^{n-1} est un fermé de $\mathbb{R}^n$. De plus, S^{n-1} est borné par définition. On peut donc en conclure que

La sphère unité S^{n-1} est un compact de $\mathbb{R}^n$.

L'application de $\mathbb{R}^n$ dans lui-même définie par $f : x \mapsto Mx$ est linéaire en dimension finie et par conséquent continue. Ainsi, l'application de $\mathbb{R}^n$ dans $\mathbb{R}$ définie par $x \mapsto \|Mx\|$ est continue en tant que composée d'applications continues. Sa restriction au compact S^{n-1} est donc bornée et atteint ses bornes, si bien que

La quantité $\|M\|_{\text{op}} = \max \{ \|Mx\| \mid x \in S^{n-1} \}$ est bien définie.

2 Commençons par établir que $\|\cdot\|_{\text{op}}$ est une norme sur $\mathcal{M}_n(\mathbb{R})$.

• **Homogénéité :** Soient $\lambda \in \mathbb{R}$ et $M \in \mathcal{M}_n(\mathbb{R})$. D'après la question 1, il existe $x \in S^{n-1}$ tel que $\|\lambda M\|_{\text{op}} = \|(\lambda M)x\|$. Par homogénéité de $\|\cdot\|$, il s'ensuit que $\|\lambda M\|_{\text{op}} = |\lambda| \|Mx\|$. Puis, par définition de $\|\cdot\|_{\text{op}}$ et puisque $|\lambda| \geq 0$, il vient

$$\|\lambda M\|_{\text{op}} \leq |\lambda| \|M\|_{\text{op}}$$

De même, la question 1 fournit l'existence d'un $y \in S^{n-1}$ tel que $\|M\|_{\text{op}} = \|My\|$. Ainsi, par homogénéité de $\|\cdot\|$, on a $|\lambda| \|M\|_{\text{op}} = \|(\lambda M)y\|$. Par définition de $\|\cdot\|_{\text{op}}$, cela entraîne l'inégalité

$$|\lambda| \|M\|_{\text{op}} \leq \|\lambda M\|_{\text{op}}$$

qui permet de conclure à l'égalité $\|\lambda M\|_{\text{op}} = |\lambda| \|M\|_{\text{op}}$.

• **Inégalité triangulaire :** Soient M et N dans $\mathcal{M}_n(\mathbb{R})$. Pour tout $x \in S^{n-1}$, en utilisant l'inégalité triangulaire pour $\|\cdot\|$ et la définition de $\|\cdot\|_{\text{op}}$, il vient

$$\|(M+N)x\| = \|Mx + Nx\| \leq \|Mx\| + \|Nx\| \leq \|M\|_{\text{op}} + \|N\|_{\text{op}}$$

En passant au maximum sur tous les $x \in S^{n-1}$, on aboutit à l'inégalité triangulaire

$$\|M+N\|_{\text{op}} \leq \|M\|_{\text{op}} + \|N\|_{\text{op}}$$

• **Caractère défini :** Soit M dans $\mathcal{M}_n(\mathbb{R})$ telle que $\|M\|_{\text{op}} = 0$. Par définition de la norme opérationnelle, on a donc pour tout $x \in S^{n-1}$, $\|Mx\| = 0$. Par caractère défini de $\|\cdot\|$, on en déduit que pour tout $x \in S^{n-1}$, $Mx = 0$. Or, quel que soit $z \in \mathbb{R}^n$ non nul, $z/\|z\| \in S^{n-1}$. D'où, $M(z/\|z\|) = 0$ puis par linéarité $Mz = 0$. Finalement, cela entraîne que $M = 0_{\mathcal{M}_n(\mathbb{R})}$.

L'application qui à $M \in \mathcal{M}_n(\mathbb{R})$ associe $\|M\|_{\text{op}}$ est une norme sur $\mathcal{M}_n(\mathbb{R})$.

Soient $M \in \mathcal{M}_n(\mathbb{R})$ et x, y dans $\mathbb{R}^n$. Lorsque $x = y$, l'inégalité demandée est évidente. Supposons donc $x \neq y$. Le vecteur $\frac{x-y}{\|x-y\|}$ est alors dans S^{n-1} . Ainsi, par définition de $\|\cdot\|_{\text{op}}$, on obtient

$$\left\| M \left(\frac{x-y}{\|x-y\|} \right) \right\| \leq \|M\|_{\text{op}}$$

Par linéarité de M et homogénéité de $\|\cdot\|$, cela se réécrit

$$\|M(x-y)\| \leq \|M\|_{\text{op}} \|x-y\|$$

d'où $\forall M \in \mathcal{M}_n(\mathbb{R}) \quad \forall (x, y) \in (\mathbb{R}^n)^2 \quad \|Mx - My\| \leq \|M\|_{\text{op}} \|x - y\|$

Mines Informatique commune MP 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`.

Mines Informatique optionnelle MP 2015 — Corrigé

Ce corrigé est proposé par Charles-Pierre Astolfi (ENS Cachan) ; il a été relu par Benjamin Monmege (ENS Cachan) et Guillaume Batog (Professeur en CPGE).

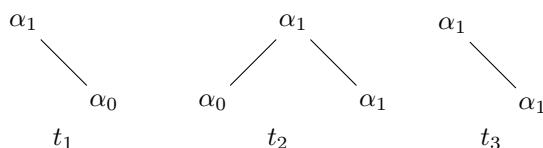
Ce sujet étudie deux généralisations des automates de mots dans le cas des arbres. Ces automates d'arbres sont par exemple utiles pour analyser une page HTML, qui peut être vue comme une structure arborescente de balises (une page qui contient plusieurs paragraphes, chacun contenant du texte ou des images...) ou pour calculer la valeur de vérité d'une formule de la logique propositionnelle. Le sujet comporte six parties mêlant programmation et théorie des automates.

- La première partie demande de coder des fonctions sur les listes. La complexité est systématiquement demandée.
- La deuxième partie introduit les arbres binaires étiquetés et fait programmer quelques fonctions pour les manipuler.
- La troisième partie définit les langages d'arbres, qui sont des généralisations aux arbres des langages de mots. Les questions sont relativement simples afin de permettre de prendre en main cette nouvelle notion.
- La quatrième partie introduit les automates d'arbres descendants déterministes, qui sont une première généralisation des automates de mots. On traite le cas de quelques automates particuliers, puis on code l'exécution d'un tel automate sur un arbre.
- La cinquième partie fait le lien entre les automates d'arbres descendants déterministes et les langages de mots.
- La sixième partie est de loin la plus longue (16 questions sur 32) et la plus difficile. Y sont introduits les automates d'arbres ascendants, qui sont une deuxième généralisation des automates de mots. Les questions reprennent la structure d'un cours sur les automates : déterminisation, clôture par union, complémentaire, intersection, etc. Les questions théoriques sont toutes suivies d'une question d'implémentation.

Ce sujet est d'une difficulté et d'une longueur peu communes pour le concours Mines-Ponts. Il est toujours possible d'avancer sans répondre à une question, mais à la condition d'avoir compris toutes les définitions des parties précédentes.

INDICATIONS

- 3 Appliquer récursivement la fonction **union**.
- 4 Utiliser une fonction auxiliaire qui à une variable **x** et une liste **l** renvoie une liste de couples ayant **x** comme premier élément et un élément de **l** comme second.
- 9 Remarquer que $|S| = 1 + |g(S_g)| + |d(S_d)|$ dans ce cas.
- 10 Définir un automate à trois états dont l'un, non final, marque le fait qu'un nœud possède un ancêtre fils droit.
- 11 Considérer les arbres suivants et montrer qu'un automate acceptant t_1 et t_2 accepte t_3 .



- 14 Construire une exécution de $\mathcal{A}^\downarrow$ sur $chaîne(x)$ à partir d'une exécution de $\mathcal{A}$ sur x et réciproquement. Vérifier ensuite que $\mathcal{L}(\mathcal{A}^\downarrow) \subset L_{chaîne}$.
- 15 Construire un automate de mots qui reconnaît L . La fonction de transition est presque la projection sur la première composante de la fonction de transition de l'automate d'arbres reconnaissant $chaîne(L)$.
- 16 Si un automate à k états reconnaît un mot à $2k$ états, il existe alors au moins un état qui est visité deux fois.
- 17 Remarquer qu'on reste dans l'état q_0 tant qu'aucune lettre α_1 n'est rencontrée.
- 18 Poser pour la fonction de transition de l'automate ascendant

$$\Delta(q_g, q_d, \alpha) = \{q \mid (q, \alpha) \in \delta^{-1}(q_g, q_d)\}$$

C'est l'ensemble des états menant à (q_g, q_d) en lisant α dans l'automate descendant où l'état initial est $\{q_0\}$ et l'ensemble des états finals est F .

- 23 Considérer un automate dont l'ensemble des états est $\mathcal{P}(Q)$, avec pour état initial, l'ensemble des états initiaux de l'automate de départ et comme états finals les parties de Q qui contiennent au moins un état final.
- 24 Utiliser le fait que les entiers de 0 à $2^n - 1$ sont en bijection avec leur représentation binaire sur n bits.
- 26 Prendre un automate déterministe et échanger les états finals et non finals.
- 28 Prendre l'union (disjointe) des deux automates.
- 30 Utiliser les questions 26 et 28.
- 32 Considérer l'ensemble des arbres dont le seul nœud binaire est à la racine puis suivi d'une branche à gauche avec uniquement des fils gauches, et une branche à droite avec uniquement des fils droits.

FONCTIONS UTILITAIRES

1 La liste `li` contient `x` si et seulement si `x` est le premier élément de `li` ou si la liste `li` privée de son premier élément contient `x`. Parcourons ainsi récursivement la liste `li` à la recherche de `x`.

```
let rec contient li x =
  match li with
  | [] -> false
  | t::q -> t = x || contient q x;;
```

La fonction examine chacun des éléments de la liste au plus une fois. En notant ℓ la liste codée par `li` et $|\ell|$ sa longueur, on en conclut que

La complexité de `contient` est $O(|\ell|)$.

La librairie standard de Caml contient la fonction `mem : 'a -> 'a list -> bool` qui réalise la même opération que `contient`, à l'ordre des arguments près. Cependant, il serait maladroit de l'appeler ici, l'énoncé demandant plutôt de recoder cette fonction.

2 On parcourt récursivement la liste `l1` pour en extraire tous les éléments qui sont dans `l1` sans être dans `l2` et lorsqu'on a parcouru entièrement `l1`, on adjoint la liste des éléments extraits à `l2`. Les éléments de cette liste sont donc des éléments qui sont dans `l1` ou dans `l2` et apparaissent exactement une fois, puisque les listes sont supposées sans doublon.

```
let rec union l1 l2 =
  match l1 with
  | [] -> l2
  | t::q -> if contient l2 t then union q l2 else t::(union q l2);;
```

Notons ℓ_1 (respectivement ℓ_2) la liste codée par `l1` (respectivement `l2`). La fonction réalise $|\ell_1|$ appels à `contient`, dont la complexité est ici $O(|\ell_2|)$. Ainsi,

La complexité de `union` est $O(|\ell_1||\ell_2|)$.

Comme pour la question précédente, la librairie standard de Caml contient la fonction `union : 'a list -> 'a list -> 'a list`, qui fait la même chose. À nouveau, il est clair que le sujet demande de la réimplémenter.

Une façon plus efficace de procéder aurait été de d'abord trier chacune des deux listes puis de les fusionner. En effet, fusionner deux listes triées peut s'effectuer en $O(|\ell_1| + |\ell_2|)$. Le tri s'effectuant en $O(|\ell_1| \log |\ell_1| + |\ell_2| \log |\ell_2|)$, la complexité totale est alors $O(|\ell_1| \log |\ell_1| + |\ell_2| \log |\ell_2|)$. Cependant, l'énoncé demande explicitement d'utiliser la fonction `contient` et n'oriente donc pas vers cette autre possibilité.

3 Il suffit d'appliquer récursivement la fonction `union` à chacune des listes de 1.

```
let rec fusion l =
  match l with
  | [] -> []
  | li::q -> union li (fusion q);;
```

Avec les notations de l'énoncé, le i -ème appel récursif à **fusion** réalise l'union de la liste ℓ_i avec une liste de longueur majorée par $\sum_{j=i+1}^k |\ell_j| \leq L$. D'après la question 2, cet appel a une complexité en $O(|\ell_i| L)$. En sommant sur i entre 1 et k ,

La complexité de **fusion** est en $O(L^2)$.

4 Définissons tout d'abord une fonction auxiliaire **couplage**: `'a -> 'b list -> ('a * 'b) list`. Étant donnés un élément x et une liste l , **couplage** x l renvoie une liste de couples dont le premier élément est x et le second un des éléments de l . En appliquant **couplage** à chaque élément de l_1 et sur la liste l_2 , on obtient une liste de listes dont la concaténation donne le produit cartésien.

```
let rec couplage x l2 =
  match l2 with
  | [] -> []
  | t::q -> (x,t)::(couplage x q);;
let rec concat l1 l2 = match l1 with
  | [] -> l2
  | t::q -> t::(concat q l2);;
let rec produit l1 l2 =
  match l1 with
  | [] -> []
  | x::q -> concat (couplage x l2) (produit q l2);;
```

En notant ℓ_1 et ℓ_2 les listes codées par l_1 et l_2 , la fonction **couplage** a une complexité en $O(|\ell_2|)$ et retourne une liste de taille ℓ_2 . Chacun des $|\ell_1|$ appels récursifs de **produit** coûte $O(|\ell_2|)$ (qui est le coût de l'appel à **couplage** et **concat**). Ainsi,

La complexité de **produit** est en $O(|\ell_1||\ell_2|)$.

ARBRES BINAIRES ÉTIQUETÉS

5 La fonction **arbre**: `int -> Arbre -> Arbre -> Arbre` fait utiliser la syntaxe de la déclaration d'enregistrements en Caml.

```
let arbre x ag ad = Noeud { etiquette=x; gauche=ag; droite=ad };;
```

6 On procède par récurrence sur l'arbre: si celui-ci est vide, sa taille est 0. Sinon, sa taille est égale à la somme des tailles de chacun de ses sous-arbres gauche et droit plus un pour sa racine.

```
let rec taille_arbre t =
  match t with
  | Vide -> 0
  | Noeud n -> 1 + taille_arbre n.gauche + taille_arbre n.droite;;
```

X/ENS Maths A MP 2015 — Corrigé

Ce corrigé est proposé par Sophie Rainero (Professeur en CPGE) ; il a été relu par Alexandre Le Meur (ENS Cachan) et Nicolas Martin (Professeur agrégé).

Ce sujet porte principalement sur les valeurs propres des matrices symétriques réelles. Il se compose de questions préliminaires et de quatre parties, chacune utilisant les résultats précédents.

- Dans les questions préliminaires, on commence par démontrer des résultats très classiques : l'ensemble $O_n(\mathbb{R})$ des matrices orthogonales de taille n est un sous-groupe du groupe $GL_n(\mathbb{R})$ des matrices inversibles et constitue une partie compacte de l'espace vectoriel normé $\mathcal{M}_n(\mathbb{R})$ des matrices réelles de taille n . On manipule ensuite des spectres ordonnés et « enlacés » : un $(n+1)$ -uplet $(\lambda_1 \geq \dots \geq \lambda_{n+1})$ et un n -uplet $(\mu_1 \geq \dots \geq \mu_n)$ sont dits enlacés lorsque

$$\lambda_1 \geq \mu_1 \geq \lambda_2 \geq \mu_2 \geq \dots \geq \lambda_n \geq \mu_n \geq \lambda_{n+1}$$

- Dans la première partie, on établit des résultats sur les racines de polynômes ; ils seront utilisés plus loin pour étudier les valeurs propres de matrices symétriques réelles. Cette partie ne fait appel qu'à des connaissances de première année, notamment l'arithmétique des polynômes et le théorème des valeurs intermédiaires.
- Dans la deuxième partie, on commence par prouver un résultat sur les déterminants de matrices par blocs ; puis on s'en sert pour relier le polynôme caractéristique d'une matrice symétrique réelle M de taille $n+1$ et celui de sa sous-matrice A , notée aussi $M_{\leq n}$, obtenue en enlevant la dernière ligne et la dernière colonne. On démontre ensuite que les spectres de M et de A sont enlacés. On prouve enfin, et c'est sans doute le passage le plus difficile du sujet, que l'ensemble des spectres des matrices extraites $(UMU^{-1})_{\leq n}$ quand U parcourt $O_{n+1}(\mathbb{R})$ est exactement l'ensemble des n -uplets $\hat{\mu}$ tels que $\text{Sp}(M)$ et $\hat{\mu}$ sont enlacés.
- La troisième partie, assez courte, établit des résultats sur les diagonales de matrices semblables à une matrice symétrique donnée, avec une matrice de passage orthogonale.
- Enfin, dans la quatrième partie, qui commence par des questions très simples, on donne des illustrations géométriques des résultats obtenus précédemment.

Ce sujet est assez long ; il comporte de nombreuses questions abordables mais aussi quelques-unes qui sont délicates ou longues. À l'exception de la première partie, qui peut être traitée dès la MPSI, il nécessite d'avoir étudié les chapitres d'algèbre euclidienne et de réduction de deuxième année.

INDICATIONS

- 1.b Penser à utiliser la dimension finie de $\mathcal{M}_n(\mathbb{R})$. Pour montrer que $O_n(\mathbb{R})$ est fermé, on pourra introduire la fonction $\mathcal{M}_n(\mathbb{R}) \rightarrow \mathcal{M}_n(\mathbb{R}), A \mapsto {}^t A A$.
- 3 Examiner attentivement tous les cas possibles pour la position de x dans $\hat{\lambda}$ et $\hat{\mu}$.
- 5.a Quel est le degré de $X Q_0 - P$?
- 5.b Appliquer plusieurs fois le théorème des valeurs intermédiaires à P en utilisant le résultat de la question 4.b.
- 5.c Pour tout $j \in \llbracket 1 ; n \rrbracket$, calculer $P(\mu_j)$ de deux manières.
- 7.a Commencer par écrire que $P \wedge Q_1$ est un diviseur de Q_1 , puis caractériser l'ordre des racines à l'aide de l'expression de P , en introduisant l'ensemble J défini dans l'énoncé.
- 7.b En notant $N + 1$ le degré de P , il suffit de trouver $N + 1$ racines réelles de P (comptées avec leurs multiplicités). Utiliser le résultat de la question 5.b après avoir factorisé le polynôme P à l'aide de la question 7.a.
- 8 On peut procéder par analyse-synthèse.
- 9.a Appliquer le théorème de réduction des matrices symétriques réelles à la matrice A .
- 9.b Utiliser le résultat admis dans la question 8.
- 9.c Introduire les racines distinctes de Q_0 pour ensuite appliquer la question 7.b.
- 10 Cette question est difficile. Commencer par montrer que $\mathcal{C}_M$ est borné à l'aide de la question 9. Pour prouver que $\mathcal{C}_M$ est fermé, procéder par étapes : prouver que
- $$\mathcal{E}_M = \{ (U M U^{-1})_{\leq n} \mid U \in O_{n+1}(\mathbb{R}) \}$$
- est compact, puis que l'ensemble des polynômes caractéristiques des matrices de $\mathcal{E}_M$ est compact et enfin utiliser la caractérisation séquentielle des fermés.
- 11.a Faire appel aux questions 5.c, 9.b et 2.
- 11.b On peut montrer que si $\text{Sp}(M)$ et $\hat{\mu}$ sont enlacés, alors $\hat{\mu}$ est la limite d'une suite $(\hat{\mu}_p)_{p \in \mathbb{N}}$ telle que, pour tout $p \in \mathbb{N}$, $\text{Sp}(M)$ et $\hat{\mu}_p$ sont strictement enlacés.
- 12 Justifier que $\mathcal{D}_M = \mathcal{D}_{\Delta(\mu_1, \mu_2)}$ puis décrire les éléments de $O_2(\mathbb{R})$ pour calculer les diagonales de toutes les matrices de la forme $U \Delta(\mu_1, \mu_2) U^{-1}$ lorsque U parcourt $O_2(\mathbb{R})$.
- 13.c Dans l'hérédité, utiliser les questions 9.c et 13.a.
- 15.a Justifier que $\varphi(H^+)$ est l'ensemble des points à coordonnées positives dans le repère $(O; \omega_1, \omega_2)$.
- 15.c Montrer que $\varphi(\mathcal{Q}_{\hat{\lambda}})$ est définie par quatre inégalités portant sur les coordonnées dans le repère $(O; \omega_1, \omega_2)$. Faire un dessin.
- 16.a Pour la réciproque, appliquer la première implication avec σ^{-1} .
- 16.c Suivre l'indication de l'énoncé pour la première partie de la question. Pour la seconde partie, utiliser la question 16.a et la première partie de la question.
- 16.d Faire appel à la question 16.c pour montrer que tout élément de $\mathcal{Q}_{\hat{\lambda}}$ est dans $\mathcal{D}_M$.
- 16.e Décomposer H comme la réunion de H^+ et d'ensembles obtenus à partir de H^+ par permutations des éléments. En déduire que $\varphi(\mathcal{D}_M)$ est la réunion de 6 quadrilatères, $\varphi(\mathcal{Q}_{\hat{\lambda}})$ et ses images par s_1, s_2 et leurs composées.

Notons une toute petite erreur dans l'introduction de l'énoncé : la notation $\mathcal{M}_{r,n}(\mathbb{R})$ désigne bien sûr l'espace vectoriel des matrices à r lignes et n colonnes à coefficients réels et non pas à coefficients complexes.

QUESTIONS PRÉLIMINAIRES

1.a Démontrons à l'aide de la caractérisation des sous-groupes que $O_n(\mathbb{R})$ est un sous-groupe du groupe $GL_n(\mathbb{R})$ des matrices inversibles.

- Pour toute matrice M de $O_n(\mathbb{R})$, ${}^t M M = I_n$ donc M est inversible (d'inverse sa transposée) et ainsi $O_n(\mathbb{R})$ est inclus dans $GL_n(\mathbb{R})$.
- La matrice identité I_n vérifie ${}^t I_n I_n = I_n^2 = I_n$ donc c'est un élément de $O_n(\mathbb{R})$. L'ensemble $O_n(\mathbb{R})$ est non vide.
- Soit $(A, B) \in O_n(\mathbb{R})^2$. Posons $C = A B$. Alors

$${}^t C C = {}^t (A B) A B = {}^t B {}^t A A B = {}^t B I_n B = {}^t B B = I_n$$

puisque A et B sont dans $O_n(\mathbb{R})$. Donc $C \in O_n(\mathbb{R})$.

- Soit $A \in O_n(\mathbb{R})$. Alors $A \in GL_n(\mathbb{R})$ et $A^{-1} = {}^t A$ donc

$${}^t (A^{-1}) A^{-1} = A A^{-1} = I_n$$

Ceci prouve que A^{-1} appartient à $O_n(\mathbb{R})$.

On en conclut grâce au théorème de caractérisation des sous-groupes que

$$O_n(\mathbb{R}) \text{ est un sous-groupe du groupe } GL_n(\mathbb{R}).$$

1.b $O_n(\mathbb{R})$ est une partie de $\mathcal{M}_n(\mathbb{R})$, qui est un espace vectoriel de dimension finie, donc on peut munir $\mathcal{M}_n(\mathbb{R})$ de n'importe quelle norme (elles sont toutes équivalentes) et démontrer que $O_n(\mathbb{R})$ est une partie compacte revient à démontrer qu'il s'agit d'une partie fermée et bornée de $\mathcal{M}_n(\mathbb{R})$.

Soit N la norme euclidienne sur $\mathcal{M}_n(\mathbb{R})$, associée au produit scalaire

$$\begin{aligned} \varphi: \begin{cases} \mathcal{M}_n(\mathbb{R})^2 \longrightarrow \mathbb{R} \\ (A, B) \longmapsto \text{Tr} \left({}^t A B \right) \end{cases} \\ \text{Autrement dit, } N: \begin{cases} \mathcal{M}_n(\mathbb{R}) \longrightarrow \mathbb{R}_+ \\ A \longmapsto \sqrt{\text{Tr} \left({}^t A A \right)} \end{cases} \end{aligned}$$

Pour toute matrice $M \in O_n(\mathbb{R})$, ${}^t M M = I_n$ donc

$$N(M) = \sqrt{\text{Tr} \left({}^t M M \right)} = \sqrt{\text{Tr} (I_n)} = \sqrt{n}$$

Il en découle que $O_n(\mathbb{R})$ est une partie bornée de l'espace vectoriel normé $(\mathcal{M}_n(\mathbb{R}), N)$.

Introduisons la fonction $f : M \mapsto {}^t M M$. Remarquons que $O_n(\mathbb{R})$ est l'image réciproque du singleton $\{I_n\}$ par f , c'est-à-dire $O_n(\mathbb{R}) = f^{-1}(\{I_n\})$. La fonction $M \mapsto {}^t M$ est linéaire en dimension finie, et par conséquent continue. De plus, la fonction $(M, N) \mapsto MN$ est également continue puisque bilinéaire en dimension finie. Par composition, la fonction f est continue. Le singleton $\{I_n\}$ étant un fermé de $\mathcal{M}_n(\mathbb{R})$, son image réciproque $O_n(\mathbb{R})$ par la fonction continue f est un fermé de $\mathcal{M}_n(\mathbb{R})$.

Comme $O_n(\mathbb{R})$ est une partie fermée et bornée de l'espace vectoriel normé de dimension finie $\mathcal{M}_n(\mathbb{R})$, on peut conclure que

$$\boxed{O_n(\mathbb{R}) \text{ est une partie compacte de } \mathcal{M}_n(\mathbb{R}).}$$

| Cette question est très classique, il faut savoir la traiter rapidement.

2 Soient M et N dans l'ensemble $S_n(\mathbb{R})$ des matrices symétriques réelles. Démontrons par double implication qu'elles ont le même polynôme caractéristique si, et seulement si, il existe $U \in O_n(\mathbb{R})$ tel que $N = U M U^{-1}$.

- On suppose qu'il existe $U \in O_n(\mathbb{R})$ tel que $N = U M U^{-1}$. Alors N et M sont semblables et, d'après le cours, elles ont alors le même polynôme caractéristique.
- Réciproquement, supposons que M et N ont le même polynôme caractéristique. Les matrices M et N sont symétriques réelles donc, en notant $(\lambda_1, \lambda_2, \dots, \lambda_n)$ et $(\mu_1, \mu_2, \dots, \mu_n)$ leurs spectres ordonnés, il existe $(P, Q) \in O_n(\mathbb{R})^2$ tel que

$$M = P \Delta(\lambda_1, \lambda_2, \dots, \lambda_n) P^{-1} \quad \text{et} \quad N = Q \Delta(\mu_1, \mu_2, \dots, \mu_n) Q^{-1}$$

$$\text{où} \quad \Delta(\lambda_1, \lambda_2, \dots, \lambda_n) = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & \lambda_n \end{pmatrix}$$

$$\text{et} \quad \Delta(\mu_1, \mu_2, \dots, \mu_n) = \begin{pmatrix} \mu_1 & 0 & \dots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & \mu_n \end{pmatrix}$$

Par hypothèse, $\chi_M = \chi_N$. De plus, M et $\Delta(\lambda_1, \lambda_2, \dots, \lambda_n)$ sont semblables donc $\chi_M = \chi_{\Delta(\lambda_1, \lambda_2, \dots, \lambda_n)}$. De même, les matrices N et $\Delta(\mu_1, \mu_2, \dots, \mu_n)$ ont le même polynôme caractéristique. Ainsi, les deux matrices diagonales $\Delta(\lambda_1, \lambda_2, \dots, \lambda_n)$ et $\Delta(\mu_1, \mu_2, \dots, \mu_n)$ ont le même polynôme caractéristique. Comme on a ordonné les valeurs propres, elles sont donc égales. Posons alors $U = Q P^{-1}$.

$$\begin{aligned} N &= Q \Delta(\mu_1, \mu_2, \dots, \mu_n) Q^{-1} \\ &= Q (P^{-1} M P) Q^{-1} \\ &= (Q P^{-1}) M (Q P^{-1})^{-1} \\ N &= U M U^{-1} \end{aligned}$$

Les matrices P et Q sont orthogonales et on a démontré que $O_n(\mathbb{R})$ est un groupe, donc U est également un élément de $O_n(\mathbb{R})$.

Finalement, pour toutes matrices symétriques réelles M et N ,

$$\boxed{\chi_M = \chi_N \quad \Longleftrightarrow \quad \exists U \in O_n(\mathbb{R}) \quad N = U M U^{-1}}$$

X Maths B MP 2015 — Corrigé

Ce corrigé est proposé par Pierre-Yves Bienvenu (ENS Ulm, doctorant) ; il a été relu par Sophie Rainero (Professeur en CPGE) et Benjamin Monmege (Enseignant-chercheur à l'université).

Ce sujet mobilise tout le programme d'analyse, tout particulièrement les intégrales dépendant d'un paramètre, les développements limités et le calcul différentiel.

- La première partie recherche un développement asymptotique en $+\infty$ de la fonction Gamma d'Euler définie par

$$\Gamma(y) = \int_0^{+\infty} e^{-t} t^{y-1} dt$$

- La deuxième partie recherche, en certains points, des approximations convaincantes de la fonction F définie sur $\mathbb{R}_+^*$ par

$$F(x) = \int_1^{+\infty} e^{-t/x} t^{-1} dt$$

Elle utilise pour cela une série divergente.

- La troisième partie démontre que dans l'espace des fonctions de $\mathbb{R}^2$ dans $\mathbb{R}$ continues sur $\mathbb{R}^2$ et 1-périodiques en chacune de leurs variables, muni de la norme infinie, le sous-espace des polynômes trigonométriques, engendré par les fonctions de la forme $(x_1, x_2) \mapsto e^{2i\pi(k_1 x_1 + k_2 x_2)}$ avec $(k_1, k_2) \in \mathbb{Z}^2$, est dense.
- La quatrième partie s'intéresse à un problème différentiel non linéaire impliquant des fonctions 1-périodiques en deux variables. On le résout d'abord dans le cas où un terme est nul, puis on réintroduit ce terme comme une perturbation afin de produire une approximation valable de la solution correspondante.

Ce sujet est inhabituel pour ce concours dans la mesure où les parties sont presque indépendantes ; seule une question de la quatrième partie nécessite un résultat de la troisième. L'avantage est que cela permet de ne pas rester bloqué, d'autant que les questions sont bien guidées.

INDICATIONS

Première partie

1a Intégrer par parties.

2a Appliquer soigneusement le théorème de convergence dominée et la caractérisation séquentielle de la limite. Pour obtenir le résultat avec ρ_N , écrire

$$\rho_N(t) = f(t) - \sum_{k=0}^N a_k t^{(k+\lambda-\mu)/\mu}$$

et analyser les deux termes séparément.

2c Décomposer l'intégrale sur $]0; +\infty[$ en une intégrale sur $]0; \delta]$ et une intégrale sur $[\delta; +\infty[$ pour un $\delta > 0$ convenable. Utiliser la question 2b pour traiter la première et la question 2a pour la seconde intégrale.

3c Calculer explicitement les premiers termes de la série entière $\phi(\sum b_k q^k) = q^2$ et identifier terme à terme les deux membres pour obtenir les premiers coefficients b_i et c_i . Ensuite, appliquer ϕ_+^{-1} et ϕ_-^{-1} à l'égalité

$$q = \phi\left(\sum_{k=1}^{+\infty} b_k q^{k/2}\right)$$

3d Partir du résultat de la question 1b, décomposer l'intégrale sur $] -1; +\infty[$ en une intégrale sur $] -1; 0[$ et une intégrale sur $]0; +\infty[$, puis procéder sur chacune à un changement de variable judicieux.

3e Appliquer le résultat final de la question 2 et le fait que $\Gamma(1/2) = \sqrt{\pi}$.

Deuxième partie

5 Raisonner par récurrence et intégrer par parties.

6c Utiliser la question précédente.

7a Remarquer que $S_N(x)$ est une somme de termes dont les signes alternent et dont les valeurs absolues décroissent.

Troisième partie

8 Pour une fonction de la forme $(\theta_1, \theta_2) \mapsto f(\theta_1)g(\theta_2)$ avec f et g continues et périodiques, fournir une approximation en multipliant deux polynômes approximant respectivement f et g .

10b Montrer que $\psi_{j,\ell_1}(\theta_1) = 0$ pour $\theta \in]0; j-1] \setminus \{k_1, k_1+1\}$, et un résultat analogue sur θ_2 . Ensuite, partitionner intelligemment le domaine $[0; 1]^2$ et utiliser le théorème de Heine pour montrer que $\|S_j(f) - f\|_\infty$ est petit pour j grand.

11 Combiner les questions 8 et 10b.

Quatrième partie

15a Intégrer sur le carré unité les deux membres de l'égalité différentielle satisfaite par h pour trouver ν . Pour exhiber une solution, utiliser le fait que ω est non résonnant et le principe de superposition (c'est-à-dire la linéarité).

15c Intégrer le résultat de la question 15b et employer la relation fixée par l'énoncé entre α et $\tilde{\alpha}$. Pour borner les termes d'erreur, utiliser l'inégalité de la moyenne et le théorème de Heine (pour traiter le terme $h(\alpha(t)) - h(\omega t)$).

PREMIÈRE PARTIE

1a Fixons $y > 0$. La fonction $\gamma : t \mapsto e^{-t}t^{y-1}$ est définie et continue sur $]0; +\infty[$. En outre, elle est positive.

Notons que l'inégalité $e^{-t}t^{y-1} \leq t^{y-1}$, valable pour tout $t \geq 0$, et l'intégrabilité de la fonction $t \mapsto t^{y-1}$ sur $]0; 1]$ assurent, par comparaison de fonctions positives, que γ est intégrable au voisinage de 0.

Au voisinage de $+\infty$, observons que $t^{y-1} = O(e^{t/2})$ par croissances comparées. Par conséquent, $e^{-t}t^{y-1} = O(e^{-t/2})$. Or, $t \mapsto e^{-t/2}$ est intégrable sur $[1; +\infty[$. Par comparaison, γ est intégrable au voisinage de $+\infty$.

Finalement, l'intégrale $\int_0^{+\infty} e^{-t}t^{y-1} dt$ converge, autrement dit

Pour tout $y > 0$, $\Gamma(y)$ est bien définie.

Pour obtenir la relation demandée, introduisons $\varepsilon > 0$ et $T > 0$ pour faire une intégration par parties sur $[\varepsilon; T]$, ainsi que les fonctions $u : t \mapsto t^y$ et $v : t \mapsto -e^{-t}$ de classe $\mathcal{C}^1$ sur $\mathbb{R}_+^*$. Alors

$$\begin{aligned} \int_{\varepsilon}^T e^{-t}t^y dt &= \int_{\varepsilon}^T u(t)v'(t) dt \\ &= [u(t)v(t)]_{\varepsilon}^T - \int_{\varepsilon}^T u'(t)v(t) dt \\ \int_{\varepsilon}^T e^{-t}t^y dt &= \left[-e^{-t}t^y\right]_{\varepsilon}^T + \int_{\varepsilon}^T e^{-t}yt^{y-1} dt \end{aligned}$$

Or, $e^{-t}t^y$ tend vers 0 quand y tend vers 0 et quand y tend vers $+\infty$. Donc en faisant tendre ε vers 0 et T vers $+\infty$, on fait tendre le crochet vers 0 et les deux intégrales vers $\Gamma(y+1)$ à gauche et $y\Gamma(y)$ à droite, de sorte que

$$\forall y > 0 \quad \Gamma(y+1) = y\Gamma(y)$$

Enfin, la valeur en $y = 1$ étant clairement

$$\Gamma(1) = \int_0^{+\infty} e^{-t} dt = 1$$

une récurrence évidente donne le résultat final.

$$\forall n \in \mathbb{N} \quad \Gamma(n+1) = n!$$

À titre culturel, signalons que la fonction Γ est en fait l'unique prolongement logarithmiquement convexe de la factorielle à $]0; +\infty[$, c'est-à-dire que Γ est la seule fonction dont le logarithme est convexe qui coïncide avec la factorielle sur les entiers : c'est le théorème de Bohr-Mollerup.

1b Pour tout $y > 0$, on a vu que $\Gamma(y) = y^{-1}\Gamma(y+1)$ ce qui se reformule en

$$\Gamma(y) = y^{-1} \int_0^{+\infty} e^{-t}t^y dt$$

Écrivons ensuite $t^y = e^{y \ln t}$ et opérons le changement de variable affine $t = y(s+1)$:

$$\begin{aligned}
\Gamma(y) &= y^{-1} \int_0^{+\infty} e^{-t} t^y \, dt \\
&= y^{-1} \int_0^{+\infty} e^{-t} e^{y \ln t} \, dt \\
&= y^{-1} \int_{-1}^{+\infty} e^{-y(s+1)} e^{y \ln(y(s+1))} y \, ds \\
\Gamma(y) &= e^{-y} e^{y \ln y} \int_{-1}^{+\infty} e^{-y(s - \ln(1+s))} \, ds
\end{aligned}$$

On reconnaît en exposant la fonction ϕ introduite par l'énoncé, de sorte que

$$\Gamma(y) = e^{-y} y^y \int_{-1}^{+\infty} e^{-y\phi(s)} \, ds$$

2a Fixons $x > 0$. La fonction $t \mapsto e^{-t/x} t^\alpha$ est définie et continue sur $[\delta; +\infty[$. En outre, $e^{-t/x} t^\alpha = O(t^{-2})$ par croissances comparées et la fonction de Riemann $t \mapsto t^{-2}$ est intégrable au voisinage de $+\infty$. Par comparaison,

$$\text{La fonction } t \mapsto e^{-t/x} t^\alpha \text{ est intégrable sur } [\delta; +\infty[.$$

Fixons à présent $n \in \mathbb{N}$; on peut en fait supposer $n > 1 - \alpha$, puisque pour k entier $x^{k+1} = O(x^k)$. Il faut prouver que l'expression

$$x^{-n} \int_\delta^{+\infty} e^{-t/x} t^\alpha \, dt = \int_\delta^{+\infty} e^{-(t/x + n \ln x)} t^\alpha \, dt$$

converge vers 0 lorsque x tend vers 0^+ . Pour cela, invoquons la caractérisation séquentielle de la limite et le théorème de convergence dominée. Soit donc $(x_k)_{k \in \mathbb{N}}$ une suite de réels strictement positifs qui converge vers 0. Notons $f_k(t) = e^{-(t/x_k + n \ln x_k)} t^\alpha$ l'intégrande, pour $k \in \mathbb{N}, t \geq \delta$. Rassemblons les hypothèses.

- Pour tout $k \in \mathbb{N}$, la fonction $t \mapsto f_k(t)$ est continue par morceaux et intégrable. Accessoirement, elle est positive.
- Pour tout $t \geq \delta$, $\lim_{k \rightarrow +\infty} f_k(t) = 0$. C'est une conséquence de la formule de croissance comparée $e^{-t/x} = o(x^n)$ quand x tend vers 0.
- Il faut majorer f_k par une fonction g intégrable indépendante de k (hypothèse de domination). Pour ce faire, fixons $t \geq \delta$ et notons que la fonction $\phi_t: x \mapsto e^{-(t/x + n \ln x)}$ est dérivable sur $\mathbb{R}_+^*$, de dérivée

$$\phi'_t(x) = \frac{t}{x^2} e^{-t/x} x^{-n} - n x^{-(n+1)} e^{-t/x} = e^{-t/x} x^{-(n+2)} (t - nx)$$

qui est positive pour $x \leq t/n$, et négative pour $x \geq t/n$, de sorte que ϕ_t atteint son maximum en $x = t/n$. Pour $f_k(t)$, ceci implique que

$$\forall t \geq \delta \quad \forall k \in \mathbb{N} \quad 0 \leq f_k(t) \leq e^{-n} \left(\frac{t}{n} \right)^{-n} t^\alpha$$

ce qui fournit effectivement la domination recherchée, puisque $t \mapsto t^{\alpha-n}$ est intégrable sur $[\delta; +\infty[$, grâce à l'hypothèse $n > 1 - \alpha$.

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 $p_j \preceq p_j$.

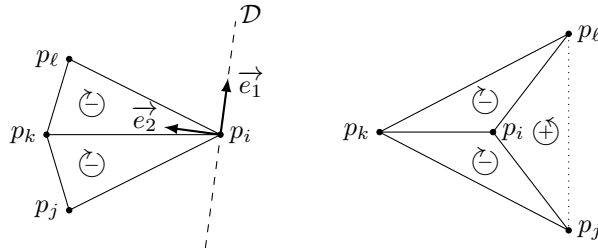
Pour l'antisymétrie, supposons fixés j et k tels que $p_j, p_k \in P \setminus \{p_i\}$, $p_j \preceq p_k$ et $p_k \preceq p_j$. Par définition,

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

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 fausse 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 $p_j \preceq p_\ell$.