
Vers le supérieur

Version prépa PCSI

AYOUB HAJLAOUI



Avant-propos

Ce document comporte une vingtaine de problèmes mathématiques corrigés, tous extraits de sujets de concours de fin de CPGE (concours d'entrée aux grandes écoles), en filière PC. Ils sont remaniés de sorte que le programme de Terminale (spécialité et maths expertes) est le seul prérequis pour en venir à bout. À cette fin, des questions intermédiaires ont parfois été ajoutées au sujet original. Les notions, définitions ou résultats qui pourraient manquer à l'élève pour parvenir à la résolution lui sont donnés en hypothèse, ainsi que des indications supplémentaires et remarques utiles.

La première originalité de ce document réside donc dans sa nature de pont entre le rivage du lycée et celui des concours. En se concentrant principalement sur des notions vues par l'élève l'an passé (suites, fonctions, intégrales, entre autres), il lui propose de les mobiliser pour faire tomber des questions de sujets de concours, montrant ainsi de manière pratique l'intérêt de bien maîtriser ces notions, de les voir comme des alliées de taille qui l'épauleront tout au long de la prépa, plutôt que de présenter leur révision comme une corvée estivale. En outre, de par les nouveautés introduites ça et là pour affiner ces dernières, de par un grand nombre de questions abordant des notions générales de raisonnement et d'ensembles, ce recueil donne au lycéen un avant-goût de la glorieuse chevauchée qui l'attend, en lui mettant le pied à l'étrier. Il lui rappellera certaines de ses errances de l'an passé, tout en le préparant à l'année à venir - sans prétendre constituer une liste exhaustive de ce qui y sera vu.

Si ces exercices ne nécessitent pas d'autre prérequis que le programme de Terminale, ils s'avéreront, en pratique, difficiles pour un élève au sortir du lycée. Le degré d'atsuce nécessaire, les idées qu'il faut avoir pour faire tomber telle ou telle question, détonnent avec la plupart des exercices rencontrés en Terminale, dont les pistes étaient en général plus claires, et les résolutions plus « téléphonées ».

C'est ici qu'intervient la seconde originalité de ce document par rapport à d'autres recueils d'exercices corrigés, originalité qui constitue son atout majeur : la correction, très détaillée, insiste autant que possible sur « le pourquoi de l'idée », la question de l'apparition de la première étincelle : à quoi telle situation nous fait-elle penser ? pourquoi est-il judicieux d'emprunter telle voie, de penser à telle astuce à ce moment-là plutôt qu'à un autre ? quel écueil faut-il éviter et comment voir que c'est un piège ? Vous trouverez de telles considérations en italique, en parallèle de la correction à proprement parler. Autant de didascalies rythmant la pièce de théâtre mathématique aux premières loges de laquelle vous êtes convié. Lever de rideau.

L'auteur en quelques mots

Lauréat de l'agrégation externe de mathématiques en 2020 (64^{ème} sur 323 admis et 3069 inscrits), docteur en mathématiques appliquées (université Paris VI), diplômé de l'école d'ingénieurs des Mines de Nancy et du master recherche MVA (Maths Vision Apprentissage) de l'ENS Cachan, je donne des cours de mathématiques (particuliers et en groupe, niveau lycée à prépa/L3) depuis 2008. Je suis également colleur en MPSI au lycée Charlemagne (Paris), lycée où j'ai moi-même effectué mes années de prépa MPSI/MP.

Parallèlement, je rédige des exercices de mathématiques (principalement niveau prépa et Terminale) ainsi que des corrigés particulièrement détaillés, que vous pouvez consulter librement sur www.ayoub-et-les-maths.com. J'y explique à l'élève non seulement le cheminement, mais aussi et surtout pourquoi il doit avoir telle idée à tel moment, pourquoi telle autre idée n'est pas appropriée, quel piège il faut éviter ; cela, dans l'optique de l'entraîner au raisonnement réel, et non à la mimique maladroite qui est l'apanage du plus grand nombre.

Sur ma chaîne youtube « Ayoub et les maths », vous trouverez également bon nombre de vidéos d'exercices corrigés niveau Terminale et prépa (parmi ceux que j'ai donnés en colle notamment), des séries thématiques pour vous familiariser avec des notions spécifiques comme le signe somme, le calcul matriciel, la valeur absolue, ainsi que des conseils plus généraux pour vous améliorer en mathématiques. Et même, une série dénommée « Où est l'arnaque » : ce sont des exercices corrigés avec des erreurs de raisonnement volontaires - révélées en fin de vidéo, bien sûr - pour tester votre vigilance mathématique...

Table des matières

Introduction	page 1
Exercice 1 Unicité si existence <i>Fonctions, raisonnement</i>	(***) page 3 CCINP 2024 PC
Exercice 2 Dérivée seconde et inégalité logarithmique <i>Fonctions, raisonnement</i>	(**) page 9 X-ENS-ESPCI 2019 PC
Exercice 3 Suite de Héron <i>Suites, limites, raisonnement</i>	(**) page 13 Centrale 2024 PC Maths 1
Exercice 4 Intégrales, limites, factorielles <i>Fonctions, suites, limites, intégrales</i>	(****) page 19 Mines-Ponts 2017 PC Maths 2
Exercice 5 Bornes atteintes et quotient d'intégrales <i>Fonctions, intégrales, raisonnement</i>	(****) page 28 Mines-Ponts 2014 PC Maths 2
Exercice 6 Transformation d'Abel <i>Suites, limites, raisonnement</i>	(****) page 32 Centrale 2013 PC Maths 1
Exercice 7 Ordres de formules de quadrature <i>Polynômes, intégrales, raisonnement</i>	(***) page 38 Centrale 2021 PC Maths 2
Exercice 8 Fonctions à croissance lente <i>Fonctions, raisonnement, ensembles</i>	(***) page 44 Mines-Ponts 2024 PC Maths 1
Exercice 9 Continuité d'une intégrale à paramètre <i>Fonctions, trigonométrie, intégrales</i>	(****) page 49 CCINP 2020 PC
Exercice 10 Partie entière, sommes, limites <i>Suites, limites, raisonnement</i>	(***) page 55 Mines-Ponts 2016 PC Maths 1

Exercice 11	Fonction homographique et suites	(***)	page 59
	<i>Fonctions, suites, limites</i>		X-ENS-ESPCI 2022 PC
Exercice 12	Complexes et inégalités trigonométriques	(***)	page 65
	<i>Trigonométrie, complexes</i>		Mines-Ponts 2022 PC Maths 1
Exercice 13	Convexité, variables aléatoires, espérance	(****)	page 71
	<i>Fonctions, probabilités, raisonnement</i>		Centrale 2017 PC Maths 2
Exercice 14	Récurrence forte et lois de probabilités	(**)	page 78
	<i>Probabilités, dénombrement, raisonnement</i>		CCINP 2021 PC
Exercice 15	Cas d'égalité de l'inégalité triangulaire	(****)	page 84
	<i>Complexes, raisonnement</i>		Centrale 2023 PC Maths 1
Exercice 16	Somme de coefficients binomiaux	(****)	page 89
	<i>Dénombrement, raisonnement</i>		Centrale 2022 PC Maths 2
Quelques rappels de calcul matriciel			page 92
Exercice 17	Traces de puissances d'une matrice	(***)	page 98
	<i>Matrices, suites, trigonométrie</i>		X-ESPCI 2008 PC
Exercice 18	Déplacement aléatoire d'un pion	(**)	page 103
	<i>Probabilités, matrices</i>		CCINP 2019 PC
Exercice 19	Polynômes stables	(***)	page 109
	<i>Polynômes, complexes, raisonnement</i>		CCINP 2014 PC Maths 1
Exercice 20	Matrices symétriques positives	(****)	page 113
	<i>Matrices, raisonnement</i>		X-ENS-ESPCI 2011 PC

Introduction

Les problèmes présentés dans ce document sont de longueur et difficulté variables. Une évaluation de cette dernière, subjective, est précisée à titre indicatif en début de chaque problème :

(*) facile (**) moyen (***) difficile (****) particulièrement difficile

J'insiste sur le caractère subjectif de cette évaluation. Si elle s'appuie sur des paramètres tels la complexité des notions mises en oeuvre et le fait que les astuces à voir pour résoudre les questions de chaque exercice soient plus ou moins cachées, plus ou moins évidentes, l'observation empirique l'influence grandement : une sorte de moyenne vague de la difficulté ressentie par les nombreux élèves que j'ai pu côtoyer (cours particuliers, TD, colles...) face à un genre de question ou d'enchaînement de questions. Il s'agit aussi, pour moi, de comparer, en termes de difficulté, les exercices de ce document les uns aux autres.

Pas de panique, donc, si vous butez face à un exercice classé comme « moyen » ou « facile ». Faites de votre mieux dans le temps imparti, puis lisez attentivement la correction.

En parlant de temps : pour chaque problème, un temps de travail préconisé est indiqué. Il ne correspond pas forcément au temps au bout duquel l'exercice doit être résolu, mais plutôt au temps de recherche au bout duquel il devient raisonnable de commencer à regarder la correction. Il peut en effet être pertinent de vous imposer de réfléchir en temps limité, pour vous entraîner aux conditions d'examen. Mais si tel énoncé vous intrigue particulièrement, si vous vous sentez prêt de le faire tomber, si vous aimeriez en venir à bout (peut-être pour des raisons de fierté personnelle, peut-être pas), quitte à lui consacrer plus de temps que les autres, n'hésitez surtout pas. C'est en jouant de ces deux modes temporels que l'on peut s'améliorer durablement en mathématiques. Une telle idée est développée plus en détail [dans cet article](#).

Chaque énoncé est suivi de « remarques sur l'énoncé ». Il ne s'agit nécessairement d'indications sur la résolution, mais plutôt de précisions sur les notations introduites par l'énoncé, ou de rappels de concepts mathématiques utiles.

La correction à proprement parler est en caractères normaux. Les passages en italique correspondent à des commentaires sur cette correction. Principalement, le fameux « pourquoi de l'idée », cette substance fugace décrite tant bien que mal dans l'avant-propos. Mais aussi, par moments, des réflexions sur d'autres méthodes que celle choisie dans la correction, des analogies avec d'autres situations, des discussions sur telle ou telle erreur courante commise par les élèves à tel endroit. Plus rarement, deux ou trois confidences sur mes choix éditoriaux : pourquoi ai-je reformulé ainsi la question originellement présente dans le sujet ? Pourquoi ai-je ajouté telle question ? Vous mettre à la place du « concepteur » du sujet (même si le titre qui me conviendrait le mieux ici serait celui de reformulateur), tenter d'en comprendre la cohérence, vous attacher aux liens entre les questions et à leur fil conducteur peut vous aider à vous sortir du labyrinthe.

Place, maintenant, à de brèves considérations d'ordre général, qui seront complétées au besoin par les remarques sur chaque énoncé. Votre aventure mathématique dans le supérieur consistera en grande partie en la manipulation d'**assertions**. Ce sont des phrases syntaxiquement correctes (autrement dit, qui ont du sens) et qui sont soit vraies, soit fausses.

Comment ça, « soit vraies, soit fausses » ? N'est-ce pas trop général comme définition ? Toute phrase ne deviendrait-elle pas une assertion ?

Certainement pas, voyez plutôt :

- $A : « 3^2 + 12 \times 7 »$ n'est pas une assertion, car dire qu'elle serait vraie ou fausse n'aurait aucun sens.
- $B : « \forall x \in \mathbb{R}, x^2 \leq 2 »$ n'est pas une assertion, car elle n'est pas syntaxiquement correcte.
- $C : « \forall x \in \mathbb{R}^*, x^2 > 0 »$ est une assertion. C'est même une assertion vraie.
- $D : « \exists x \in \mathbb{R}, e^x \leq 0 »$ est une assertion. C'est une assertion fausse.
- $D : « x + 7 \geq 2 »$ est une assertion. C'est une assertion dont la véracité dépend du choix du paramètre x . Pour l'anecdote, une telle assertion est appelée un prédicat.

Rappelons la signification des symboles $\forall$ et $\exists$:

- $\forall$ signifie : « pour tout », ou encore « quel que soit ». L'assertion C se lit donc : « pour tout réel x non nul, $x^2 > 0$ ». Ou encore, de manière équivalente : « quel que soit le réel x non nul, $x^2 > 0$ »
- $\exists$ signifie « il existe ». L'assertion D se lit donc : « il existe un réel x tel que $e^x \leq 0$ ». Voyez comment j'ai intercalé un « tel que » à la place de la virgule, pour que la phrase ait du sens en français.

Si A et B sont deux assertions mathématiques, « $A \implies B$ » (se lit « A implique B ») veut dire que si A est vrai, alors B est vrai.

Autrement dit, A est une condition suffisante à B .

Il suffit que A soit vrai pour que B soit vrai.

Autrement dit, B est une condition nécessaire à A .

A ne peut pas être vrai sans que B ne soit vrai.

« $A \iff B$ » (se lit « A équivaut à B » ou « A est équivalent à B ») veut dire que nous avons à la fois $A \implies B$ et $B \implies A$. Autrement dit : A est vrai si et seulement si B est vrai.

La négation de « $A \implies B$ », notée $\text{non}(A \implies B)$, est « A et $\text{non}(B)$ ».

Autrement dit, A n'entraîne pas B puisque A est réalisé mais pas B

Exercice 1

*Par sadisme enseignant, l'auteur serait fort aise
de te voir bégayant entre les hypothèses.*

Énoncé : (temps conseillé : 50 min) (***) d'après CCINP 2024 PC

On admet qu'il existe une fonction $f :]0 ; +\infty[\rightarrow \mathbb{R}$ vérifiant les quatre conditions suivantes :

- (i) f est dérivable sur $]0 ; +\infty[$.
- (ii) $\forall x \in]0 ; +\infty[, f(x+1) - f(x) = \ln(x)$
- (iii) la fonction f' est croissante sur $]0 ; +\infty[$
- (iv) $f(1) = 0$

Dans la suite, on note (C) l'ensemble de ces quatre conditions.

On cherche à montrer que f est l'unique fonction vérifiant les conditions de (C).

On considère une fonction $g :]0 ; +\infty[\rightarrow \mathbb{R}$ vérifiant les conditions de (C), et on pose $h = f - g$. Les questions 1) et 2) sont indépendantes.

1) Montrer que pour tout $x > 0$, on a $h(x+1) = h(x)$ et $h'(x+1) = h'(x)$

2) Soient $x \in]0 ; 1]$ et $p \in \mathbb{N}^*$. Montrer que : $f'(p) - g'(1+p) \leq h'(x+p) \leq f'(1+p) - g'(p)$
puis que $f'(p) - g'(1+p) = h'(p) - \frac{1}{p}$. En déduire que $|h'(x+p) - h'(p)| \leq \frac{1}{p}$

3) Déduire des questions précédentes que h' est constante sur $]0 ; +\infty[$

4) Conclure.

Remarques sur l'énoncé :

Rappelons la signification de ces symboles, mis en jeu dans cet énoncé et la correction :

- $\forall$ signifie : « pour tout », ou encore « quel que soit ».
- $\exists$ signifie : « il existe »
- $\in$ signifie « appartient à ».
- $\mathbb{N}^*$ désigne l'ensemble des entiers naturels non nuls, autrement dit l'ensemble des entiers naturels supérieurs ou égaux à 1

Correction de l'exercice 1 :

1) Pour tout $x > 0$, $h(x+1) = f(x+1) - g(x+1)$. f et g vérifiant les conditions de (C), et en particulier (i), il s'ensuit : $h(x+1) = f(x) + \ln(x) - (g(x) + \ln(x)) = f(x) - g(x)$

Nous avons bien montré : $\forall x > 0, h(x+1) = h(x)$

Le résultat sur les dérivées est à portée de main. Même si nous sommes dans un cas simple, explicitons un minimum la démarche.

La fonction $x \mapsto x+1$ est bien dérivable sur $]0; +\infty[$ car polynomiale. *Et même affine.*

De plus : $\forall x \in]0; +\infty[, x+1 \in]0; +\infty[$, et h est dérivable sur $]0; +\infty[$.

Par composition, $x \mapsto h(x+1)$ est dérivable sur $]0; +\infty[$, et en dérivant membre à membre l'encadré précédent, nous obtenons bien : $\forall x > 0, h'(x+1) = h'(x)$

D'aucuns pourraient trouver les justifications précédentes un peu longues. Ce sont elles, notamment, qui permettent de se prémunir d'erreurs de dérivations qui mèneraient certains, par exemple, à se tromper en écrivant que la dérivée de la fonction $x \mapsto h(x^2)$ est la fonction $x \mapsto h'(x^2)$, alors que c'est $x \mapsto 2xh'(x^2)$.

Au-delà de prévenir d'éventuelles erreurs calculatoires, ces justifications permettent aussi de se rendre compte que « tout se passe bien au niveau des intervalles », c'est-à-dire qu'au sens de la composition, la fonction « à l'intérieur » (la première qu'on applique à x , dans notre cas la fonction $x \mapsto x+1$) est bien dérivable sur l'intervalle demandé, et qu'elle nous renvoie bien vers un intervalle sur lequel la fonction « extérieure » (h dans notre cas) est dérivable.

En parlant d'intervalles, certains élèves au sortir de la Terminale pourraient s'étonner de me voir écrire « la fonction $x \mapsto x+1$ est bien dérivable sur $]0; +\infty[$ car polynomiale ». Ils aimeraient rétorquer : pourquoi sur $]0; +\infty[$, et pas sur $\mathbb{R}$ tout entier ? Rappelons à cet effet que dire d'une fonction qu'elle est dérivable sur un intervalle I n'interdit pas qu'elle le soit sur un intervalle plus grand. Dire : « f dérivable sur $]0; +\infty[$ » n'interdit pas à f d'être aussi dérivable sur $] -\infty; 0]$. « Mais alors, vu qu'ici, $x \mapsto x+1$ est dérivable sur $\mathbb{R}$ tout entier, pourquoi ne pas avoir dit ça plutôt qu'un intervalle plus petit ? » Parce qu'ici, c'est la dérivabilité de $x \mapsto h(x+1)$ sur $]0; +\infty[$ qui m'intéresse. Et le sens de ce pavé, c'est d'appuyer l'idée qu'une part de la prise d'initiative en mathématiques, c'est d'être capable de faire le tri, dans les informations dont nous disposons, entre celles qui nous servent et celles dont nous pouvons nous délester. J'ai en tête notamment des situations - que nous croiserons peut-être plus loin dans ce document - où il est judicieux de ne travailler que sur une partie des deux inégalités d'un encadrement fourni par l'énoncé...

2) L'énoncé est assez sympathique pour nous indiquer qu'il ne faut pas utiliser les résultats de 1) ici (ce qui aurait pu nous tenter, et nous faire perdre du temps). Comment parvenir à l'encadrement demandé ? Revenons à l'expression de h en fonction de f et g ...

Pour tout $y > 0$, $h(y) = f(y) - g(y)$. Puis $h'(y) = f'(y) - g'(y)$.

En particulier : $h'(x+p) = f'(x+p) - g'(x+p)$.

Or, par hypothèse (condition (iii)), f' et g' sont croissantes sur $]0; +\infty[$. Et $p < x+p \leq 1+p$

Donc : $f'(p) \leq f'(x+p) \leq f'(1+p)$ et $g'(p) \leq g'(x+p) \leq g'(1+p)$

Nous voulons encadrer $f'(x+p) - g'(x+p)$, mais attention à ne pas faire la bêtise de soustraire membre à membre les deux encadrements précédents. Rappelons en effet que $[(a \leq b \text{ et } c \leq d)]$ entraîne $a + c \leq b + d$, mais n'entraîne pas $a - c \leq b - d$

Mais que faire dans ce cas ? Tout simplement, multiplier le second encadrement par -1 , en n'oubliant pas de changer le sens des inégalités ($-1 < 0$)

Le dernier encadrement fournit : $-g'(p) \geq -g'(x+p) \geq -g'(1+p)$,

c'est-à-dire : $-g'(1+p) \leq -g'(x+p) \leq -g'(p)$. Et rappelons : $f'(p) \leq f'(x+p) \leq f'(1+p)$

Maintenant, nous pouvons additionner membre à membre ces deux encadrements.

Puis : $f'(p) - g'(1+p) \leq f'(x+p) - g'(x+p) \leq f'(1+p) - g'(p)$

Nous avons bien montré : $f'(p) - g'(1+p) \leq h'(x+p) \leq f'(1+p) - g'(p)$

Ensuite : puisque g vérifie (ii), nous savons : $\forall x \in]0; +\infty[, g(x+1) - g(x) = \ln(x)$

g étant dérivable sur $]0; +\infty[$, nous pouvons dériver membre à membre pour obtenir :

$\forall x \in]0; +\infty[, g'(x+1) - g'(x) = \frac{1}{x}$ Une justification plus sobre qu'à la question 1)

En particulier (comme $p > 0$) : $g'(p+1) - g'(p) = \frac{1}{p}$. Autrement dit : $g'(1+p) = \frac{1}{p} + g'(p)$

Donc $f'(p) - g'(1+p) = f'(p) - \left(\frac{1}{p} + g'(p)\right) = f'(p) - g'(p) - \frac{1}{p} = h'(p) - \frac{1}{p}$.

Nous avons bien établi : $f'(p) - g'(1+p) = h'(p) - \frac{1}{p}$

Comment, de ces informations, pouvons-nous déduire que $|h'(x+p) - h'(p)| \leq \frac{1}{p}$? Nous voyons le $h'(x+p)$ au centre de l'encadrement établi précédemment. Et, forts de la dernière égalité établie, nous pouvons remplacer, dans l'encadrement, le membre de gauche par $h'(p) - \frac{1}{p}$. Mais que faire du membre de droite, $f'(1+p) - g'(p)$? Procédons donc de manière

similaire à celle que nous avons suivie pour transformer $f'(p) - g'(1+p)$

De même que pour g , nous avons : $f'(1+p) = \frac{1}{p} + f'(p)$

D'où : $f'(1+p) - g'(p) = \frac{1}{p} + f'(p) - g'(p) = \frac{1}{p} + h'(p)$

L'encadrement obtenu précédemment peut donc se réécrire :

$$h'(p) - \frac{1}{p} \leq h'(x+p) \leq h'(p) + \frac{1}{p}. \text{ Puis : } -\frac{1}{p} \leq h'(x+p) - h'(p) \leq \frac{1}{p}$$

Autrement dit : $|h'(x+p) - h'(p)| \leq \frac{1}{p}$

Rappelons en effet cette équivalence de base dans l'utilisation de la valeur absolue : pour tout $a \geq 0$ et pour tout réel t , dire « $|t| \leq a$ » revient à dire « $-a \leq t \leq a$ »

Vous êtes souvent très prompts à vous souvenir de sa définition : pour tout réel x , $|x| = x$ lorsque x est positif, et $|x| = -x$ lorsque x est négatif. Mais bien moins prompts à garder en tête les propriétés qui en découlent. Reproche un peu sévère à l'endroit d'élèves sortant de Terminale, vu qu'on n'a pas souvent l'occasion de la manipuler sérieusement pendant cette année-là. Mais en prépa, vous y aurez souvent affaire... Tenez, d'ailleurs, voici un lien vers une playlist de vidéos passant en revue l'essentiel des propriétés de la valeur absolue - que vous êtes censés avoir vues pour la plupart - et dont vous aurez besoin.

3) Nous savons que pour tout $x > 0$, $h'(x+1) = h'(x)$. Par récurrence immédiate, nous pouvons en déduire que pour tout entier naturel p : $h'(x+p) = h'(x)$ (*)

Leur balancer un « par récurrence immédiate » dès l'exercice 1... Risqué... Le but n'est pas de vous habituer à la nonchalance rédactionnelle : il ne s'agira pas d'écrire ça dès que vous aurez la flemme de rédiger une récurrence. Il s'agit plutôt de constater honnêtement que l'initialisation d'une telle récurrence est immédiate ($h'(x) = h'(x)$), et que son hérédité est honnêtement triviale : si, pour un certain entier naturel p , $h'(x+p) = h'(x)$ (hypothèse de récurrence), alors $h'(x+p+1) = h'(x+p)$ (par hypothèse sur h') et donc $h'(x+p+1) = h'(x)$ (par l'hypothèse de récurrence)

Nous en déduisons aussi : pour tout entier naturel non nul p , $h'(p) = h'(1)$

$h'(p) = h'(1+p-1) = h'(1)$ en vertu de (*), car $1 > 0$ et $p-1$ est un entier naturel

Or, nous avons établi en 2) : $\forall x \in]0; 1], \forall p \in \mathbb{N}^*, |h'(x+p) - h'(p)| \leq \frac{1}{p}$

Ce qui donne aussi, en vertu de ce qui précède : $\forall x \in]0 ; 1], \forall p \in \mathbb{N}^*, |h'(x) - h'(1)| \leq \frac{1}{p}$
C'est marrant, il n'y a pas de p dans le membre de gauche, et l'inégalité est vraie pour tout entier naturel non nul p. Aussi grand que l'on voudrait...

Soit $x \in]0 ; 1]$. Nous savons : $\forall p \in \mathbb{N}^*, 0 \leq |h'(x) - h'(1)| \leq \frac{1}{p}$. Et : $\lim_{p \rightarrow +\infty} \frac{1}{p} = 0$

Le théorème des gendarmes (ou d'encadrement, si vous préférez l'appeler comme ça, mais personnellement, je l'appellerai gendarmes dans ce document) nous permet donc d'affirmer : $\lim_{p \rightarrow +\infty} |h'(x) - h'(1)| = 0$

Mais $|h'(x) - h'(1)|$ est constant vis-à-vis de p. D'où : $\lim_{p \rightarrow +\infty} |h'(x) - h'(1)| = |h'(x) - h'(1)| \dots$

Autrement dit : $|h'(x) - h'(1)| = 0$, c'est-à-dire : $h'(x) = h'(1)$.

Nous avons donc montré : $\forall x \in]0 ; 1], h'(x) = h'(1)$ (**)

Nous y sommes presque ! Il « suffit » d'élargir l'égalité précédente - pour l'instant valable pour tout $x \in]0 ; 1]$ - à tout $x > 0$, et le tour est joué...

Soit $x > 0$. Il existe un entier naturel p tel que $p < x \leq p + 1$.

L'existence d'un tel entier naturel p est assez évidente pour qu'à ce stade, son introduction vous soit permise sans justification. Attention, toutefois, au fait que les inégalités de cet encadrement ne soient pas toutes les deux strictes (sinon, cela interdirait à x d'être entier). Pour les fêrus de partie entière, en notant $[x]$ la partie entière supérieure de x, cet entier p, en fait unique, est $p = [x] - 1$ (ce qui revient à dire que $p + 1$ est la partie entière supérieure de x).

Autrement dit, il existe un entier naturel p tel que $0 < x - p \leq 1$ Nous avons alors : $h'(x) = h'(x - p + p) = h'(x - p)$ d'après (*).

Et comme $x - p \in]0 ; 1]$, (**) nous permet d'affirmer : $h'(x - p) = h'(1)$.

Nous avons enfin établi : $\forall x > 0, h'(x) = h'(1)$

Cela nous permet de conclure que h' est constante sur $]0 ; +\infty[$

4) Comment ça, « conclure » ? Il faut revenir au but annoncé par l'énoncé : « on cherche à montrer que f est l'unique fonction vérifiant les conditions de (C) ». A cette fin, on a introduit une fonction g vérifiant ces mêmes conditions, puis $h = f - g$. Nous voulons en fait montrer que $g = f$, c'est-à-dire que h est la fonction nulle sur $]0 ; +\infty[$

Nous savons : $\exists a \in \mathbb{R}, \forall x > 0, h'(x) = a$.

h est donc une primitive sur $]0; +\infty[$ de la fonction $x \mapsto a$. Or, les primitives de $x \mapsto a$ sont les fonctions $x \mapsto ax + b$, où b est une constante. h étant une de ces primitives, il existe donc $b \in \mathbb{R}$ tel que : $\forall x > 0, h(x) = ax + b$ Reste à prouver que a et b sont nuls...

f et g vérifient la condition (iv) On l'avait oubliée, celle-là...

Donc $h(1) = f(1) - g(1) = 0 - 0 = 0$. D'autre part, $h(1) = a + b$. D'où : $a + b = 0$.

Une autre équation liant a et b , et le tour est joué.

Nous savons aussi (cf 1) : $\forall x > 0, h(x+1) = h(x)$.

Autrement dit : $\forall x > 0, a(x+1) + b = ax + b$ ou encore $ax + a + b = ax + b$. Donc : $a = 0$.

Nous pouvons en conclure : $a = b = 0$, puis : $\forall x > 0, h(x) = 0$.

h est donc la fonction nulle sur $]0; +\infty[$, ce qui revient à dire que les fonctions g et f sont égales.

f est bien l'unique fonction vérifiant les conditions de (C).

Exercice 2

*Pour voir le résultat, quelle idée bien féconde
parfois d'aller jusqu'à la dérivée seconde.*

Énoncé : (temps conseillé : 35 min) (**) *d'après X-ENS-ESPCI 2019 PC*

Soit $p \in [0 ; 1]$, et soit $g : \mathbb{R}_+ \rightarrow \mathbb{R}$ la fonction définie par $g(x) = \ln(1 - p + pe^x)$ pour tout $x \geq 0$.

1) Montrer que g est bien définie et deux fois dérivable sur $\mathbb{R}_+$. Pour $x \geq 0$, exprimer $g''(x)$ sous la forme $\frac{\alpha\beta}{(\alpha + \beta)^2}$ où α et β sont des réels positifs pouvant dépendre de x .

2) Montrer que $g''(x) \leq \frac{1}{4}$ pour tout $x \geq 0$.

3) Montrer que $\ln(1 - p + pe^x) \leq px + \frac{x^2}{8}$ pour tout $x \geq 0$.

Remarques sur l'énoncé :

Rappelons au cas où que $\mathbb{R}_+$ est l'ensemble des réels positifs. C'est donc $[0; +\infty[$

Quant à $\mathbb{R}_+^*$, c'est l'ensemble des réels positifs non nul. C'est donc l'ensemble des réels strictement positifs, c'est-à-dire $]0; +\infty[$.

Correction de l'exercice 2 :

1) Pour montrer que g est définie sur $\mathbb{R}_+$, il suffit de montrer : $\forall x \in \mathbb{R}_+, 1 - p + pe^x > 0$
Par hypothèse, $p \in [0 ; 1]$, donc $1 - p \geq 0$. De plus, pour tout $x \geq 0$, $e^x > 0$, donc $pe^x \geq 0$
Comment ça, « pour tout $x \geq 0$ » ? N'est-ce pas vrai pour tout x réel ? Si, bien sûr, mais là, j'avais juste besoin de le dire pour les réels positifs.

Nous voilà donc avec une somme de deux réels positifs, $1 - p$ et pe^x . Mais il nous faut une somme strictement positive...

En tant que somme de deux réels positifs, $1 - p + pe^x$ est nul si et seulement si $1 - p$ et pe^x sont simultanément nuls, c'est-à-dire si et seulement si, simultanément, $p = 1$ et $p = 0$, ce qui est impossible. Donc : $\forall x \in \mathbb{R}_+, 1 - p + pe^x > 0$.

Enfin, g est bien définie sur $\mathbb{R}_+$.

Remarque : on aurait aussi pu constater : $\forall x \in \mathbb{R}_+, 1 - p + pe^x = 1 + p(e^x - 1)$

Et, par croissance de la fonction exponentielle sur $\mathbb{R}_+ : \forall x \geq 0, e^x \geq e^0 = 1$, donc $e^x - 1 \geq 0$
Puis, comme $p \geq 0$, $p(e^x - 1) \geq 0$, et enfin $1 + p(e^x - 1) \geq 1 > 0$

La fonction $x \mapsto 1 - p + pe^x$ est dérivable sur $\mathbb{R}_+$ par produit et somme de fonctions dérivables. Et nous avons établi : $\forall x \in \mathbb{R}_+, 1 - p + pe^x \in \mathbb{R}_+^*$.

La fonction $\ln$ étant dérivable sur $\mathbb{R}_+^*$, nous pouvons en déduire, par composée de fonctions dérivables, que g est dérivable sur $\mathbb{R}_+$.

Nous obtenons : $\forall x \geq 0, g'(x) = \frac{pe^x}{1 - p + pe^x}$ Du $\ln(u(x))$, qui se dérive en $\frac{u'(x)}{u(x)}$

Puis : g' est dérivable sur $\mathbb{R}_+$ en tant que quotient, dont le dénominateur ne s'annule pas sur $\mathbb{R}_+$, de fonctions dérivables. Autrement dit, g est deux fois dérivable sur $\mathbb{R}_+$.

Pour tout $x \geq 0$, $g'(x) = p \times \frac{e^x}{1 - p + pe^x}$

Je mets la constante p en facteur en évidence pour ne pas me la trimballer tout le long de ma dérivation

D'où : $\forall x \geq 0, g''(x) = p \times \frac{e^x(1 - p + pe^x) - e^x \times pe^x}{(1 - p + pe^x)^2} = \frac{pe^x(1 - p)}{(1 - p + pe^x)^2}$. Enfin :

$g''(x) = \frac{\alpha\beta}{(\alpha + \beta)^2}$ avec $\alpha = pe^x \geq 0$ et $\beta = 1 - p \geq 0$

Ou inversement, si vous préférez...

2) La forme demandée précédemment nous aidera peut-être à établir cette majoration...

Montrons : $\forall \alpha, \beta \in \mathbb{R}_+, (\alpha + \beta)^2 \geq 4\alpha\beta$.

Pour tous α et β positifs : $(\alpha + \beta)^2 - 4\alpha\beta = \alpha^2 + 2\alpha\beta + \beta^2 - 4\alpha\beta = \alpha^2 - 2\alpha\beta + \beta^2 = (\alpha - \beta)^2 \geq 0$

Nous avons donc bien : $\forall \alpha, \beta \in \mathbb{R}_+, (\alpha + \beta)^2 \geq 4\alpha\beta$. Autrement dit : $\alpha\beta \leq \frac{1}{4} \times (\alpha + \beta)^2$

A deux doigts d'aboutir au résultat, mais restons rigoureux jusqu'au bout. Attention à la division par zéro...

α et β sont deux réels positifs. Si $\alpha + \beta > 0$, et donc $(\alpha + \beta)^2 > 0$. Dans ce cas, nous obtenons donc : $\frac{\alpha\beta}{(\alpha + \beta)^2} \leq \frac{1}{4}$

En 1), nous avons établi : $\forall x \geq 0, g''(x) = \frac{\alpha\beta}{(\alpha + \beta)^2}$ avec $\alpha = pe^x \geq 0$ et $\beta = 1 - p \geq 0$

$\alpha + \beta = 1 - p + pe^x > 0$ (établi en 1 pour la définition de g). Ouf, condition respectée !

Finalement : $\text{pour tout } x \geq 0, g''(x) \leq \frac{1}{4}$

L'inégalité $(\alpha + \beta)^2 \geq 4\alpha\beta$, valable pour tous réels α et β , est équivalente à : $\alpha\beta \leq \frac{1}{4}(\alpha^2 + \beta^2)$
Comme elle est vraie quels que soient les signes de α et β , elle fournit plus généralement : $|\alpha\beta| \leq \frac{1}{4}(\alpha^2 + \beta^2)$. Cette inégalité, relativement simple à obtenir, s'avérera très pratique (tellement pratique que mon prof de maths en MP l'appelait « la clé ») pour pas mal de majorations dont vous aurez besoin au cours de votre cursus...

Ici, j'ai un peu forcé le lien avec les inégalités de 2), mais cette clé s'obtient tout simplement à partir de $(|\alpha| - |\beta|)^2 \geq 0$, équivalente à $|\alpha|^2 - 2|\alpha||\beta| + |\beta|^2 \geq 0$, c'est-à-dire $\alpha^2 + \beta^2 \geq 2|\alpha\beta|$

3) Montrons que pour tout $x \geq 0, \ln(1 - p + pe^x) \leq px + \frac{x^2}{8}$, c'est-à-dire : $px + \frac{x^2}{8} - g(x) \geq 0$

Reconnaître g , c'est bien la moindre des choses. Et maintenant... Etude de fonction, avec dérivations successives, pour prouver que cette quantité est positive (le fait de nous avoir fait pencher longtemps sur $g''(x)$ nous y fait penser).

Soit h la fonction définie sur $\mathbb{R}_+$ par $h(x) = px + \frac{x^2}{8} - g(x)$.

h est deux fois dérivable sur $\mathbb{R}_+$, et : $\forall x \in \mathbb{R}_+, h'(x) = p + \frac{2x}{8} - g'(x) = p + \frac{x}{4} - g'(x)$

Puis : $\forall x \in \mathbb{R}_+, h''(x) = \frac{1}{4} - g''(x)$. Oh toi, je connais ton signe...

D'après 2), pour tout réel positif x , $h''(x) \geq 0$. Donc h' est croissante sur $\mathbb{R}_+$.

D'accord, mais je préférerais le signe de h' sur $\mathbb{R}_+$ plutôt que ses variations.

$$\text{De plus, } h'(0) = p - g'(0) = p - \frac{pe^0}{1-p+pe^0} = p - \frac{p}{1} = 0.$$

La croissance de h' sur $\mathbb{R}_+$ entraîne donc : $\forall x \in \mathbb{R}_+, h'(x) \geq h'(0) = 0$.

h est donc croissante sur $\mathbb{R}_+$.

D'accord, mais je préférerais le signe de h sur $\mathbb{R}_+$ plutôt que ses variations...

$$h(0) = p \times 0 + 0 - g(0) = -\ln(1 - p + pe^0) = -\ln(1) = 0$$

La croissance de h sur $\mathbb{R}_+$ entraîne donc : $\forall x \in \mathbb{R}_+, h(x) \geq h(0) = 0$.

Nous avons donc montré : $\forall x \geq 0, px + \frac{x^2}{8} - g(x) \geq 0$

Autrement dit : $\text{pour tout } x \geq 0, \ln(1 - p + pe^x) \leq px + \frac{x^2}{8}$

Ces dérivations successives d'une fonction différence bien choisie - ici h - pourront vous tirer d'affaire dans pas mal de situations où il faut établir une inégalité difficile à obtenir par opérations « simples », et où le fait de dériver arrange l'expression (dans notre cas, le fait de dériver deux fois nous a permis de tomber sur une fonction dont nous connaissons le signe grâce à une information préalable sur g).

Bien sûr, pour que les choses se passent bien, il faut un minimum de « chance », ou du moins un pari sur le fait qu'en rebroussant chemin de h'' à h , les signes des dérivées s'obtiendront relativement simplement.

Vous avez peut-être (ou pas) déjà croisé ce procédé en Terminale, notamment dans le cadre d'inégalités avec des fonctions trigonométriques.

Voici un lien vers un cas classique de son utilisation, qui permet de démontrer une croissance comparée mettant en jeu la fonction exponentielle.

Exercice 3

*Un point dans ce ciel gris. Concentrez-vous pour voir
Héron d'Alexandrie perché sur son grand phare.*

Énoncé : (temps conseillé : 1 h 10 min) (**) d'après Centrale 2024 PC Maths 1

Soit $a \in \mathbb{R}_+$. On définit la suite $(c_n)_{n \in \mathbb{N}}$ par :
$$\begin{cases} c_0 &= 1 \\ \forall n \in \mathbb{N}, c_{n+1} &= \frac{1}{2} \left(c_n + \frac{a}{c_n} \right) \end{cases}$$

- 1) Montrer par récurrence que pour tout $n \in \mathbb{N}$, c_n est bien défini et que $c_n > 0$
- 2) Pour tout $n \in \mathbb{N}$, donner une expression de $c_{n+1}^2 - a$ faisant intervenir $(c_n^2 - a)^2$
- 3) En déduire que pour tout $n \geq 1$, $c_n \geq \sqrt{a}$
- 4) Montrer que (c_n) converge vers $\sqrt{a}$
- 5) Dans cette question, on suppose que $a = 2$.
 - a) Calculer c_1 , puis montrer que pour tout $n \in \mathbb{N}^*$, $c_n^2 - 2 \leq 8 \left(\frac{1}{32} \right)^{2^{n-1}}$
 - b) En déduire qu'il existe un réel positif M tel que : $\forall n \in \mathbb{N}^*, |c_n - \sqrt{2}| \leq M \left(\frac{1}{32} \right)^{2^{n-1}}$

Remarques sur l'énoncé :

L'énoncé originel appelait $(c_n(a))$ la suite que j'appelle ici (c_n) pour alléger vos calculs (et les miens). De rien.

Dans la question 1, il faut montrer que pour tout $n \in \mathbb{N}$, c_n est bien défini. Comment ça ? Posez-vous la question de ce qui pourrait mal se passer avec la définition de la suite (c_n) .

Correction de l'exercice 3 :

1) Bon, ben, « par récurrence », les choses sont dites...

Soit, pour tout $n \in \mathbb{N}$, la propriété P_n : « c_n est bien défini et $c_n > 0$ »

Montrons par récurrence que pour tout entier naturel n , P_n est vraie.

Initialisation : c_0 est bien défini, $c_0 = 1 > 0$, donc P_0 est vraie.

Hérédité : Supposons que pour un certain $n \in \mathbb{N}$, P_n soit vraie, et montrons que P_{n+1} aussi est vraie.

Supposons donc que c_n soit bien défini, et que $c_n > 0$.

c_n est en particulier non nul, donc $\frac{1}{2}(c_n + \frac{a}{c_n})$ est bien défini. Autrement dit, c_{n+1} est bien défini. Et comme $c_n > 0$ et $a \geq 0$: $\frac{a}{c_n} \geq 0$, puis $c_n + \frac{a}{c_n} \geq c_n > 0$. D'où : $c_{n+1} > 0$.

P_{n+1} est donc vraie.

Pas l'hérédité la plus dure...

Conclusion : Le principe de raisonnement par récurrence nous permet de conclure que pour tout entier naturel n , P_n est vraie.

Autrement dit, pour tout entier naturel n , c_n est bien défini et $c_n > 0$.

$$2) \text{ Pour tout } n \in \mathbb{N}, c_{n+1}^2 - a = \frac{1}{4}\left(c_n + \frac{a}{c_n}\right)^2 - a = \frac{1}{4}\left(c_n^2 + 2c_n \times \frac{a}{c_n} + \frac{a^2}{c_n^2}\right) - a$$

$$\text{Donc : } c_{n+1}^2 - a = \frac{1}{4} \times \left(c_n^2 + 2a + \frac{a^2}{c_n^2}\right) - a = \frac{c_n^2}{4} + \frac{a}{2} + \frac{a^2}{4c_n^2} - a = \frac{c_n^2}{4} - \frac{a}{2} + \frac{a^2}{4c_n^2}$$

Que faire de tout ce bazar ? Une mise sous le même dénominateur, pour y voir clair...

$$\text{Puis : } c_{n+1}^2 - a = \frac{c_n^4 - 2ac_n^2 + a^2}{4c_n^2} \quad \text{Ah. Et là, que remarque-t-on ?}$$

$$\text{Enfin : } \forall n \in \mathbb{N}, c_{n+1}^2 - a = \frac{(c_n^2 - a)^2}{4c_n^2}$$

On a vraiment fini, là ? L'énoncé - qui demande une expression « faisant intervenir $(c_n^2 - a)^2$ » - est certes un peu vague mais, d'une part, notre expression fait bien intervenir $(c_n^2 - a)^2$, et d'autre part, je ne vois pas comment on pourrait faire plus simple en respectant cette contrainte. Un ou deux petits malins auraient pu se dire : « ben : $c_{n+1}^2 - a = (c_n^2 - a)^2 - (c_n^2 - a)^2 + c_{n+1}^2 - a$ ». Qu'ils essayent aux concours, et qu'ils nous donnent de leurs nouvelles...

3) Pour tout $n \in \mathbb{N}$, $c_{n+1}^2 - a = \frac{(c_n^2 - a)^2}{4c_n^2} \geq 0$ Un carré de réel est toujours positif

Donc : $c_{n+1}^2 \geq a$. Puis, par croissance de la fonction carré sur $\mathbb{R}_+$, $\sqrt{c_{n+1}^2} \geq \sqrt{a}$.

De plus, $c_{n+1} > 0$, donc $\sqrt{c_{n+1}^2} = c_{n+1}$

Rappelons que pour x réel, $\sqrt{x^2} = |x| = \begin{cases} x & \text{si } x \geq 0 \\ 0 & \text{si } x < 0 \end{cases}$. Donc on n'a pas toujours $\sqrt{x^2} = x$

Nous avons donc montré : $\forall n \in \mathbb{N}$, $c_{n+1} \geq \sqrt{a}$. Autrement dit : $\forall n \geq 1, c_n \geq \sqrt{a}$

Dire : « pour tout $n \geq 0$, une propriété P_{n+1} est vraie » revient effectivement à dire, en décalant : « pour tout $n \geq 1$, P_n est vraie » De manière explicite (mais vous n'aurez pas à le redémontrer à chaque fois que vous aurez recours à cette reformulation) :

- si pour tout $n \geq 0$, P_{n+1} est vraie : comme, pour tout $n \geq 1$, $n-1 \geq 0$, $P_{(n-1)+1}$ est vraie. D'où : pour tout $n \geq 1$, P_n est vraie.

- réciproquement : si, pour tout $n \geq 1$, P_n est vraie : comme, pour tout $n \geq 0$, $n+1 \geq 1$, on a bien : $\forall n \geq 0$, P_{n+1} est vraie

4) On vient de nous faire prouver qu'à partir du rang 1, (u_n) est minorée par $\sqrt{a}$... Etudions donc les variations de cette suite, en espérant très fort qu'elle soit décroissante (au moins à partir d'un certain rang).

Pour tout $n \in \mathbb{N}$, $c_{n+1} - c_n = \frac{1}{2} \left(c_n + \frac{a}{c_n} \right) - c_n = -\frac{c_n}{2} + \frac{a}{2c_n} = \frac{-c_n^2 + a}{2c_n}$, avec $2c_n > 0$

$c_{n+1} - c_n$ est donc du signe de $a - c_n^2$. Or : $\forall n \geq 1$, $c_n \geq \sqrt{a}$. D'où, par croissance de la fonction carré sur $\mathbb{R}_+$: $\forall n \geq 1$, $c_n^2 \geq a$

Oui, c'est vrai, on détricote un peu ce qu'on a tricoté en début de 3)...

Donc : $\forall n \geq 1$, $a - c_n^2 \leq 0$, et donc $c_{n+1} - c_n \leq 0$

A partir du rang 1, la suite (u_n) est décroissante et minorée (par $\sqrt{a}$), donc d'après le théorème de convergence monotone, elle converge vers un réel $l \geq \sqrt{a}$.

Aucune raison, à ce stade, de conclure « vers $\sqrt{a}$ ».

De plus, en considérant la fonction $f : x \mapsto \frac{1}{2} \left(x + \frac{a}{x} \right)$, nous avons : $\forall n \in \mathbb{N}$, $c_{n+1} = f(c_n)$

Et là, je veux me servir de ce sympathique théorème que vous voyez en Terminale, parfois sous le nom de « théorème du point fixe », qui vous assure que si (u_n) converge vers l , et si l appartient à un intervalle I sur lequel f est continue, alors $f(l) = l$. Mais j'ai un petit souci d'appartenance... Je sais que l appartient à $[\sqrt{a}; +\infty[$, mais si $a = 0$, cet intervalle contient 0, et f n'est même pas définie en 0... Si $a = 0$... Traitons ce cas à part.

Nous savons que $a \in \mathbb{R}_+$.

Si $a > 0$: (c_n) converge vers un réel $l \in [\sqrt{a}; +\infty[$. Or : $[\sqrt{a}; +\infty[\subset]0; +\infty[$.

Pour A et B deux ensembles, « $A \subset B$ » se lit : « A est inclus dans B »

Donc : $l \in]0; +\infty[$. Et la fonction f est bien continue sur $]0; +\infty[$. Le théorème du point fixe nous permet de conclure que $f(l) = l$. Autrement dit : $\frac{1}{2}(l + \frac{a}{l}) = l$

Ou encore : $2l = l + \frac{a}{l}$, c'est-à-dire : $l = \frac{a}{l}$

Nous obtenons donc : $l^2 = a$, et enfin (comme $l \geq 0$) : $l = \sqrt{a}$

Le théorème cité permet d'aller vite, mais dans ce cas où f est relativement simple (ne fait par exemple pas intervenir $\exp, \ln...$), on peut aussi en faire l'économie, en reprenant : $\forall n \in \mathbb{N}, c_{n+1} = \frac{1}{2}(c_n + \frac{a}{c_n})$. Et comme nous savons que (c_n) converge vers $l \neq 0$, un passage à la limite (faire tendre n vers $+\infty$) de chaque côté - avec quotient, somme et produit de limites du côté droit - donne, par unicité de la limite, cette même égalité : $l = \frac{1}{2}(l + \frac{a}{l})$

Si $a = 0$: la relation de récurrence vérifiée par (c_n) devient : $\forall n \in \mathbb{N}, c_{n+1} = \frac{1}{2} \times c_n$

La suite (c_n) est donc géométrique de raison $\frac{1}{2} \in]-1; 1[$. D'où : (c_n) converge vers 0.

Dans ce cas aussi, (c_n) converge bien vers $\sqrt{a}$, c'est-à-dire $\sqrt{0}$

Nous avons bien établi, dans tous les cas, que (c_n) converge vers $\sqrt{a}$

$$5)a) c_1 = \frac{1}{2}(c_0 + \frac{2}{c_0}) = \frac{1}{2}(c_0 + \frac{2}{c_0}) = \frac{1}{2}(1 + \frac{2}{1}) \text{ donc } c_1 = \frac{3}{2}$$

D'après 2) : pour tout $n \in \mathbb{N}$, $c_{n+1}^2 - 2 = \frac{(c_n^2 - 2)^2}{4c_n^2}$

Ce serait sympa d'avoir une inégalité faisant intervenir $c_n^2 - 2$ et $c_{n+1}^2 - 2$...

Or, d'après 3) : pour tout $n \in \mathbb{N}^*$, $c_n^2 \geq 2$, puis $4c_n^2 \geq 8$ et, par décroissance de la fonction inverse sur $\mathbb{R}_+$, $\frac{1}{4c_n^2} \leq \frac{1}{8}$. D'où (comme $(c_n^2 - 2)^2 \geq 0$) : $\frac{(c_n^2 - 2)^2}{4c_n^2} \leq \frac{(c_n^2 - 2)^2}{8}$

Autrement dit : $\forall n \in \mathbb{N}^*, c_{n+1}^2 - 2 \leq \frac{1}{8} \times (c_n^2 - 2)^2$ (*)

Bon, pas mal ce petit 8... Mais quid du $(\frac{1}{32})^{2^{n-1}}$ à faire apparaître ? Cette puissance de 2 en exposant... Peut-être qu'une récurrence... Avec l'inégalité obtenue précédemment qui nous servirait dans l'hérédité...

Soit, pour tout $n \in \mathbb{N}^*$, la propriété Q_n : « $c_n^2 - 2 \leq 8(\frac{1}{32})^{2^{n-1}}$ »

Montrons par récurrence que pour tout $n \in \mathbb{N}^*$, Q_n est vraie.

Initialisation : $c_1^2 - 2 = \frac{9}{4} - \frac{8}{4} = \frac{1}{4}$ et $8(\frac{1}{32})^{2^{1-1}} = 8(\frac{1}{32})^1 = \frac{1}{4}$ donc Q_1 est vraie.

L'inégalité au sens large est bien vraie.

Hérédité : Supposons que pour un certain $n \in \mathbb{N}^*$, Q_n soit vraie, et montrons que Q_{n+1} aussi est vraie.

Supposons donc que $c_n^2 - 2 \leq 8(\frac{1}{32})^{2^{n-1}}$. Et faisons intervenir $\frac{1}{8} \times (c_n^2 - 2)^2$...

Nous savons en fait (d'après 3) : $0 \leq c_n^2 - 2 \leq 8(\frac{1}{32})^{2^{n-1}}$

Par croissance de la fonction carré sur $\mathbb{R}_+$: (c'est pour utiliser ça que j'ai rappelé $0 \leq \dots$)

$(c_n^2 - 2)^2 \leq 64 \times \left[(\frac{1}{32})^{2^{n-1}} \right]^2 = 64 \times (\frac{1}{32})^{2^{n-1} \times 2}$ Donc : $(c_n^2 - 2)^2 \leq 64 \times (\frac{1}{32})^{2^n}$

L'occasion de vérifier que vos règles de calcul sur les puissances sont bien maîtrisées...

Puis : $\frac{1}{8} \times (c_n^2 - 2)^2 \leq \frac{64}{8} \times (\frac{1}{32})^{2^n} = 8 \times (\frac{1}{32})^{2^n}$. Nous avons établi par ailleurs (cf (*)) :

$c_{n+1}^2 - 2 \leq \frac{1}{8} \times (c_n^2 - 2)^2$. Nous en déduisons : $c_{n+1}^2 - 2 \leq 8(\frac{1}{32})^{2^n}$.

« Par transitivité de la relation d'ordre $\leq$, si on voulait frimer...

Q_{n+1} est donc vraie.

Conclusion : Le principe de raisonnement par récurrence nous permet de conclure que pour tout entier naturel non nul n , Q_n est vraie.

Autrement dit, pour tout entier naturel non nul n , $c_n^2 - 2 \leq 8(\frac{1}{32})^{2^{n-1}}$.

5)b) Vu ce que nous avons établi en 5)a), ce qu'ils nous demandent maintenant n'a pas

l'air si difficile...

D'après 5)a), nous savons : $\forall n \in \mathbb{N}^*, (c_n - \sqrt{2})(c_n + \sqrt{2}) \leq 8\left(\frac{1}{32}\right)^{2^{n-1}}$, avec $c_n + \sqrt{2} > 0$

Donc : $c_n - \sqrt{2} \leq \frac{8}{c_n + \sqrt{2}}\left(\frac{1}{32}\right)^{2^{n-1}} \leq \frac{8}{\sqrt{2}} \times \left(\frac{1}{32}\right)^{2^{n-1}}$ (car $c_n + \sqrt{2} > \sqrt{2}$ et $8 \times \left(\frac{1}{32}\right)^{2^{n-1}} \geq 0$)

Enfin, comme $c_n - \sqrt{2} \geq 0$, $|c_n - \sqrt{2}| = c_n - \sqrt{2} \leq \frac{8}{\sqrt{2}} \times \left(\frac{1}{32}\right)^{2^{n-1}}$

Avec le réel positif $M = \frac{8}{\sqrt{2}} = 4\sqrt{2}$, nous avons bien établi :

$$\forall n \in \mathbb{N}^*, |c_n - \sqrt{2}| \leq M\left(\frac{1}{32}\right)^{2^{n-1}}$$

Remarque : j'aurais pu obtenir une majoration plus fine avec un M plus petit : à l'étape, $\frac{8}{c_n + \sqrt{2}}\left(\frac{1}{32}\right)^{2^{n-1}} \leq \frac{8}{\sqrt{2}} \times \left(\frac{1}{32}\right)^{2^{n-1}}$, j'aurais pu écrire : $\frac{8}{c_n + \sqrt{2}}\left(\frac{1}{32}\right)^{2^{n-1}} \leq \frac{8}{\sqrt{2} + \sqrt{2}} \times \left(\frac{1}{32}\right)^{2^{n-1}}$ puisque $c_n \geq \sqrt{2}$

Et j'aurais donc établi l'inégalité encadrée, cette fois-ci avec $M = 2\sqrt{2}$. Mais bon, puisqu'on ne m'a rien imposé sur M mis à part qu'il soit un réel positif..

Voici un lien vers un exercice sur les suites auquel, par certains aspects, celui que vous venez de subir me fait penser.

Exercice 4

*Factorielle puissance intégrale à la suite :
je vois dans tous les sens des gens prendre la fuite...*

Énoncé : (temps conseillé : 1 h 15 min) (****) d'après Mines-Ponts 2017 PC Maths 2

On rappelle qu'une suite $(u_n)_{n \in \mathbb{N}}$ converge vers un réel l si et seulement si :
 $\forall \epsilon > 0, \exists N_\epsilon \in \mathbb{N}, \forall n \geq N_\epsilon, |u_n - l| < \epsilon$

Soit $b \in]0 ; 1[$. Pour tout entier naturel non nul n , on pose $I_n = \int_0^b (te^{-t})^n dt$

Soit y un réel strictement positif. On pose, pour tout $n \in \mathbb{N}^*$, $a_n = \frac{n^{n+1}}{n!} y^n$.

1) Montrer que : $\lim_{x \rightarrow 0} \frac{\ln(1+x)}{x} = 1$

2) Calculer $\lim_{n \rightarrow +\infty} \frac{a_{n+1}}{a_n}$. On pourra s'intéresser à : $\lim_{n \rightarrow +\infty} \left(1 + \frac{1}{n}\right)^n$

3) En déduire que, si $y < e^{-1}$, alors $\lim_{n \rightarrow +\infty} a_n = 0$. On pourra, après l'avoir justifié, utiliser le fait qu'il existe un entier naturel non nul N tel que : $\forall n \geq N, a_n \leq ye + \frac{1-ye}{2}$

4) Etudier les variations de la fonction $f : u \mapsto ue^{-u}$ sur $\mathbb{R}_+$

5) Montrer enfin que : $\lim_{n \rightarrow +\infty} \frac{n^{n+1}}{n!} I_n = 0$

Remarques sur l'énoncé :

Rappelons au cas où que $\mathbb{R}_+$ est l'ensemble des réels positifs. C'est donc $[0; +\infty[$
Quant à $\mathbb{R}_+^*$, c'est l'ensemble des réels positifs non nul. C'est donc l'ensemble des réels strictement positifs, c'est-à-dire $]0; +\infty[$.

Rappelons aussi la signification de ces symboles mis en jeu :

- $\forall$ signifie : « pour tout », ou encore « quel que soit ».
- $\exists$ signifie : « il existe »
- $\in$ signifie « appartient ».

La définition de convergence avec les quantificateurs rappelée par l'énoncé est rarement présentée en Terminale (ou alors, de manière succincte, sans insister dessus). A moins que vous ne soyez particulièrement familier avec cette définition, et que vous n'ayez eu l'occasion de la manipuler assez, cet énoncé, bien que court, risque de vous donner du fil à retordre...

« $\forall \epsilon > 0, \exists N_\epsilon \in \mathbb{N}, \forall n \geq N_\epsilon, |u_n - l| < \epsilon$ » donne, en français : « pour tout réel epsilon strictement positif, il existe un entier naturel N_ϵ tel que pour tout (entier naturel) n supérieur ou égal à ce N_ϵ , $|u_n - l| < \epsilon$ (autrement dit la distance entre u_n et l est inférieure à ϵ) ».

Dire $|u_n - l| < \epsilon$ revient à dire $l - \epsilon < u_n < l + \epsilon$

L'idée est donc la suivante : les termes de la suite peuvent être rendus aussi proches que l'on veut de l à condition de prendre des rangs n assez grands.

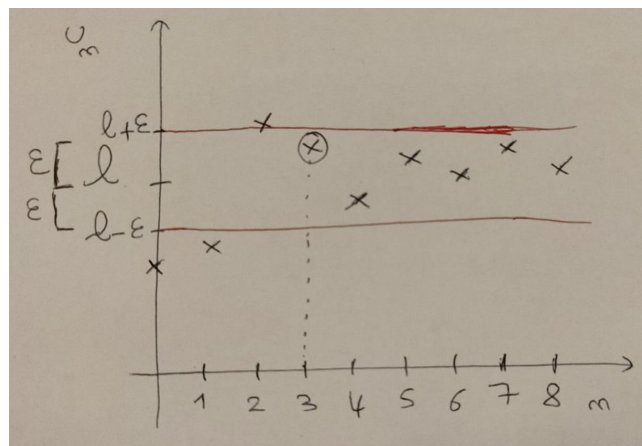


Schéma à main levée (je pense que ça se voit...) représentant les premiers termes d'une suite (u_n) convergeant vers l

Le schéma ci-dessus est censé représenter les premiers termes d'une suite convergeant vers l . Quel que soit le « faisceau » de rayon non nul ($\epsilon > 0$) centré autour de l (correspondant à l'intervalle $]l - \epsilon; l + \epsilon[$ en ordonnée), il existe un rang N_ϵ à partir duquel tous les termes de la suite sont emprisonnés dans le faisceau.

Et ce, quel que soit le rayon ϵ du faisceau, du moment qu'il est non nul...

Dans cet exemple, pour le ϵ de la figure, $N_\epsilon = 3$ semble convenir (tous les entiers supé-

rieurs ou égaux à 3 aussi, d'ailleurs...). Pour tous les n supérieurs ou égaux à ce N_ϵ , on voit effectivement : $l - \epsilon < u_n < l + \epsilon$ (autrement dit, $-\epsilon < u_n - l < \epsilon$, ce qui équivaut à dire : $|u_n - l| < \epsilon$)

Une variante légèrement plus rigoureuse de l'écriture avec les quantificateurs (bien que la précédente soit couramment utilisée et tolérée) est :

$$\forall \epsilon > 0, \exists N_\epsilon \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq N_\epsilon \implies |u_n - l| < \epsilon$$

Pourquoi plus rigoureuse ? Parce que dans la première, lorsqu'on écrit $\forall n \geq N_\epsilon$, on n'indique pas la nature du n en question. Un correcteur de mauvaise foi pourrait faire semblant de comprendre que les n dont on parle seraient des réels en général (alors que le contexte indique clairement que ce sont plus précisément des entiers naturels mais bon, que pouvons-nous contre la mauvaise foi...)

Vous pourrez rencontrer d'autres variantes de cette définition parlant plutôt de N au lieu de N_ϵ (ce qui est tout à fait permis, mais j'aime bien N_ϵ , ça nous rappelle que ce rang dépend de ϵ), ou encore mettant un signe $\leq$ à la place de $<$ (ce qui ne change rien à la définition, grâce au $\forall \epsilon$ du début)

Finissons ces remarques avec un rappel sur la notion de divergence vers $+\infty$.

Dire : $\lim_{n \rightarrow +\infty} \lambda_n = +\infty$ revient à dire : $\forall M > 0, \exists n_0 \in \mathbb{N}, \forall n \geq n_0, \lambda_n \geq M$

En français : « pour tout réel M strictement positif, il existe un entier naturel n_0 entier naturel tel que pour tout (entier naturel) n supérieur ou égal à ce n_0 , $\lambda_n \geq M$ »

Autrement dit, λ_n peut être rendu aussi grand que l'on veut (plus grand que n'importe quel M) pourvu que n soit assez grand.

Correction de l'exercice 4 :

1) Une limite classique que pas mal d'entre vous ont dû voir en Terminale. Mais ici, il faut la redémontrer. Comment faire ? Un calcul direct donne une forme indéterminée : le numérateur et le dénominateur tendant tous deux vers 0 lorsque x tend vers 0. Confrontés à ce genre de situation, beaucoup d'élèves de Terminale lancent « factorisation ! » par réflexe. Ici, pas de factorisation particulière à se mettre sous la dent...

Une autre considération - boudée par les élèves en général, parce qu'elle fait appel à une notion plus souvent manipulée en Première qu'en Terminale - permet parfois de lever ce genre d'indétermination.

Rappelons que si f est une fonction définie sur un intervalle I et si a est un réel de I : f est dérivable en a si et seulement si le taux d'accroissement $\frac{f(x)-f(a)}{x-a}$ admet une limite finie en zéro. Dans ce cas : $\lim_{x \rightarrow a} \frac{f(x)-f(a)}{x-a} = f'(a)$. En Première, vous vous en êtes souvent servis pour prouver l'existence du nombre dérivé $f'(a)$ et le déterminer. Mais dans l'autre sens, si vous savez que $f'(a)$ existe et que vous avez sa valeur, vous obtenez l'existence et la valeur de $\lim_{x \rightarrow a} \frac{f(x)-f(a)}{x-a}$...

Soit f la fonction définie sur $] -1 ; +\infty[$ par $f(x) = \ln(1+x)$. f est dérivable sur $] -1 ; +\infty[$ (et en particulier, en 0) en tant que composée de fonctions dérivables : en effet, $x \mapsto 1+x$ est bien dérivable sur $] -1 ; +\infty[$ (fonction affine) et, pour tout x appartenant à $] -1 ; +\infty[$, $1+x$ appartient à $] 0 ; +\infty[$, intervalle sur lequel $\ln$ est bien dérivable.

Pour tout $x > -1$, $f'(x) = \frac{1}{1+x}$, et en particulier : $f'(0) = \frac{1}{1+0} = 1$

D'où : $\lim_{x \rightarrow 0} \frac{f(x)-f(0)}{x-0} = f'(0) = 1$. Or, $f(0) = \ln(1+0) = \ln(1) = 0$.

Nous avons donc bien établi : $\lim_{x \rightarrow 0} \frac{\ln(1+x)}{x} = 1$

Cette démo classique en vidéo

2) Voyons voir ce que donne l'expression de $\frac{a_{n+1}}{a_n}$, avant de nous intéresser à la limite suggérée par l'énoncé.

Pour tout entier naturel non nul n , $a_n \neq 0$ (en particulier car $y \neq 0$), donc $\frac{a_{n+1}}{a_n}$ est bien défini. Et, pour tout $n \in \mathbb{N}^*$, $\frac{a_{n+1}}{a_n} = \frac{(n+1)^{n+2}}{(n+1)!} y^{n+1} \times \frac{n!}{n^{n+1}} \times \frac{1}{y^n}$

Oui pardon, plutôt que d'écrire un gros quotient immonde de deux fractions, j'ai préféré multiplier la première par l'inverse de la seconde... Place aux simplifications, maintenant.

Rappelons que : $(n+1)! = n! \times (n+1)$. Nous obtenons donc : $\frac{a_{n+1}}{a_n} = \frac{(n+1)^{n+2}}{n+1} \times \frac{1}{n^{n+1}} \times y$

$$\text{Puis : } \frac{a_{n+1}}{a_n} = \frac{(n+1)^{n+1}}{n^{n+1}} \times y = \left(\frac{n+1}{n}\right)^{n+1} \times y = \left(1 + \frac{1}{n}\right)^{n+1} \times y$$

Et nous voulons la limite de cette expression lorsque n tend vers $+\infty$.

Premier jet : $\lim_{n \rightarrow +\infty} \frac{1}{n} = 0$ donc $\lim_{n \rightarrow +\infty} 1 + \frac{1}{n} = 1$. Or, pour tout entier naturel n , $1^{n+1} = 1$.

Donc : $\lim_{n \rightarrow +\infty} \left(1 + \frac{1}{n}\right)^{n+1} = 1$. Vous êtes d'accord avec ça ? Vous ne devriez pas.

Pourquoi donc ? Au-delà du fait qu'il puisse paraître trop simple - et nous dispense même d'utiliser l'indication de l'énoncé - quel reproche objectif formuler à un tel raisonnement ?

Il est tout à fait exact d'affirmer : $\lim_{n \rightarrow +\infty} \frac{1}{n} = 0$. Il est tout aussi exact de dire que pour

tout entier naturel n , $1^{n+1} = 1$. Il est parfaitement faux d'en déduire $\lim_{n \rightarrow +\infty} \left(1 + \frac{1}{n}\right)^{n+1} = 1$.

Pourquoi ? Parce que la variable n , que nous faisons tendre vers $+\infty$, est présente à la fois dans $1 + \frac{1}{n}$ et dans l'exposant $n+1$.

Dire : $\lim_{n \rightarrow +\infty} \left(1 + \frac{1}{n}\right)^{1000} = 1$ eût été exact (remplacez 1000 par l'entier fixé de votre choix, ça

reste vrai). D'ailleurs, qu'est-ce qui vous permet de l'affirmer ? Le fait que : $\lim_{n \rightarrow +\infty} 1 + \frac{1}{n} = 1$,

et le fait que la fonction $g : x \mapsto x^{1000}$ est continue (car polynomiale) sur $\mathbb{R}$ (donc en particulier en 1), nous permettent de conclure, en composant : $\lim_{n \rightarrow +\infty} g\left(1 + \frac{1}{n}\right) = g(1) = 1^{1000}$.

Autrement dit : $\lim_{n \rightarrow +\infty} \left(1 + \frac{1}{n}\right)^{1000} = 1$

Mais avec n en exposant, qui serait votre fonction g ? Ce serait plutôt la fonction $g_n : x \mapsto x^n$, qui change à chaque rang n ... Vous ne pouvez plus faire semblant d'effectuer une composée de limites licite. Bon, reprenons notre calcul, et tentons autre chose...

Pour tout $n \in \mathbb{N}^*$: $\frac{a_{n+1}}{a_n} = \left(1 + \frac{1}{n}\right)^n \times \left(1 + \frac{1}{n}\right) \times y$ avec $\lim_{n \rightarrow +\infty} \left(1 + \frac{1}{n}\right) \times y = 1 \times y = y$

Quid de $\left(1 + \frac{1}{n}\right)^n$? Nous allons lui faire subir une transformation à laquelle vous aurez souvent recours face à une quantité dépendant de votre variable et élevée à une puissance dépendant aussi de votre variable...

Pour tout $n \in \mathbb{N}^*$, $\left(1 + \frac{1}{n}\right)^n = \exp \left[\ln \left(\left(1 + \frac{1}{n}\right)^n \right) \right] = \exp \left[n \ln \left(1 + \frac{1}{n} \right) \right]$

Euh, d'accord... Et en quoi ça nous avance pour déterminer sa limite ? Voyez plutôt :

Pour tout $n \in \mathbb{N}^*$, $n \ln \left(1 + \frac{1}{n} \right) = \frac{\ln \left(1 + \frac{1}{n} \right)}{\frac{1}{n}}$ Oh, cette forme me rappelle quelque chose...

Or, : $\lim_{n \rightarrow +\infty} \frac{1}{n} = 0$. Et nous avons établi en 1) : $\lim_{x \rightarrow 0} \frac{\ln(1+x)}{x} = 1$

Donc, par composée de limites : $\lim_{n \rightarrow +\infty} \frac{\ln \left(1 + \frac{1}{n} \right)}{\frac{1}{n}} = 1$. Autrement dit : $\lim_{n \rightarrow +\infty} n \ln \left(1 + \frac{1}{n} \right) = 1$

Puis, par continuité de la fonction exponentielle sur $\mathbb{R}$ (et en particulier en 1) :

$\lim_{n \rightarrow +\infty} \exp \left[\ln \left(\left(1 + \frac{1}{n}\right)^n \right) \right] = \exp(1) = e$. D'où : $\lim_{n \rightarrow +\infty} \left(1 + \frac{1}{n} \right)^n = e$.

Eh oui... Si vous trouvez ça peu intuitif, voire bizarre, n'hésitez pas à regarder, à la calculatrice, $\left(1 + \frac{1}{1000} \right)^{1000}$. Et voyez si ça vous semble plus se rapprocher de e ou de 1...

Enfin, par produit de limites : $\lim_{n \rightarrow +\infty} \frac{a_{n+1}}{a_n} = ye$

3) Supposons : $y < e^{-1}$ (ce qui revient à supposer $ye < 1$)

Comment, à partir de la limite obtenue précédemment, établir ce qui est suggéré par l'énoncé ? Cet entier N à partir duquel une certaine inégalité doit être vraie ne vous rappelle-t-il pas quelque chose ?

La suite $\left(\frac{a_{n+1}}{a_n} \right)$ converge vers ye . Ce qui se traduit ainsi :

$\forall \epsilon > 0, \exists N_\epsilon \in \mathbb{N}^*, \forall n \geq N_\epsilon, \left| \frac{a_{n+1}}{a_n} - ye \right| < \epsilon$

Dit autrement : $\forall \epsilon > 0, \exists N_\epsilon \in \mathbb{N}^*, \forall n \geq N_\epsilon, ye - \epsilon < \frac{a_{n+1}}{a_n} < ye + \epsilon$

J'ai écrit « $N_\epsilon \in \mathbb{N}^$ » et pas « $N_\epsilon \in \mathbb{N}$ » comme dans le rappel en début d'énoncé, parce qu'ici, la suite $\left(\frac{a_{n+1}}{a_n} \right)$ est définie à partir du rang $n = 1$*

Certes, mais que faire de tout ça ? La définition de la convergence commence par : $\forall \epsilon > 0...$ Autrement dit, nous pouvons appliquer la propriété pour n'importe quel $\epsilon > 0...$

En particulier, en choisissant $\epsilon = \frac{1 - ye}{2}$ (qui est bien strictement positif car $ye < 1$) :

il existe $N \in \mathbb{N}^*$ tel que : $\forall n \geq N, ye - \frac{1 - ye}{2} < \frac{a_{n+1}}{a_n} < ye + \frac{1 - ye}{2}$

J'ai pris la liberté de l'appeler N et pas $N_{\frac{1-ye}{2}}$...

Nous avons donc en particulier : $\forall n \geq N, \frac{a_{n+1}}{a_n} < ye + \frac{1-ye}{2}$

Bon. Nous sommes parvenus à établir l'indication. Mais qu'en faire ? D'ailleurs, ce $ye + \frac{1-ye}{2}$, à quoi bon ?

Par ailleurs : $ye + \frac{1-ye}{2} = \frac{2ye + 1 - ye}{2} = \frac{ye + 1}{2} < \frac{1+1}{2}$. Donc $ye + \frac{1-ye}{2} < 1$

A partir du rang N , $\frac{a_{n+1}}{a_n}$ est toujours inférieur à cette constante $\frac{ye + 1}{2}$ strictement inférieure à 1...

Posons $q = \frac{ye + 1}{2}$. Nous avons établi : $\forall n \geq N, \frac{a_{n+1}}{a_n} < q$

Autrement dit, puisque $a_n > 0 : \forall n \geq N, a_{n+1} < qa_n < a_n$

(a_n) est donc (s)trictement décroissante) à partir du rang N . Ah, si elle était minorée...

Et, pour tout $n \in \mathbb{N}$, $a_n > 0$. Donc (a_n) est minorée par 0.

Le théorème de convergence monotone nous permet de conclure que (a_n) converge vers un réel $l \geq 0$.

Par passage à la limite, les inégalités se conservent au sens large. $\geq$ a donc remplacé $>$

De plus : $\forall n \geq N, a_{n+1} < qa_n$. Un passage à la limite ($n \rightarrow +\infty$) dans cette inégalité fournit : $l \leq ql$

Passage à la limite permis car les deux membres a_{n+1} et qa_n possèdent bien une limite finie en $+\infty$. Rappelons aussi que si $a_n \xrightarrow{n \rightarrow +\infty} l$, alors on a aussi : $a_{n+1} \xrightarrow{n \rightarrow +\infty} l$

Même suite avec un simple décalage d'indice, donc même limite.

Puis : $l - ql \leq 0$, c'est-à-dire : $l(1 - q) \leq 0$. En divisant par $1 - q$, qui est strictement négatif, nous obtenons : $l \leq 0$. Enfin, $l = 0$.

Nous avons bien établi que si $y < e^{-1}$, alors $\lim_{n \rightarrow +\infty} a_n = 0$.

À partir de : « $\forall n \geq N, a_{n+1} < qa_n$ », une autre méthode était possible pour établir la convergence de (a_n) vers 0. On pouvait montrer, par une récurrence relativement simple (mais y penser ne l'était pas forcément...) : $\forall n \geq N, 0 < a_n \leq q^{n-N} a_N$. Une fois cet encadrement établi, comme $0 \leq q < 1$, nous savons : $\lim_{n \rightarrow +\infty} q^{n-N} a_N = 0$, et le théorème des

gendarmes nous permet enfin de conclure : $\lim_{n \rightarrow +\infty} a_n = 0$

Voici [un lien vers une vidéo d'exercice corrigé](#) où cette méthode (à peu de choses près) est employée pour montrer qu'une certaine suite diverge vers $+\infty$

4) On souffle un peu avec une question plus classique au sortir de la Terminale.

f est dérivable sur $\mathbb{R}_+$ par composée et produit de fonctions dérivables. De plus, pour tout réel positif u : $f'(u) = e^{-u} + u \times (-e^{-u}) = (1-u)e^{-u}$, qui est du signe de $1-u$ car $e^{-u} > 0$. On obtient le tableau de signe suivant pour $f'(x)$, ainsi que le tableau de variations de f :

x	0	1	$+\infty$
$f'(x)$	+	0	-
f	0	e^{-1}	0

f est strictement croissante sur $[0 ; 1]$ et strictement décroissante sur $[1 ; +\infty[$.

$f(0) = 0e^{-0} = 0$, $f(1) = 1e^{-1} = e^{-1}$, et $f(x) = \frac{x}{e^x} \xrightarrow{x \rightarrow +\infty} 0$ par croissance comparée.

Plus précisément, en utilisant la croissance comparée : $\frac{e^x}{x} \xrightarrow{x \rightarrow +\infty} +\infty$, puis par quotient de limites (« $\frac{1}{+\infty} = 0$ »). Bon, ici, ils ne demandent pas explicitement le tableau de variations complet (avec les valeurs / limites) mais comme elles sont relativement simples à obtenir. Et puis, qui sait, peut-être nous serviront-elles...

5) Pour tout entier naturel non nul n , posons $u_n = \frac{n^{n+1}}{n!} I_n$

$$\forall n \in \mathbb{N}^*, u_n = \frac{n^{n+1}}{n!} I_n = \frac{n^{n+1}}{n!} \times \int_0^b (te^{-t})^n dt = \frac{n^{n+1}}{n!} \times \int_0^b (f(t))^n dt$$

f est positive sur $\mathbb{R}_+$, et en particulier sur $[0 ; b]$. Donc : $\forall t \in [0 ; b], (f(t))^n \geq 0$

Par positivité de l'intégrale, $I_n \geq 0$, puis $u_n \geq 0$

Maintenant, si j'arrivais à majorer intelligemment u_n ...

$b \in]0 ; 1[$ donc : $\forall t \in [0 ; b], 0 \leq t \leq b < 1$

De plus, f est croissante sur $[0 ; 1]$ (cf tableau de variations)

Donc : $\forall t \in [0 ; b], f(t) \leq f(b)$

Puis, par croissance de la fonction $x \mapsto x^n$ sur $\mathbb{R}_+$ (intervalle dans lequel se situent bien $f(t)$ et $f(b)$) : $(f(t))^n \leq (f(b))^n$

Il s'ensuit, par croissance de l'intégrale : (les bornes étant dans le bon sens)

$$\int_0^b (f(t))^n dt \leq \int_0^b (f(b))^n dt$$

La variable d'intégration est t . Dans l'intégrale, $(f(b))^n$ est constante. Rappelons que pour a, b, K réels, $\int_a^b K dt = (b-a)K$ (ce qui se retrouve en primitivant : $[Kt]_a^b = Kb - Ka$)

$$\text{Puis : } \int_0^b (f(t))^n dt \leq b \times (f(b))^n$$

Par suite : $0 \leq u_n \leq \frac{n^{n+1}}{n!} \times b \times (f(b))^n$ Si je m'arrangeais avec ça pour utiliser 3)...

$$\text{Or, } 0 < b < 1 \text{ donc (comme } \frac{n^{n+1}}{n!} \times (f(b))^n) : \frac{n^{n+1}}{n!} \times b \times (f(b))^n \leq \frac{n^{n+1}}{n!} \times (f(b))^n$$

$$\text{Nous avons donc établi : } \forall n \in \mathbb{N}^*, 0 \leq u_n \leq \frac{n^{n+1}}{n!} \times (f(b))^n$$

Et, par stricte croissance de f sur $[0; 1]$, $f(0) < f(b) < f(1)$. Autrement dit : $0 < f(b) < e^{-1}$.

$$\text{En posant } y = f(b), \text{ on peut donc déduire de 3) : } \lim_{n \rightarrow +\infty} \frac{n^{n+1}}{n!} \times (f(b))^n = 0.$$

Le résultat de 3) est utilisable pour tout réel y strictement positif tel que $y < e^{-1}$. On est donc bien dans les clous en prenant $y = f(b)$. Soit dit en passant, c'est pour garder cela que j'ai temporisé dans mes majoration : plutôt que de majorer $f(b)$ par $f(1)$ dès le début, je l'ai gardé, parce que précisément, j'avais besoin d'un y strictement inférieur à e^{-1}

Le théorème des gendarmes permet enfin de conclure : $\lim_{n \rightarrow +\infty} u_n = 0$.

$$\text{Autrement dit : } \boxed{\lim_{n \rightarrow +\infty} \frac{n^{n+1}}{n!} I_n = 0}$$

Ouf! Bien qu'il ne mobilise que des prérequis de Terminale, cet exercice n'est pas évident du tout au sortir de cette année. En effet, il s'appuie sur un certain nombre de réflexes (de majoration, notamment) pas encore acquis, mais qui vous sembleront plus naturels dans les mois à venir.

Exercice 5

*Ma feuille est un chaos car ma plume combat
une intégrale en haut, une intégrale en bas.*

Énoncé : (temps conseillé : 30 min) (****) d'après Mines-Ponts 2014 PC Maths 2

On admet le résultat suivant :

si $[a; b]$ est un segment de $\mathbb{R}$, et si f est une fonction continue sur $[a; b]$, alors f est bornée et atteint ses bornes sur $[a; b]$. C'est le théorème des bornes atteintes.

Autrement dit, f admet un minimum et un maximum sur $[a; b]$, c'est-à-dire qu'il existe $x_0, x_1 \in [a; b]$ tels que : $\forall x \in [a; b], f(x_0) \leq f(x) \leq f(x_1)$.

Soit f une fonction continue sur $\mathbb{R}$, et soit ω une fonction continue sur $\mathbb{R}$ à valeurs dans $\mathbb{R}_+^*$.

Pour $x \in \mathbb{R}^*$, on pose $g(x) = \frac{1}{\int_0^x \omega(t) dt} \times \int_0^x f(t)\omega(t) dt$

1) Montrer que g est bien définie sur $\mathbb{R}^*$

2) Déterminer $\lim_{x \rightarrow 0} g(x)$

Remarques sur l'énoncé :

Juste au cas où, dans la traduction du théorème des bornes atteintes par :
« $\exists x_0, x_1 \in \mathbb{R}, \forall x \in [a; b], f(x_0) \leq f(x) \leq f(x_1)$ », le minimum est $f(x_0)$ et le maximum est $f(x_1)$ (à ne pas confondre avec x_0 et x_1 , qui sont des points de $[a; b]$ où respectivement, le minimum et le maximum sont atteints)

Correction de l'exercice 5 :

1) Les fonctions f et ω sont continues sur $\mathbb{R}$, donc par produit, la fonction $t \mapsto f(t)\omega(t)$ aussi. En particulier, pour tout réel non nul x , ces fonctions sont continues sur le segment $[0; x]$ (si $x > 0$) ou $[x; 0]$ (si $x < 0$). Les intégrales $\int_0^x \omega(t) dt$ et $\int_0^x f(t)\omega(t) dt$ sont bien définies.

Rappelons au passage que : $\int_0^x \omega(t) dt = - \int_x^0 \omega(t) dt$

Reste à établir le fait que $\int_0^x \omega(t) dt \neq 0$, pour que $g(x)$ soit bien défini. Que savons-nous de potentiellement utile sur la fonction ω ? Qu'elle est à valeurs dans $\mathbb{R}_+^*$... Autrement dit, que pour tout réel t , $\omega(t) > 0$...

Il y a aussi ce théorème des bornes atteintes donné par l'énoncé, et applicable à toute fonction continue sur un segment...

Pour tout réel $x > 0$: ω est continue sur $\mathbb{R}$, donc en particulier sur $[0; x]$. D'après le théorème des bornes atteintes, ω admet un minimum sur $[0; x]$. Autrement dit, il existe un réel $c \in [0; x]$ tel que : $\forall t \in [0; x], \omega(c) \leq \omega(t)$.

Puis, par croissance de l'intégrale (bornes dans le bon sens) : $\int_0^x \omega(c) dt \leq \int_0^x \omega(t) dt$
 $\omega(c)$ étant constant vis-à-vis de t , nous obtenons : $x \times \omega(c) \leq \int_0^x \omega(t) dt$

Revenez ici si vous voulez vous rafraîchir la mémoire sur l'intégrale d'une constante.

De plus, ω est à valeurs dans $\mathbb{R}_+^*$ donc $\omega(c) > 0$. Notre ticket de sortie !

D'où : $\forall x > 0, 0 < x \times \omega(c) \leq \int_0^x \omega(t) dt$

Donc $\int_0^x \omega(t) dt > 0$ et a fortiori : $\int_0^x \omega(t) dt \neq 0$. Enfin, $g(x)$ est bien défini.

En conclusion : g est bien définie sur $\mathbb{R}_+^*$

De même, pour tout réel $x < 0$: ω est continue sur $[x; 0]$. D'après le théorème des bornes atteintes, ω admet un minimum sur $[x; 0]$. Autrement dit, il existe un réel $c \in [x; 0]$ tel que : $\forall t \in [x; 0], \omega(c) \leq \omega(t)$.

Puis, par croissance de l'intégrale (bornes dans le bon sens) : $\int_x^0 \omega(c) dt \leq \int_x^0 \omega(t) dt$

Nous obtenons donc : $-x \times \omega(c) \leq \int_x^0 \omega(t) dt$, c'est-à-dire : $-x \times \omega(c) \leq - \int_0^x \omega(t) dt$

Il s'ensuit : $x \times \omega(c) \geq \int_0^x \omega(t) dt$. Or, $\omega(c) > 0$ et $x < 0$.

Donc : $x \times \omega(c) < 0$ et, a fortiori $\int_0^x \omega(t) dt < 0$. Enfin, dans ce cas aussi : $\int_0^x \omega(t) dt \neq 0$, et $g(x)$ est bien défini pour x appartenant à $\mathbb{R}_-^*$.

En conclusion : g est bien définie sur $\mathbb{R}^*$

Vous pourrez bientôt, sans passer par le théorème des bornes atteintes, utiliser directement le résultat suivant : si h est une fonction continue sur un segment $[a; b]$ ($a < b$), et si, pour tout réel $t \in [a; b]$, $h(t) \geq 0$, alors : $\int_a^b h(t) dt = 0 \implies \forall t \in [a; b], h(t) = 0$ (L'implication dans l'autre sens étant vraie et évidente)

Autrement dit, pour une fonction h fonction continue et positive sur $[a; b]$: si $\int_a^b h(t) dt = 0$, alors h est identiquement nulle sur $[a; b]$.

Dans notre cas, ω , continue et positive sur $[0; x]$, et n'étant pas identiquement nulle sur $[0; x]$ (mieux encore : ω ne s'annule en fait jamais sur $[0; x]$), $\int_0^x \omega(t) dt$ ne peut être nulle.

*En disant cela, on se servirait de la **contraposée** de l'implication :*

$$\int_a^b h(t) dt = 0 \implies h \text{ identiquement nulle sur } [a; b]$$

*Lorsque p et q sont deux assertions mathématiques, l'implication $p \implies q$ est équivalente à sa **contraposée**, à savoir l'implication : $\text{Non}(q) \implies \text{Non}(p)$*

Il semble en effet relativement intuitif que si q est une condition nécessaire à p , l'absence de q implique l'absence de p .

Une erreur courante est de considérer que $p \implies q$ entraînerait que $\text{Non}(p) \implies \text{Non}(q)$

Ca ne tient pas la route : « s'il y a un incendie, alors il y a de la fumée » n'entraîne pas « s'il y a de la fumée, alors il y a un incendie ».

Par contre, « s'il y a un incendie, alors il y a de la fumée » entraîne bien « s'il n'y a pas de fumée, alors il n'y a pas d'incendie »

Dans notre situation, la contraposée de l'implication :

$$\int_a^b h(t) dt = 0 \implies h \text{ identiquement nulle sur } [a; b]$$

est l'implication : h non identiquement nulle sur $[a; b] \implies \int_a^b h(t) dt \neq 0$

2) On nous demande maintenant la limite du quotient d'intégrales $g(x)$ lorsque x tend vers 0. Intuitivement, $\int_0^x f(t)\omega(t) dt$ et $\int_0^x \omega(t) dt$ tendraient tous deux vers 0, ce qui nous ferait une belle forme indéterminée. Que faire ? Nous ne pouvons « calculer » explicitement ces deux intégrales, puisque nous ne connaissons ni f ni ω . Mais nous savons qu'elles sont continues sur $\mathbb{R}$... Ce qui nous donne le droit d'introduire des primitives, qui nous permettraient d'exprimer $g(x)$ sans le signe intégral.

Les fonctions ω et $t \mapsto f(t)\omega(t)$ sont continues sur $\mathbb{R}$. Elles admettent donc des primitives sur $\mathbb{R}$. Soient F une primitive de ω et G une primitive de $t \mapsto f(t)\omega(t)$ sur $\mathbb{R}$.

Par définition : $\forall x \in \mathbb{R}^*$, $\int_0^x \omega(t) dt = F(x) - F(0)$ et $\int_0^x f(t)\omega(t) dt = G(x) - G(0)$

D'où : $\forall x \in \mathbb{R}^*$, $g(x) = \frac{G(x) - G(0)}{F(x) - F(0)}$ *Quid de la limite de ce quotient ?*

Si nous faisons apparaître des taux d'accroissement, pour ensuite procéder comme dans l'exercice 4 ?

Donc : $\forall x \in \mathbb{R}^*$, $g(x) = \frac{G(x) - G(0)}{x - 0} \times \frac{x - 0}{F(x) - F(0)}$

Division-multiplication, pour faire apparaître des quantités qui m'arrangent.

Or, G est dérivable sur $\mathbb{R}$ (et en particulier en 0), car c'est une primitive de la fonction $t \mapsto f(t)\omega(t)$ sur $\mathbb{R}$. Nous savons plus précisément : $\forall a \in \mathbb{R}$, $G'(a) = f(a)\omega(a)$

Cela nous permet d'établir : $\lim_{x \rightarrow 0} \frac{G(x) - G(0)}{x - 0} = G'(0) = f(0)\omega(0)$

De même, F est dérivable sur $\mathbb{R}$ (et en particulier en 0), car c'est une primitive de la fonction ω sur $\mathbb{R}$. Et : $\forall a \in \mathbb{R}$, $F'(a) = \omega(a)$

Cela nous permet d'établir : $\lim_{x \rightarrow 0} \frac{F(x) - F(0)}{x - 0} = F'(0) = \omega(0)$

Or, $\omega(0) \neq 0$, car ω est à valeurs dans $\mathbb{R}_+^*$. Donc, par quotient de limites (*passage à l'inverse, tout simplement*) :

$\lim_{x \rightarrow 0} \frac{x - 0}{F(x) - F(0)} = \frac{1}{\omega(0)}$. Puis, par produit de limites : $\lim_{x \rightarrow 0} g(x) = f(0)\omega(0) \times \frac{1}{\omega(0)}$

Enfin, nous avons établi : $\lim_{x \rightarrow 0} g(x) = f(0)$

Exercice 6

*Je vous choisis ce jour ma plume la plus belle
dûment aiguisée pour mes sommations d'Abel.*

Énoncé : (temps conseillé : 50 min) (****) d'après Centrale 2013 PC Maths 1

On admet que pour toute suite réelle (u_n) , si $(\sum_{k=0}^n |u_k|)_{n \in \mathbb{N}}$ converge, alors $(\sum_{k=0}^n u_k)_{n \in \mathbb{N}}$ converge aussi.

Soit (a_n) une suite réelle décroissante qui converge vers 0, et soit (b_n) une suite réelle telle que la suite (B_n) définie pour tout $n \in \mathbb{N}$ par $B_n = \sum_{k=0}^n b_k = b_0 + b_1 + \dots + b_n$ est bornée

Soit (S_n) la suite définie pour tout $n \in \mathbb{N}$ par $S_n = \sum_{k=0}^n a_k b_k = a_0 b_0 + a_1 b_1 + \dots + a_n b_n$

On cherche à montrer que (S_n) converge.

- 1) Montrer que pour tout entier $n \geq 1$, $S_n = a_n B_n + \sum_{k=0}^{n-1} (a_k - a_{k+1}) B_k$
- 2) Montrer que $(a_n B_n)_{n \in \mathbb{N}}$ converge.
- 3) Montrer que $(\sum_{k=0}^{n-1} (a_k - a_{k+1}) B_k)_{n \in \mathbb{N}^*}$ converge.
- 4) Conclure.

Remarques sur l'énoncé :

Le résultat de convergence admis (si $(\sum_{k=0}^n |u_k|)_{n \in \mathbb{N}}$ converge, alors $(\sum_{k=0}^n u_k)_{n \in \mathbb{N}}$ converge aussi), traité ici sous l'angle des suites, mais que vous aborderez l'an prochain sous l'angle des séries, revient à dire que la convergence absolue d'une série numérique entraîne sa convergence. En voici [une démonstration en vidéo](#).

Rappelons, à toutes fins utiles, que pour $n \in \mathbb{N}$, $\sum_{k=0}^n u_k$ se lit « somme pour k allant de 0 à n des u_k », et désigne donc tout simplement $u_0 + u_1 + \dots + u_n$.

Plus généralement, pour $p, n \in \mathbb{N}$ tels que $p \leq n$, $\sum_{k=p}^n a_k = a_p + a_{p+1} + \dots + a_n$

Voici [un lien vers une playlist de vidéos courtes](#), que je vous conseille vivement de consulter pour vous familiariser avec le signe $\sum$, en particulier si vous ne l'avez pas suffisamment manipulé en Terminale. Certaines manipulations de sommes dans la correction,

que vous pourrez certes trouver intuitives sans ce visionnage préalable, vous sembleront plus naturelles après avoir regardé ladite playlist.

Entre autres propriétés intuitives qui vous serviront peut-être pour résoudre ce problème : lorsque C est une constante (indépendante de k), on a : $\sum_{k=0}^n C a_k = C \sum_{k=0}^n a_k$ (linéarité)
C'est une simple factorisation : trivialement, $C a_1 + C a_2 + \dots + C a_n = C(a_1 + a_2 + \dots + a_n)$

La linéarité permet aussi d'écrire $\sum_{k=0}^n a_k + b_k = \sum_{k=0}^n a_k + \sum_{k=0}^n b_k$

Ajoutons cette précision sur les changements d'indice, qui pourra vous être utile : pour $n \geq 1$, la somme $\sum_{k=1}^n u_{k-1}$ est égale à la somme $\sum_{k=0}^{n-1} u_k$.

En effet : $\sum_{k=1}^n u_{k-1} = u_{1-1} + u_{2-1} + \dots + u_{n-1} = u_0 + u_1 + \dots + u_{n-1} = \sum_{k=0}^{n-1} u_k$

De manière plus formelle, en partant de $\sum_{k=1}^n u_{k-1}$, on pose le changement d'indice $j = k - 1$. Puisque k allait de 1 jusqu'à n , j va de 0 jusqu'à $n - 1$. Et u_{k-1} est remplacé par u_j . On se retrouve donc avec $\sum_{j=0}^{n-1} u_j$, qui n'est autre que $\sum_{k=0}^{n-1} u_k$.

Les variables k et j dans ces deux dernières sommes sont en effet « muettes ». Elles ne sont là que pour indiquer le parcours de l'indice de sommation. De la même manière que $\int_1^2 \ln(t) dt = \int_1^2 \ln(u) du$.

Mais alors pourquoi nous être embêté à poser ce fameux j ? Pour ne pas nous emmêler les pinceaux entre l'ancien et le nouvel indice de sommation.

Pour plus de précisions sur les changements d'indice, n'hésitez pas à consulter [cette vidéo](#). Oui, oui, encore une vidéo de la playlist dont je vous parle sans arrêt, mais elle est bien, vous verrez.

Enfin, la transformation de sommes que cet exercice nous fait effectuer (en question 1) est un grand classique, connu sous le nom de transformation d'Abel. En ce sens, l'annale de Centrale était plus un prétexte pour vous la présenter qu'autre chose. D'ailleurs, par rapport au sujet originel, il y a eu une petite modification au niveau des indices (k démarrant à 0 plutôt qu'à 1), et j'ai pris (b_n) réelle alors que pour eux, elle était complexe (ce qui ne change pas grand-chose au niveau de la résolution).

Correction de l'exercice 6 :

1) Pour tout entier $n \geq 1$, $S_n = \sum_{k=0}^n a_k b_k$. Comment faire entrer les B_k dans la danse ?

De plus, pour tout $k \in \mathbb{N}^*$, $b_k = B_k - B_{k-1}$.

En effet : $B_k - B_{k-1} = b_0 + b_1 + \dots + b_{k-1} + b_k - (b_0 + b_1 + \dots + b_{k-1}) = b_k$

Pas si simple d'y penser la première fois... Si chaque B_k est la somme cumulée des b_j pour j de 0 à k , chaque b_k est la différence entre deux sommes cumulées successives.

Ecrit avec Σ , cela donne : $B_k - B_{k-1} = \sum_{j=0}^k b_j - \sum_{j=0}^{k-1} b_j = \sum_{j=0}^{k-1} b_j + b_k - \sum_{j=0}^{k-1} b_j = b_k$

Attention, dans l'expression $B_k = \sum_{j=0}^k b_j$, je ne pouvais pas choisir k comme indice muet

dans la somme, et écrire une folie du style $\sum_{k=0}^k b_k$. k est la borne du haut de cette somme, bien fixée, et ne peut représenter l'indice qui se balade de 0.. à k . D'où le choix d'une lettre différente (j par exemple).

Dernière considération avant de reprendre la correction. Nous avons bien écrit : pour tout $k \in \mathbb{N}^$, $b_k = B_k - B_{k-1}$. Pour tout k entier naturel non nul. Pas pour $k = 0$, parce que ça donnerait $b_0 = B_0 - B_{-1}$, et on n'a pas défini de B_{-1} . Remarque, on pourrait poser $B_{-1} = 0$, mais je préfère « expulser » le premier terme de la somme S_n comme suit :*

$$\begin{aligned} \text{Pour tout entier } n \geq 1, S_n &= a_0 b_0 + \sum_{k=1}^n a_k b_k = a_0 b_0 + \sum_{k=1}^n a_k (B_k - B_{k-1}) \\ &= a_0 b_0 + \sum_{k=1}^n (a_k B_k - a_k B_{k-1}) = a_0 b_0 + \sum_{k=1}^n a_k B_k - \sum_{k=1}^n a_k B_{k-1} \text{ par linéarité.} \end{aligned}$$

Vu ce qu'on nous demande d'obtenir, il serait judicieux de transformer la seconde somme.

Le changement d'indice $j = k - 1$ (cf remarques sur l'énoncé) fournit :

$$\sum_{k=1}^n a_k B_{k-1} = \sum_{j=0}^{n-1} a_{j+1} B_j = \sum_{k=0}^{n-1} a_{k+1} B_k. \text{ Donc } S_n = a_0 b_0 + \sum_{k=1}^n a_k B_k - \sum_{k=0}^{n-1} a_{k+1} B_k$$

J'aimerais bien jouer de la linéarité pour regrouper les deux dernières sommes, mais les indices ne partent pas de la même valeur et ne s'arrêtent pas à la même valeur... Et je voudrais une somme avec k allant de 0 à $n - 1$...

Puis : $S_n = a_0 b_0 + \sum_{k=1}^n a_k B_k - \sum_{k=0}^{n-1} a_{k+1} B_k = a_0 b_0 + \sum_{k=0}^{n-1} a_k B_k - a_0 B_0 + a_n B_n - \sum_{k=0}^{n-1} a_{k+1} B_k$

À la dernière étape, nous avons juste remplacé $\sum_{k=1}^n a_k B_k$ par $\sum_{k=0}^{n-1} a_k B_k - a_0 B_0 + a_n B_n$

Tout simplement parce que, par rapport à $\sum_{k=0}^{n-1} a_k B_k$, $\sum_{k=1}^n a_k B_k$ a le terme correspondant à $k = n$ en plus, et celui correspondant à $k = 0$ en moins.

Remarquons aussi : $B_0 = b_0$ Ben oui, une « somme » où il n'y a que b_0 ...

Par suite : $S_n = a_0 b_0 + \sum_{k=0}^{n-1} a_k B_k - a_0 b_0 + a_n B_n - \sum_{k=0}^{n-1} a_{k+1} B_k = \sum_{k=0}^{n-1} a_k B_k + a_n B_n - \sum_{k=0}^{n-1} a_{k+1} B_k$

Enfin : pour tout entier $n \geq 1$, $S_n = a_n B_n + \sum_{k=0}^{n-1} (a_k - a_{k+1}) B_k$

2) Par hypothèse, la suite (B_n) est bornée. Il existe donc deux réels m et M tels que : $\forall n \in \mathbb{N}, m \leq B_n \leq M$. En multipliant par $a_n \geq 0$ (car (a_n) tend en décroissant vers 0) :

$$\forall n \in \mathbb{N}, m a_n \leq a_n B_n \leq M a_n$$

Or : $\lim_{n \rightarrow +\infty} a_n = 0$. Donc : $\lim_{n \rightarrow +\infty} m a_n = 0$ et $\lim_{n \rightarrow +\infty} M a_n = 0$

Le théorème des gendarmes nous permet de conclure que $(a_n B_n)$ converge vers 0.

Bon, l'énoncé demandait juste « converge », mais le théorème des gendarmes nous donne aussi la limite, alors...

Vous avez vu en Terminale qu'une suite réelle est dite bornée si et seulement si elle est à la fois minorée et majorée.

Même si je ne m'en suis pas servi ici, habituez-vous à cette autre formulation, qui pourra être plus commode par moments, et qui a l'avantage d'être aussi valable sur $\mathbb{C}$ (ça n'a pas de sens a priori de dire qu'une suite complexe est majorée ou minorée) :

$$(u_n) \text{ bornée} \iff \exists A \geq 0, \forall n \in \mathbb{N}, |u_n| \leq A$$

Autrement dit, (u_n) est bornée si et seulement si $(|u_n|)$ est majorée.

En effet, si (u_n) est bornée, avec votre définition de Terminale, il existe deux réels m et M tels que pour tout $n \in \mathbb{N}$, $m \leq u_n \leq M$. En notant A la plus grande des valeurs absolues de m et M - autrement dit, $A = \max(|m|, |M|)$ - nous avons bien : $\forall n \in \mathbb{N}, |u_n| \leq A$

Réciproquement, s'il existe un réel positif A tel que : $\forall n \in \mathbb{N}, |u_n| \leq A$, alors :

$\forall n \in \mathbb{N}, -A \leq u_n \leq A$. (u_n) est donc minorée (par $-A$) et majorée (par A), donc bornée.

J'aurais pu user de cette reformulation dans une correction alternative, qui ne se sert pas

de la positivité de (a_n) :

Par hypothèse, la suite (B_n) est bornée. Il existe donc un réel $A \geq 0$ tel que : $\forall n \in \mathbb{N}, |B_n| \leq A$.

En multipliant par $|a_n| \geq 0 : \forall n \in \mathbb{N}, 0 \leq |a_n| |B_n| \leq A |a_n|$. C'est-à-dire : $0 \leq |a_n B_n| \leq A |a_n|$

Or : $\lim_{n \rightarrow +\infty} |a_n| = 0$. Donc : $\lim_{n \rightarrow +\infty} A |a_n| = 0$. Le théorème des gendarmes nous permet de conclure que $(|a_n B_n|)$ converge vers 0, ce qui revient à dire que $(a_n B_n)$ converge vers 0.

3) Les choses se compliquent... Convergence d'une suite définie à l'aide d'une somme pas commode. À ce sujet, je crois me souvenir qu'on m'a fait admettre un truc...

Posons, pour tout $n \in \mathbb{N}^*$: $v_n = \sum_{k=0}^{n-1} (a_k - a_{k+1}) B_k$ et $w_n = \sum_{k=0}^{n-1} |(a_k - a_{k+1}) B_k|$

Nous cherchons à montrer que (v_n) converge, et un résultat donné par l'énoncé nous permet d'affirmer que si (w_n) converge, alors (v_n) converge aussi.

Il suffirait donc de montrer que (w_n) converge pour parvenir à nos fins.

Par décroissance de (a_n) , nous avons : pour tout entier naturel k , $a_k - a_{k+1} \geq 0$.

Donc $|a_k - a_{k+1}| = a_k - a_{k+1}$. D'où : $\forall n \in \mathbb{N}^* : w_n = \sum_{k=0}^{n-1} |a_k - a_{k+1}| |B_k| = \sum_{k=0}^{n-1} (a_k - a_{k+1}) |B_k|$

Toujours ça de pris de se débarrasser d'une valeur absolue encombrante...

(B_k) étant bornée, nous savons : $\exists A \geq 0, \forall k \in \mathbb{N}, |B_k| \leq A$. Habituez-vous, disais-je...

Puis : $\forall n \in \mathbb{N}^*, \forall k \in \llbracket 0; n-1 \rrbracket, (a_k - a_{k+1}) |B_k| \leq (a_k - a_{k+1}) \times A$

$(a_k - a_{k+1})$ étant bien positif

La notation $\llbracket 0; n-1 \rrbracket$ (intervalle fermé avec une double barre) désigne juste l'ensemble des entiers compris au sens large entre 0 et $n-1$. L'ensemble $\llbracket 0; n-1 \rrbracket$ est donc l'ensemble $\{0, 1, \dots, n-1\}$ (notation moins formelle à cause des trois petits points)

En sommant, nous obtenons : $\forall n \in \mathbb{N}^*, \sum_{k=0}^{n-1} (a_k - a_{k+1}) |B_k| \leq \sum_{k=0}^{n-1} A \times (a_k - a_{k+1})$

Puis, par linéarité : $\forall n \in \mathbb{N}^*, w_n \leq A \sum_{k=0}^{n-1} (a_k - a_{k+1})$

Et alors, où allons-nous avec ça ? - Et alors, ceux qui auront eu la curiosité de consulter la playli.. - ARRÊTE AVEC TA FICHUE PLAYLIST! - Bon, si vous voulez, on va vous faire voir un télescope en explicitant la somme. Blague à part, j'ai déjà croisé ça (même si rarement) dans des exercices de type bac, où l'on attendait du candidat d'explicitier la somme pour y voir plus clair.

Or :
$$\sum_{k=0}^{n-1} (a_k - a_{k+1}) = a_0 - a_1 + a_1 - a_2 + \dots + a_{n-2} - a_{n-1} + a_{n-1} - a_n = a_0 - a_n$$

En effet, tous les termes se simplifient deux à deux sauf a_0 et $-a_n$

Une simplification peu formelle, à laquelle le programme de Terminale nous cantonne.

En toute rigueur, l'expression explicite que je donne de la somme (avec assez de termes pour que vous voyiez bien la simplification) n'est valable que pour $n > 3$. Mais l'égalité

$\sum_{k=0}^{n-1} (a_k - a_{k+1}) = a_0 - a_n$ qui suit est bien valable pour tout entier naturel n non nul. Et, si vous voulez vous en convaincre, n'hésitez pas à la vérifier pour $n = 1, 2$ et 2 .

Une utilisation du principe de télescopage nous aurait permis d'écrire directement :

$$\sum_{k=0}^{n-1} (a_k - a_{k+1}) = a_0 - a_{n-1+1} = a_0 - a_n$$

Nous avons donc montré : $\forall n \in \mathbb{N}^*, w_n \leq A(a_0 - a_n)$.

Et alors (bis) ? Et alors, on n'est pas loin d'établir que (w_n) est majorée... Mais un majorant ne doit pas dépendre de n .

De plus, pour tout $n \in \mathbb{N}^* : a_0 - a_n \leq a_0$ et $A \geq 0$ donc : $A(a_0 - a_n) \leq Aa_0$ puis : $w_n \leq Aa_0$

La suite (w_n) est donc majorée.

Si seulement elle était croissante, et si seulement ce n'était pas une galère à démontrer...

De plus : $\forall n \in \mathbb{N}^*, w_{n+1} - w_n = \sum_{k=0}^n |(a_k - a_{k+1})B_k| - \sum_{k=0}^{n-1} |(a_k - a_{k+1})B_k| = |(a_n - a_{n+1})B_n|$

Différence de deux sommes qui sont les mêmes à un terme près (celui correspondant à $k = n$, que la première somme possède en trop par rapport à la seconde)

Donc : $\forall n \in \mathbb{N}^*, w_{n+1} - w_n \geq 0$ $| (a_n - a_{n+1})B_n |$ est une valeur absolue...

La suite (w_n) est donc croissante.

Le théorème de convergence monotone nous permet de conclure que (w_n) est convergente.

Autrement dit : $\left(\sum_{k=0}^{n-1} (a_k - a_{k+1})B_k \right)_{n \in \mathbb{N}^*}$ converge.

4) D'après 2) et 3), les suites $(a_n B_n)_{n \in \mathbb{N}}$ et $\left(\sum_{k=0}^{n-1} (a_k - a_{k+1})B_k \right)_{n \in \mathbb{N}^*}$. Enfin, d'après 1), par somme, la suite (S_n) converge.

Et sa limite est la somme des limites des suites $(a_n B_n)_{n \in \mathbb{N}}$ et de $\left(\sum_{k=0}^{n-1} (a_k - a_{k+1})B_k \right)_{n \in \mathbb{N}^}$*

Exercice 7

*Attendez-vous encore à deux ou trois ratures
avant de trouver l'ordre de la quadrature.*

Énoncé : (temps conseillé : 40 min) (***) d'après Centrale 2021 PC Maths 2

Soit f une fonction continue sur l'intervalle $[0; 1]$. On cherche à approximer $\int_0^1 f(t) dt$ par une expression de la forme $I_n(f) = \sum_{j=0}^n \lambda_j f(x_j)$, où $n \in \mathbb{N}$, $\lambda_0, \lambda_1, \dots, \lambda_n$ sont des réels, et $x_0, x_1, \dots, x_n$ sont $n+1$ points distincts de $[0; 1]$. Une telle expression $I_n(f)$ est appelée formule de quadrature.

Etant donné un entier $m \in \mathbb{N}$, on note $\mathbb{R}_m[X]$ l'ensemble des polynômes à coefficients réels de degré inférieur ou égal à m .

On dit qu'une formule de quadrature $I_n(f)$ est exacte sur $\mathbb{R}_m[X]$ si :

$$\forall P \in \mathbb{R}_m[X], \int_0^1 P(t) dt = \sum_{j=0}^n \lambda_j P(x_j).$$

Enfin, on appelle ordre d'une formule de quadrature $I_n(f)$ le plus grand entier $m \in \mathbb{N}$ pour lequel la formule de quadrature $I_n(f)$ est exacte sur $\mathbb{R}_m[X]$.

1) Déterminer l'ordre de la formule de quadrature $I_0(f) = f(0)$

2) Faire de même avec la formule de quadrature $I_0(f) = f(\frac{1}{2})$

3) Déterminer les coefficients λ_0, λ_1 et λ_2 pour que la formule $I_2(f) = \lambda_0 f(0) + \lambda_1 f(\frac{1}{2}) + \lambda_2 f(1)$ soit exacte sur $\mathbb{R}_2[X]$. Que peut-on dire sur l'ordre de cette formule de quadrature ?

Remarques sur l'énoncé :

$\mathbb{R}[X]$ est l'ensemble des polynômes à coefficients réels et d'indéterminée X . Tout polynôme P donné de $\mathbb{R}[X]$ sera identifié à la fonction polynomiale associée.

Autrement dit, par exemple, on ne s'embêtera pas à faire la distinction entre l'objet $P = X^2 + 5X - 3$, polynôme de $\mathbb{R}[X]$ (on peut aussi noter $P(X) = X^2 + 5X - 3$), et la fonction polynomiale associée, définie sur $\mathbb{R}$ par $P(x) = x^2 + 5x - 3$. Vous ne faites normalement pas la distinction en Terminale. Vous la ferez sur le plan théorique en première année,

mais de nombreux exercices, problèmes de concours, annales, vous permettront de ne pas la faire. La même lettre (ici P) pourra désigner le polynôme et la fonction polynomiale associée. Mais mettez des grand X quand vous parlez du polynôme, et des petit x (ou autres lettres, minuscules en général) quand vous parlez de l'image que renvoie la fonction polynomiale P à un réel

Petit rappel sur les degrés des polynômes par le biais d'exemples (rappel qui dépasse un peu ce dont vous avez précisément besoin dans cet exercice) :

$X^3 - 2X + 1$ est de degré 3, et le coefficient du terme de degré 3 - ou coefficient dominant - est 1 (coefficient devant X^3). $2X - 1$ est de degré 1, et son coefficient dominant (coefficient devant X) est 2. $3 + 5X^2$ est de degré 2, et son coefficient dominant est 5

Le polynôme 7 est de degré 0. En fait, tout polynôme constant non nul est de degré 0. Et, par convention, le polynôme nul est de degré $-\infty$.

Un mot sur cette dernière convention, qui peut sembler bizarre mais qui est bien comode, en fait : si $\deg(\cdot)$ désigne le degré d'un polynôme, on veut que, pour tous polynômes P et Q , l'égalité $\deg(PQ) = \deg(P) + \deg(Q)$ soit vérifiée.

Dans le cas de deux polynômes non nuls, la puissance la plus grande obtenue par produit des deux polynômes correspond trivialement à la somme des plus grandes puissances pour chacun des deux polynômes :

Si $P = 3X^2 - 4X$ et $Q = X^3 - 2X^2 + 1$, alors $PQ = (3X^2 - 4X)(X^3 - 2X^2 + 1)$, et le terme de plus haut degré de PQ est $3X^2 \times X^3 = 3X^{2+3}$

Si le polynôme nul, que l'on notera P_0 , était de degré 0, on aurait eu, en prenant par exemple le polynôme $Q = X$: d'une part, $P_0Q = P_0$, et donc $\deg(P_0Q) = 0$. Et d'autre part, $\deg(P_0) + \deg(Q) = 0 + 1 = 1$, donc l'égalité souhaitée ($\deg(P_0Q) = \deg(P_0) + \deg(Q)$) n'est pas vérifiée. Par contre, en posant $\deg(P_0Q) = -\infty$, « $-\infty + 1 = -\infty$ » marche mieux...

Comme rappelé dans l'énoncé, pour tout entier $m \in \mathbb{N}$, on note $\mathbb{R}_m[X]$ l'ensemble des polynômes à coefficients réels de degré inférieur ou égal à m . En particulier, $\mathbb{R}_0[X]$ est l'ensemble des polynômes réels constants.

Remarquons aussi : $\forall n, p \in \mathbb{N}$, si $n \leq p$, alors $\mathbb{R}_n[X] \subset \mathbb{R}_p[X]$

Ben oui, par exemple, si P est de degré inférieur ou égal à 3, alors il est aussi de degré inférieur ou égal à 7...

Correction de l'exercice 7 :

1) Où est la somme dans cette formule de quadrature ? C'est un cas un peu trivial avec $n = 0$, $\lambda_0 = 0$ et $x_0 = 0$. Avec de telles valeurs, on retrouve l'expression $I_0(f) = \sum_{j=0}^0 \lambda_j f(x_j)$

On cherche le plus grand entier $m \in \mathbb{N}$ tel que : $\forall P \in \mathbb{R}_m[X], \int_0^1 P(t) dt = P(0)$

Comment faire ? Intuitivement, il paraît difficile que pour tout polynôme de $\mathbb{R}_m[X]$, la valeur de son intégrale de 0 à 1 coïncide avec sa valeur en 0... À moins de polynômes vraiment simples, et donc d'un m « petit » ...

Pour tout $P \in \mathbb{R}_0[X] : \exists a \in \mathbb{R}, \forall t \in \mathbb{R}, P(t) = a$. D'où : $\int_0^1 P(t) dt = \int_0^1 a dt = a(1 - 0) = a$

Et $P(0) = a$. Donc : $\forall P \in \mathbb{R}_0[X], \int_0^1 P(t) dt = P(0)$

La formule de quadrature $I_0(f) = f(0)$ est donc exacte sur $\mathbb{R}_0[X]$. Son ordre est donc supérieur ou égal à 0.

Oui, supérieur ou égal à 0. À ce stade, elle pourrait être exacte sur des $\mathbb{R}_m[X]$ avec m plus grand. Mais rien que pour $\mathbb{R}_1[X]$, je ne le sens pas trop...

Le polynôme $Q = X$ appartient à $\mathbb{R}_1[X]$. D'une part : $\int_0^1 Q(t) dt = \int_0^1 t dt = \left[\frac{t^2}{2} \right]_0^1 = \frac{1}{2}$

D'autre part : $Q(0) = 0 \neq \frac{1}{2}$. La formule de quadrature $I_0(f) = f(0)$ n'est donc pas exacte sur $\mathbb{R}_1[X]$. Il suffisait en effet d'un contre-exemple pour en casser l'exactitude.

0 est donc le plus grand entier naturel m pour lequel cette formule de quadrature est exacte sur $\mathbb{R}_m[X]$. Autrement dit, l'ordre de la formule de quadrature $I_0(f) = f(0)$ est 0.

2) Pour celle-ci, je vais être un peu plus gourmand d'entrée de jeu. Je vais voir si elle est exacte sur $\mathbb{R}_1[X]$. Si je vois que non, je reverrai mes prétentions à la baisse et me rabattrai sur $\mathbb{R}_0[X]$. Pourquoi pas regarder carrément sur $\mathbb{R}_2[X]$ ou plus ? Parce qu'un petit coup d'oeil à la question 3) me fait comprendre que cette dernière aurait peu d'intérêt si la formule de quadrature $I_0(f) = f\left(\frac{1}{2}\right)$ était déjà exacte sur $\mathbb{R}_2[X]$... (Pour la 3, $\lambda_0 = 0, \lambda_1 = 1$ et $\lambda_2 = 0$ conviendraient trivialement)

Pour tout $P \in \mathbb{R}_1[X] : \exists a, b \in \mathbb{R}, \forall t \in \mathbb{R}, P(t) = at + b$.

D'où : $\int_0^1 P(t) dt = \int_0^1 (at + b) dt = \left[a \frac{t^2}{2} + bt \right]_0^1 = \frac{1}{2} \times a + b$

Et $P\left(\frac{1}{2}\right) = \frac{1}{2} \times a + b$. Donc : $\forall P \in \mathbb{R}_1[X], \int_0^1 P(t) dt = P\left(\frac{1}{2}\right)$

La formule de quadrature $I_0(f) = f\left(\frac{1}{2}\right)$ est donc exacte sur $\mathbb{R}_1[X]$, et son ordre est supérieur ou égal à 1.

Le polynôme $Q = X^2$ appartient à $\mathbb{R}_2[X]$. D'une part : $\int_0^1 Q(t) dt = \int_0^1 t^2 dt = \left[\frac{t^3}{3} \right]_0^1 = \frac{1}{3}$

D'autre part : $Q\left(\frac{1}{2}\right) = \left(\frac{1}{2}\right)^2 = \frac{1}{4} \neq \frac{1}{3}$. La formule de quadrature $I_0(f) = f\left(\frac{1}{2}\right)$ n'est donc pas exacte sur $\mathbb{R}_2[X]$.

Enfin, l'ordre de la formule de quadrature $I_0(f) = f\left(\frac{1}{2}\right)$ est 1.

3) On cherche trois réels λ_0, λ_1 et λ_2 tels que :

$$\forall P \in \mathbb{R}_2[X], \int_0^1 P(t) dt = \lambda_0 P(0) + \lambda_1 P\left(\frac{1}{2}\right) + \lambda_2 P(1)$$

Pour tout $P \in \mathbb{R}_2[X], \exists a, b, c \in \mathbb{R}, P = aX^2 + bX + c$.

D'une part : $\int_0^1 P(t) dt = \int_0^1 (at^2 + bt + c) dt = \left[a \frac{t^3}{3} + b \frac{t^2}{2} + ct \right]_0^1 = \frac{a}{3} + \frac{b}{2} + c$

D'autre part : $\lambda_0 P(0) + \lambda_1 P\left(\frac{1}{2}\right) + \lambda_2 P(1) = \lambda_0 c + \lambda_1 \left(\frac{a}{4} + \frac{b}{2} + c\right) + \lambda_2 (a + b + c)$
 $= a\left(\frac{\lambda_1}{4} + \lambda_2\right) + b\left(\frac{\lambda_1}{2} + \lambda_2\right) + c(\lambda_0 + \lambda_1 + \lambda_2)$

On cherche donc trois réels λ_0, λ_1 et λ_2 tels que :

$$\forall a, b, c \in \mathbb{R}, \frac{a}{3} + \frac{b}{2} + c = a\left(\frac{\lambda_1}{4} + \lambda_2\right) + b\left(\frac{\lambda_1}{2} + \lambda_2\right) + c(\lambda_0 + \lambda_1 + \lambda_2) \quad (*)$$

On a bien envie, par identification, de déclarer que cela revient à imposer :

$$\frac{\lambda_1}{4} + \lambda_2 = \frac{1}{3}, \quad \frac{\lambda_1}{2} + \lambda_2 = \frac{1}{2} \quad \text{et} \quad \lambda_0 + \lambda_1 + \lambda_2 = 1 \quad \text{C'est vrai, mais il va falloir le justifier...}$$

Si (*) est vrai, alors en particulier :

- en prenant $a = 1, b = 0$ et $c = 0$: $\frac{\lambda_1}{4} + \lambda_2 = \frac{1}{3}$
- en prenant $a = 0, b = 1$ et $c = 0$: $\frac{\lambda_1}{2} + \lambda_2 = \frac{1}{2}$
- en prenant $a = 0, b = 0$ et $c = 1$: $\lambda_0 + \lambda_1 + \lambda_2 = 1$

Hein ? Mais qu'est-ce qui nous permet de prendre des valeurs de a, b et c en particulier ? Le fait que () impose une égalité vraie **pour tous** réels a, b , et c . Donc, en particulier,*

pour des valeurs de notre choix.

Réciproquement, si $\frac{\lambda_1}{4} + \lambda_2 = \frac{1}{3}$, $\frac{\lambda_1}{2} + \lambda_2 = \frac{1}{2}$ et $\lambda_0 + \lambda_1 + \lambda_2 = 1$, alors (*) est trivialement vrai.

Résolvons donc le système suivant :

(dont les deux premières lignes correspondent à un système classique de deux équations à deux inconnues...)

$$\begin{aligned} \begin{cases} \frac{\lambda_1}{4} + \lambda_2 &= \frac{1}{3} \\ \frac{\lambda_1}{2} + \lambda_2 &= \frac{1}{2} \\ \lambda_0 + \lambda_1 + \lambda_2 &= 1 \end{cases} &\iff \begin{cases} \frac{\lambda_1}{4} + \lambda_2 &= \frac{1}{3} \\ \lambda_1 &= 1 - 2\lambda_2 \\ \lambda_0 &= 1 - \lambda_1 - \lambda_2 \end{cases} &\iff \begin{cases} \frac{1}{4}(1 - 2\lambda_2) + \lambda_2 &= \frac{1}{3} \\ \lambda_1 &= 1 - 2\lambda_2 \\ \lambda_0 &= 1 - (1 - 2\lambda_2) - \lambda_2 \end{cases} \\ &\iff \begin{cases} \lambda_2 &= \frac{1}{6} \\ \lambda_1 &= 1 - 2 \times \frac{1}{6} \\ \lambda_0 &= \lambda_2 \end{cases} &\iff \begin{cases} \lambda_2 &= \frac{1}{6} \\ \lambda_1 &= \frac{2}{3} \\ \lambda_0 &= \frac{1}{6} \end{cases} \end{aligned}$$

Les coefficients λ_0, λ_1 et λ_2 pour que la formule $I_2(f) = \lambda_0 f(0) + \lambda_1 f(\frac{1}{2}) + \lambda_2 f(1)$ soit exacte sur $\mathbb{R}_2[X]$ sont donc : $\lambda_0 = \frac{1}{6}, \lambda_1 = \frac{2}{3}$ et $\lambda_2 = \frac{1}{6}$

Que peut-on dire sur l'ordre de cette formule de quadrature $I_2(f) = \frac{1}{6}f(0) + \frac{2}{3}f(\frac{1}{2}) + \frac{1}{6}f(1)$?
Pour l'instant, qu'il est supérieur ou égal à 2.

Le polynôme $Q = X^3$ appartient à $\mathbb{R}_3[X]$. D'une part : $\int_0^1 Q(t) dt = \int_0^1 t^3 dt = \left[\frac{t^4}{4} \right]_0^1 = \frac{1}{4}$

D'autre part : $\frac{1}{6}Q(0) + \frac{2}{3}Q(\frac{1}{2}) + \frac{1}{6}Q(1) = \frac{1}{6} \times 0^3 + \frac{2}{3} \times (\frac{1}{2})^3 + \frac{1}{6} \times 1^3 = \frac{2}{3} \times \frac{1}{8} + \frac{1}{6} = \frac{1}{12} + \frac{1}{6} = \frac{3}{12} = \frac{1}{4}$.

Houlà, moi qui espérais tenir un contre-exemple... Cette formule a l'air d'être exacte sur $\mathbb{R}_3[X]$. Mais attention, pour l'établir, il me faut prendre un polynôme quelconque de $\mathbb{R}_3[X]$. Pas spécifiquement X^3 (même si le calcul effectué précédemment peut me servir...)

Pour tout $P \in \mathbb{R}_3[X] : \exists a, b, c, d \in \mathbb{R}, \forall t \in \mathbb{R}, P(t) = at^3 + bt^2 + ct + d$.

D'où : $\int_0^1 P(t) dt = \int_0^1 (at^3 + bt^2 + ct + d) dt = a \int_0^1 t^3 dt + \int_0^1 (bt^2 + ct + d) dt$ par linéarité.

Quel intérêt de fragmenter notre intégrale ainsi ? Pour profiter de ce que nous avons déjà établi, à savoir l'exactitude de la formule sur $\mathbb{R}_2[X]$, et le fait qu'elle marche pour X^3 .

Avec les $\lambda_0, \lambda_1, \lambda_2$ déterminés précédemment, et en introduisant le polynôme $R = bX^2 + cX + d$, nous savons : $\int_0^1 R(t) dt = \lambda_0 R(0) + \lambda_1 R(\frac{1}{2}) + \lambda_2 R(1)$, puisque la formule de quadrature $I_2(f) = \lambda_0 f(0) + \lambda_1 f(\frac{1}{2}) + \lambda_2 f(1)$ est exacte sur $\mathbb{R}_2[X]$.

De plus, avec $Q = X^3$, nous avons aussi établi : $\int_0^1 Q(t) dt = \lambda_0 Q(0) + \lambda_1 Q(\frac{1}{2}) + \lambda_2 Q(1)$

Remarquez : $P = aQ + R$

En reprenant l'expression de $\int_0^1 P(t) dt$, nous obtenons : $\int_0^1 P(t) dt$
 $= a \int_0^1 Q(t) dt + \int_0^1 R(t) dt = a \left(\lambda_0 Q(0) + \lambda_1 Q(\frac{1}{2}) + \lambda_2 Q(1) \right) + \lambda_0 R(0) + \lambda_1 R(\frac{1}{2}) + \lambda_2 R(1)$
 $= \lambda_0 \left(aQ(0) + R(0) \right) + \lambda_1 \left(aQ(\frac{1}{2}) + R(\frac{1}{2}) \right) + \lambda_2 \left(aQ(1) + R(1) \right) = \lambda_0 P(0) + \lambda_1 P(\frac{1}{2}) + \lambda_2 P(1)$

Nous avons établi : $\forall P \in \mathbb{R}_3[X], \int_0^1 P(t) dt = \lambda_0 P(0) + \lambda_1 P(\frac{1}{2}) + \lambda_2 P(1)$

La formule de quadrature $I_2(f) = \lambda_0 f(0) + \lambda_1 f(\frac{1}{2}) + \lambda_2 f(1)$ est donc exacte sur $\mathbb{R}_3[X]$, et son ordre est supérieur ou égal à 3.

Allez, un dernier effort, en espérant qu'elle ne soit pas exacte sur $\mathbb{R}_3[X]$... En vérité, nous aurions pu nous rassurer en vérifiant qu'elle ne marche pas pour X^4 avant notre calcul général précédent (histoire, si par hasard elle avait été vraie pour X^4 , de montrer qu'elle est exacte sur $\mathbb{R}_4[X]$).

Le polynôme $S = X^4$ appartient à $\mathbb{R}_4[X]$. D'une part : $\int_0^1 S(t) dt = \int_0^1 t^4 dt = \left[\frac{t^5}{5} \right]_0^1 = \frac{1}{5}$
 D'autre part : $\frac{1}{6}S(0) + \frac{2}{3}S(\frac{1}{2}) + \frac{1}{6}S(1) = \frac{1}{6} \times 0^4 + \frac{2}{3} \times \left(\frac{1}{2}\right)^4 + \frac{1}{6} \times 1^4 = \frac{2}{3} \times \frac{1}{16} + \frac{1}{6} = \frac{1}{24} + \frac{1}{6}$
 $= \frac{1}{24} + \frac{4}{24} = \frac{5}{24} \neq \frac{1}{5}$

La formule de quadrature $I_2(f) = \lambda_0 f(0) + \lambda_1 f(\frac{1}{2}) + \lambda_2 f(1)$ n'est donc pas exacte sur $\mathbb{R}_4[X]$.

Enfin, l'ordre de cette formule de quadrature est 3.

Exercice 8

*La croissance est bien lente et le soleil se lève :
ce n'est que maintenant que ma démo s'achève.*

Énoncé : (temps conseillé : 30 min) (***) d'après Mines-Ponts 2024 PC Maths 1

On rappelle l'inégalité triangulaire : $\forall a, b \in \mathbb{R}, |a + b| \leq |a| + |b|$.

Soit F l'ensemble de fonctions de $\mathbb{R}$ dans $\mathbb{R}$. Soit A un sous-ensemble de F (c'est-à-dire, tout simplement, un ensemble d'éléments qui appartiennent tous à F : peut-être F tout entier, peut-être pas). On aurait aussi pu dire : soit A une partie de F

On dit que A est un sous-espace vectoriel de F lorsque les deux conditions suivantes sont vérifiées :

- (i) F est non vide (c'est-à-dire qu'il existe au moins un élément de F appartenant à A)
- (ii) pour tous réels a et b , pour toutes fonctions f et g appartenant à A , la fonction $af + bg$ appartient aussi à A .

On note $CL(\mathbb{R})$ l'ensemble des fonctions de $\mathbb{R}$ dans $\mathbb{R}$ à croissance lente, c'est-à-dire :

$$CL(\mathbb{R}) = \left\{ f \in F, \exists C > 0, \exists k \in \mathbb{N} \text{ tel que pour tout } x \in \mathbb{R}, |f(x)| \leq C(1 + |x|^k) \right\}$$

1) Soit $f \in F$. On suppose qu'il existe $d \in \mathbb{N}$ et des réels $a_0, a_1, \dots, a_d$ tels que :

$$\forall x \in \mathbb{R}, |f(x)| \leq \sum_{i=0}^d a_i |x|^i. \quad \text{Montrer que } f \text{ est à croissance lente.}$$

2) Montrer que $CL(\mathbb{R})$ est un sous-espace vectoriel de F .

Remarques sur l'énoncé :

L'inégalité triangulaire rappelée en début d'énoncé se généralise à une somme de plusieurs termes : $\forall N \in \mathbb{N}^*, \forall a_1, a_2, \dots, a_N \in \mathbb{R}, |a_1 + a_2 + \dots + a_N| \leq |a_1| + |a_2| + \dots + |a_N|$. Autrement dit : $\left| \sum_{k=1}^N a_k \right| \leq \sum_{k=1}^N |a_k|$. Cela s'obtient, à partir de l'inégalité triangulaire classique, par récurrence. Je ne dis pas que vous en aurez besoin dans cet exercice, et je ne dis pas non plus que vous n'en aurez pas besoin...

Ensuite, l'énoncé introduit la notion de sous-espace vectoriel (en la définissant afin qu'elle soit utilisable sans autres prérequis que le programme de Terminale), sans même

avoir introduit la définition d'espace vectoriel (plus riche, et pas nécessaire pour la résolution de l'exercice) que vous verrez en première année. Pour faire (très) simple, un espace vectoriel est un ensemble d'éléments muni de deux lois (ou opérations) satisfaisant certaines règles sur cet ensemble. En particulier, notre ensemble F muni de la loi interne $+$ (addition entre deux fonctions) et de la loi externe $\cdot$ (multiplication d'une fonction par un réel) est un espace vectoriel.

Pour f et g deux fonctions de F , pour a et b deux réels, la fonction $af + bg$ est tout simplement la fonction définie sur $\mathbb{R}$ par $(af + bg)(x) = af(x) + bg(x)$

La condition (ii) pour être un sous-espace vectoriel de F est en fait la stabilité par combinaison linéaire. Si f et g appartiennent à A , alors toute combinaison linéaire de f et g (autrement dit, toute fonction de la forme $af + bg$ avec a et b deux réels) appartient encore à A .

Enfin, je ne suis pas particulièrement fan du mélange de genres dans la notation :

$$CL(\mathbb{R}) = \left\{ f \in F, \exists C > 0, \exists k \in \mathbb{N} \text{ tel que pour tout } x \in \mathbb{R}, |f(x)| \leq C(1 + |x|^k) \right\}$$

Je lui aurais préféré :

$$CL(\mathbb{R}) = \left\{ f \in F, \exists C > 0, \exists k \in \mathbb{N}, \forall x \in \mathbb{R}, |f(x)| \leq C(1 + |x|^k) \right\}$$

Mais j'ai tout de même gardé la première, qui est celle du sujet. Je crois deviner que ce dernier a souhaité, de la sorte, éviter un enchaînement trop long de symboles qui aurait pu faire paniquer inutilement les élèves, quitte à faire grincer des dents deux ou trois puristes. Cette notation n'est pas dérangeante en soi, même si on vous incite en général à éviter, dans vos compositions, des mélanges trop abusés entre « français » et symboles mathématiques (ce qu'il m'arrive tout de même de faire, comme - il me semble - tout le monde). Quoi qu'il en soit, entraînez-vous à savoir décoder ces notations, qui, « en toutes lettres », se lisent : « l'ensemble des f appartenant à F telles qu'il existe $C > 0$ et k appartenant à $\mathbb{N}$ tels que pour tout x appartenant à $\mathbb{R}$, $|f(x)| \leq C(1 + |x|^k)$ »

Oui, l'enchaînement des « tel que » est un peu moche...

Correction de l'exercice 8 :

1) Il nous faut donc montrer l'existence d'un réel strictement positif C et d'un entier naturel k tels que : $\forall x \in \mathbb{R}, |f(x)| \leq C(1 + |x|^k)$

Nous savons que pour tout réel x , $|f(x)| \leq \sum_{i=0}^d a_i |x|^i$. Or, pour tout $i \in \llbracket 0 ; d \rrbracket$, $a_i \leq |a_i|$

Rappelons en effet que pour tout réel t , $|t|$ étant le maximum entre t et $-t$, on a : $t \leq |t|$

Puis (comme $|x|^i \geq 0$) : $a_i |x|^i \leq |a_i| |x|^i$ et donc, en sommant : $\sum_{i=0}^d a_i |x|^i \leq \sum_{i=0}^d |a_i| |x|^i$.

D'où : $|f(x)| \leq \sum_{i=0}^d |a_i| |x|^i$

L'ensemble $\{|a_0|, |a_1|, \dots, |a_d|\}$ étant un ensemble fini de réels, il admet un maximum.

Il existe donc un réel (positif) M tel que : $\forall i \in \llbracket 0 ; d \rrbracket, |a_i| \leq M$

Plus précisément, il existe $j \in \llbracket 0 ; d \rrbracket$ tel que ce M est égal à a_j . Mais l'indice j ne me sert pas à grand-chose ici. En réalité, pour établir l'existence d'un tel M , plutôt que de parler de maximum, j'aurais pu me contenter de dire que la partie $\{|a_0|, |a_1|, \dots, |a_d|\}$ est majorée. Mais bon, introduire le maximum m'a semblé plus parlant.

Nous obtenons donc : $|f(x)| \leq \sum_{i=0}^d M |x|^i = M \sum_{i=0}^d |x|^i$

Linéarité de la somme. Revenez ici si ça ne vous dit plus rien...

Bon, mais cette somme continue tout de même de me déranger... J'aimerais bien dire que tous les $|x|^i$ sont inférieurs ou égaux à $|x|^d$ (c-à-d $|x|$ élevé à la plus grande puissance possible), mais je ne tomberai pas dans ce piège. Si $|x|$ est strictement compris entre 0 et 1, c'est le contraire : plus l'entier naturel i est grand, et plus $|x|^i$ est petit...

Si $|x| \geq 1$: $\forall i \in \llbracket 0 ; d \rrbracket, |x|^i \leq |x|^d$ (Suite géométrique $(|x|^n)$ croissante dans ce cas)

Et si $|x| < 1$: $\forall i \in \llbracket 0 ; d \rrbracket, |x|^i \leq 1$

Dans tous les cas : $\forall i \in \llbracket 0 ; d \rrbracket, |x|^i \leq 1 + |x|^d$

Majoration pertinente ici, au vu de ce à quoi nous voulons aboutir... Dans d'autres situations, il eût été plus pertinent de dire : $\forall i \in \llbracket 0 ; d \rrbracket, |x|^i \leq \max(1, |x|^d)$

Il s'ensuit donc : $\forall x \in \mathbb{R}, |f(x)| \leq M \sum_{i=0}^d (1 + |x|^d) = M(d+1)(1 + |x|^d)$

Dans la dernière somme, $1 + |x|^d$ est une constante. La somme est donc égale à cette

constante multipliée par le nombre de termes de la somme, c'est-à-dire $d + 1$.

Nous y sommes presque ! Nous tenons notre entier k (ici d en l'occurrence), et potentiellement notre réel strictement positif C . Pouvons-nous prendre $C = M(d + 1)$? Pas tout à fait, car le M introduit précédemment peut être nul (pour l'éviter, nous aurions pu, au moment de l'introduire, majorer par un M strictement plus grand, quitte à ce que ce ne soit plus le maximum des $|a_i|$). Et bien, au lieu de M , prenons simplement $M + 1$...

Par suite : $\forall x \in \mathbb{R}, |f(x)| \leq M(d + 1)(1 + |x|^d) \leq (M + 1)(d + 1)(1 + |x|^d)$

En posant $C = (M + 1)(d + 1)$ et $k = d$, nous avons bien : $C > 0$, $k \in \mathbb{N}$, et :

$\forall x \in \mathbb{R}, |f(x)| \leq C(1 + |x|^k)$ Enfin, f est à croissance lente.

2) Il suffit de montrer que $CL(\mathbb{R})$ est un sous-ensemble de F vérifiant (i) et (ii).

$CL(\mathbb{R})$ est un sous-ensemble de F par définition.

C'est l'ensemble des fonctions f appartenant à F vérifiant certaines conditions, donc c'est nécessairement un sous-ensemble de F .

Soit f_0 la fonction nulle sur $\mathbb{R}$. f_0 vérifie : $\forall x \in \mathbb{R}, |f(x)| \leq 1 \times (1 + |x|^0)$

Ici, $C = 1$, $k = 0$, mais nous aurions aussi bien pu prendre $C = 13$ et $k = 7$

D'où ; $f_0 \in CL(\mathbb{R})$ et donc $CL(\mathbb{R})$ est non vide.

Soient f et g deux fonctions de $CL(\mathbb{R})$, et soient $a, b \in \mathbb{R}$. Montrons : $af + bg \in CL(\mathbb{R})$

L'appartenance de f et g à $CL(\mathbb{R})$ fournit :

$\exists C_1, C_2 \in \mathbb{R}_+^*, \exists k_1, k_2 \in \mathbb{N}, \forall x \in \mathbb{R}, |f(x)| \leq C_1(1 + |x|^{k_1})$ et $|g(x)| \leq C_2(1 + |x|^{k_2})$

Attention, il n'y a aucune raison que le C soit le même pour f et g . Idem pour k .

$\forall x \in \mathbb{R}, |(af + bg)(x)| = |af(x) + bg(x)| \leq |af(x)| + |bg(x)|$ d'après l'inégalité triangulaire.

Donc : $\forall x \in \mathbb{R}, |(af + bg)(x)| \leq |a||f(x)| + |b||g(x)|$ Mini-étape pour rappeler $|xy| = |x| \times |y|$

Or (comme $|a| \geq 0$ et $|b| \geq 0$) : $|a||f(x)| \leq |a|C_1(1 + |x|^{k_1})$ et $|b||g(x)| \leq |b|C_2(1 + |x|^{k_2})$

Donc : $\forall x \in \mathbb{R}, |(af + bg)(x)| \leq |a|C_1(1 + |x|^{k_1}) + |b|C_2(1 + |x|^{k_2})$

Faut-il refaire tout le travail de 1) pour établir que $af + bg$ est à croissance lente ? Flemme. Plutôt nous servir directement du résultat de 1)... Si nous mettons en évidence le fait que $|(af + bg)(x)|$ est majoré par une expression polynomiale en $|x|$, alors 1) nous permettra de conclure que $af + bg$ est à croissance lente. En réalité, cette majoration par

une expression polynomiale en $|x|$, nous la tenons déjà...

Nous avons donc : $\forall x \in \mathbb{R}, |(af + bg)(x)| \leq |a|C_1 + |b|C_2 + |a||x|^{k_1} + |b||x|^{k_2}$ (*)

Soit P la fonction définie sur $\mathbb{R}$ par $P(t) = |a|C_1 + |b|C_2 + |a|t^{k_1} + |b|t^{k_2}$

P est une fonction polynomiale (car somme de fonctions polynomiales), d'où l'existence d'un entier naturel d et de réels $a_0, a_1, \dots, a_d$ tels que : $\forall t \in \mathbb{R}, P(t) = \sum_{i=0}^d a_i t^i$

Cette petite pirouette m'a permis d'éviter d'explicitier d et les a_i avec une distinction de cas un peu agaçante (selon que k_1 et k_2 soient égaux ou non, selon que l'un d'entre eux - ou les deux - soit égal à 0...). Par exemple, dans le cas où $k_1 = k_2$ et $k_1 > 0$: $|a|C_1 + |b|C_2 + |a||x|^{k_1} + |b||x|^{k_2} = |a|C_1 + |b|C_2 + (|a| + |b|)|x|^{k_1} = \sum_{i=0}^d a_i |x|^i$ avec $d = k_1$, $a_0 = |a|C_1 + |b|C_2$, $a_d = |a| + |b|$, et pour tout entier i strictement compris entre 0 et k_1 (s'il en existe) : $a_i = 0$. Voyez la lourdeur juste pour un cas. D'où l'intérêt d'avoir esquissé ces tracasseries de formalisation par un argument tout à fait sérieux : le fait qu'une somme de fonctions polynomiales reste une fonction polynomiale.

Nous avons établi précédemment (cf (*)) : $\forall x \in \mathbb{R}, |(af + bg)(x)| \leq P(|x|)$

Autrement dit : $\forall x \in \mathbb{R}, |(af + bg)(x)| \leq \sum_{i=0}^d a_i |x|^i$

Le résultat de 1) nous permet alors de conclure : $af + bg \in CL(\mathbb{R})$.

Nous avons donc bien montré : $\forall f, g \in CL(\mathbb{R}), \forall a, b \in \mathbb{R}, af + bg \in CL(\mathbb{R})$

En conclusion, $CL(\mathbb{R})$ est un sous-espace vectoriel de F .

Exercice 9

*Jeune matheux je vois l'hésitation très grande
dans laquelle te place une telle intégrande.*

Énoncé : (temps conseillé : 1 h 30 min) (****) *d'après CCINP 2020 PC*

On rappelle l'inégalité triangulaire pour les intégrales : pour tous réels a et b tels que $a \leq b$, pour toute fonction f continue sur $[a ; b]$, $\left| \int_a^b f(t) dt \right| \leq \int_a^b |f(t)| dt$

1) Montrer l'inégalité suivante : $\forall x \in [-1 ; 1], |e^x - 1| \leq |x|e$

a) d'une première manière par une étude de fonction

b) d'une seconde manière en écrivant $e^x - 1$ comme une intégrale.

2) Soit $x \in \mathbb{R}$ et soit f_x la fonction définie sur $\mathbb{R}$ par : $\forall t \in \mathbb{R}, f_x(t) = \begin{cases} \frac{\sin(t)}{t} e^{-xt} & \text{si } t \neq 0 \\ 1 & \text{si } t = 0 \end{cases}$

Montrer que f_x est continue sur $\mathbb{R}$.

3) Soit F la fonction définie sur $\mathbb{R}$ par $F(x) = \int_0^1 f_x(t) dt$. Montrer que F est continue sur $\mathbb{R}$.

Remarques sur l'énoncé :

L'inégalité triangulaire pour les intégrales vous évoque peut-être celle pour les sommes rappelée dans l'exercice précédent : $\left| \sum_{k=1}^N a_k \right| \leq \sum_{k=1}^N |a_k|$

L'an prochain, l'inégalité des accroissements finis deviendra votre amie pour établir rapidement l'inégalité demandée en 1).

Correction de l'exercice 9 :

1)a) *Ce serait bien d'étudier le signe de la différence, mais ces valeurs absolues me dérangent. Peut-être qu'une distinction de cas...*

Pour tout x appartenant à $[0 ; 1]$, $e^x \geq e^0$ par croissance de la fonction $\exp$ sur $\mathbb{R}$.

Donc : $\forall x \in [0 ; 1]$, $e^x \geq 1$. Autrement dit : $e^x - 1 \geq 0$, et donc : $|e^x - 1| = e^x - 1$.

Par ailleurs, $x \geq 0$ donc $|x|e = xe$.

Sur $[0 ; 1]$, montrer l'inégalité demandée revient donc à montrer : $e^x - 1 \leq xe$

Soit la fonction g définie sur $[0 ; 1]$ par $g(x) = e^x - 1 - xe$

g est dérivable sur $[0 ; 1]$ par somme de fonctions dérivables, et nous avons :

$\forall x \in [0 ; 1]$, $g'(x) = e^x - e$. Par croissance de $\exp$ sur $\mathbb{R}$: $\forall x \in [0 ; 1]$, $e^x \leq e^1$, d'où : $e^x - e \leq 0$

Donc : $\forall x \in [0 ; 1]$, $g'(x) \leq 0$. g est donc décroissante sur $[0 ; 1]$.

De plus, $g(0) = e^0 - 1 - 0 = 0$. D'où : $\forall x \in [0 ; 1]$, $g(x) \leq g(0) = 0$.

Autrement dit : $\forall x \in [0 ; 1]$, $e^x - 1 - xe \leq 0$. Ou encore : $\forall x \in [0 ; 1]$, $|e^x - 1| \leq |x|e$

Pour tout x appartenant à $[-1 ; 0]$, $e^x \leq e^0$ par croissance de $\exp$ sur $\mathbb{R}$.

Donc : $\forall x \in [-1 ; 0]$, $e^x - 1 \leq 0$, et donc : $|e^x - 1| = -(e^x - 1) = 1 - e^x$.

Par ailleurs, $x \leq 0$ donc $|x|e = -xe$.

Sur $[-1 ; 0]$, montrer l'inégalité demandée revient donc à montrer : $1 - e^x \leq -xe$

Soit la fonction h définie sur $[-1 ; 0]$ par $h(x) = 1 - e^x + xe$

h est dérivable sur $[-1 ; 0]$ par somme de fonctions dérivables, et nous avons :

$\forall x \in [-1 ; 0]$, $h'(x) = -e^x + e$. Par croissance de $\exp$ sur $\mathbb{R}$: $\forall x \in [-1 ; 0]$, $e^x \leq e^0 < e^1$, d'où : $e^x - e \leq 0$. Donc : $-e^x + e \geq 0$. D'où : $\forall x \in [-1 ; 0]$, $h'(x) \geq 0$, et h est croissante sur $[-1 ; 0]$.

De plus, $h(0) = 1 - e^0 + 0 = 0$. D'où : $\forall x \in [-1 ; 0]$, $h(x) \leq h(0) = 0$.

Autrement dit : $\forall x \in [-1 ; 0]$, $1 - e^x \leq -xe$. Ou encore : $\forall x \in [-1 ; 0]$, $|e^x - 1| \leq |x|e$

Nous avons bien établi : $\forall x \in [-1 ; 1]$, $|e^x - 1| \leq |x|e$

1)b) « En écrivant $e^x - 1$ comme une intégrale » ? Comment une intégrale pourrait apparaître juste à partir de cette différence ? Ah, différence...

Pour tout x appartenant à $[-1 ; 1]$: $e^x - 1 = e^x - e^0 = [e^t]_0^x = \int_0^x e^t dt$

En effet, la fonction exponentielle est une primitive d'elle-même sur $\mathbb{R}$

Donc : $\forall x \in [-1 ; 1], |e^x - 1| = \left| \int_0^x e^t dt \right|$

Un petit coup d'inégalité triangulaire pour les intégrales... Attention, pas si vite ! L'énoncé nous le dit bien, la borne du bas doit être inférieure ou égale à la borne du haut, ce qui n'est pas toujours le cas ici... Donc si l'on voulait appliquer cette inégalité triangulaire maintenant, il faudrait distinguer les cas $x \geq 0$ et $x < 0$.

Mais de toute façon, cette inégalité n'est pas forcément utile ici, l'intégrande e^t étant de signe constant (strictement positive). Par contre, ça ne change rien au fait qu'il faille faire attention au signe de x ...

Si $x \geq 0$, la positivité de l'intégrale fournit : $\int_0^x e^t dt \geq 0$, donc $\left| \int_0^x e^t dt \right| = \int_0^x e^t dt$

D'où : $|e^x - 1| = \int_0^x e^t dt$. Et : $\forall t \in [0 ; x], e^t \leq e^x$ par croissance de $\exp$ sur $\mathbb{R}$

Donc, par croissance de l'intégrale (les bornes étant dans le bon sens) :

$\int_0^x e^t dt \leq \int_0^x e^x dt = xe^x \leq xe^1 = |x|e$ $\int_0^x e^x dt$ est l'intégrale d'une constante...

Nous avons établi : $\forall x \in [0 ; 1], |e^x - 1| \leq |x|e$

Si $x \leq 0$, $\left| \int_0^x e^t dt \right| = \left| - \int_x^0 e^t dt \right| = \left| \int_x^0 e^t dt \right|$.

Les bornes étant dans le bon sens, la positivité de l'intégrale fournit encore :

$\left| \int_x^0 e^t dt \right| = \int_x^0 e^t dt$. D'où : $|e^x - 1| = \int_x^0 e^t dt$

Et, par croissance de l'intégrale : $\int_x^0 e^t dt \leq \int_x^0 e^0 dt = \int_x^0 1 dt = -x = |x| \leq |x|e$

Nous avons établi : $\forall x \in [-1 ; 0], |e^x - 1| \leq |x|e$

Une fois encore, nous avons établi : $\forall x \in [-1 ; 1], |e^x - 1| \leq |x|e$

2) Dans l'expression donnée par l'énoncé, ce que la fonction f_x prend en variable, c'est bien t , et pas x , qui est constante dans ce contexte.

Pour tout $t \in \mathbb{R}^*$, $f_x(t) = \frac{\sin(t)}{t} e^{-xt}$. D'une part, la fonction $t \mapsto e^{-xt}$ est continue sur $\mathbb{R}$ (et donc a fortiori sur $\mathbb{R}^*$) en tant que composée de fonctions continues sur $\mathbb{R}$. D'autre part, la fonction $t \mapsto \frac{\sin(t)}{t}$ est continue sur $\mathbb{R}^*$ par quotient, dont le dénominateur ne s'annule pas sur $\mathbb{R}^*$, de telles fonctions. Enfin, par produit, f_x est continue sur $\mathbb{R}^*$

L'an prochain, vous pourrez vous permettre d'aller un peu plus vite sur cette continuité par opérations « classiques »

Nous avons donc établi la continuité de f_x sur $\mathbb{R}^$. Attention, cela ne veut pas dire que f n'est pas continue en 0. Cela veut juste dire que f est continue en tous les réels qui ne sont pas 0. Rien n'est dit sur 0 pour l'instant, et d'ailleurs, vu que l'énoncé nous demande de montrer la continuité de f_x sur $\mathbb{R}$, il nous faut établir cette continuité manquante en 0. Pour cela, nous devons montrer : $\lim_{t \rightarrow 0} f_x(t) = f_x(0)$*

Nous savons : $\lim_{t \rightarrow 0} -xt = 0$. Puis, par continuité de la fonction exponentielle sur $\mathbb{R}$ (et en particulier en 0) : $\lim_{t \rightarrow 0} e^{-xt} = e^0 = 1$

Pour $\frac{\sin(t)}{t}$, c'est une autre paire de manches... Ça donne une forme indéterminée « $\frac{0}{0}$ » Une idée de comment lever cette indétermination ? Pas de factorisation, pas de manipulation calculatoire immédiate en vue. Comment ? Qu'est-ce que j'apprends ? On m'accuse de faire du forcing sur les limites de taux d'accroissement ? Je plaide coupable.

Pour tout $t \neq 0$, $\frac{\sin(t)}{t} = \frac{\sin(t) - \sin(0)}{t - 0}$. Or, la fonction sin est dérivable sur $\mathbb{R}$ (et en particulier en 0) et $\sin' = \cos$. Donc : $\lim_{t \rightarrow 0} \frac{\sin(t)}{t} = \sin'(0) = \cos(0) = 1$

Enfin, par produit de limites : $\lim_{t \rightarrow 0} f_x(t) = 1 = f_x(0)$. f_x est donc continue en 0.

En conclusion, f_x est bien continue sur $\mathbb{R}$.

3) Les choses se compliquent pas mal. Les deuxième année trouveraient assez osé que je confronte de jeunes gens au sortir de la Terminale à une intégrale à paramètre qui ne dit pas son nom. Mais le programme de Terminale et les questions précédentes nous suffiront, promis.

Pour tout réel x , la fonction f_x est continue sur $\mathbb{R}$ (cf question 2) et donc en particulier sur $[0 ; 1]$. $F(x)$ est donc bien défini. Autrement dit, F est bien définie sur $\mathbb{R}$. L'énoncé le dit déjà et ne nous demande pas de le justifier. Simple vérification de courtoisie (et pour vous rassurer éventuellement).

Soit $a \in \mathbb{R}$.

Montrons que $\lim_{x \rightarrow a} F(x) = F(a)$. Autrement dit, montrons : $\lim_{x \rightarrow a} \int_0^1 f_x(t) dt = \int_0^1 f_a(t) dt$

$$\text{Pour tout réel } x : F(x) - F(a) = \int_0^1 f_x(t) dt - \int_0^1 f_a(t) dt = \int_0^1 (f_x(t) - f_a(t)) dt$$

*Si nous arrivions à montrer que cette dernière quantité tend vers 0 lorsque x tend vers a ...
 Passons à la valeur absolue, ça pourrait nous faire - enfin - profiter de l'inégalité triangulaire pour les intégrales rappelée en début d'énoncé...*

$\forall x \in \mathbb{R}, |F(x) - F(a)| = \left| \int_0^1 (f_x(t) - f_a(t)) dt \right| \leq \int_0^1 |f_x(t) - f_a(t)| dt$ d'après l'inégalité triangulaire pour les intégrales (les bornes étant dans le bon sens)

Pour tout réel $t \in]0 ; 1]$:

$$|f_x(t) - f_a(t)| = \left| \frac{\sin(t)}{t} e^{-xt} - \frac{\sin(t)}{t} e^{-at} \right| = \left| \frac{\sin(t)}{t} e^{-at} (1 - e^{(a-x)t}) \right| = |f_a(t)| |1 - e^{(a-x)t}|$$

$$\text{De même, pour } t = 0 : |f_x(0) - f_a(0)| = |1 - 1| = 0 \text{ et } |f_a(0)| |1 - e^{(a-x)0}| = 1 \times |1 - 1| = 0$$

$$\text{Nous avons donc : } \forall t \in [0 ; 1], |f_x(t) - f_a(t)| = |f_a(t)| |1 - e^{(a-x)t}|$$

Ce $|1 - e^{(a-x)t}|$ ne vous rappelle rien ? Quelque chose me dit qu'on pourrait le majorer intelligemment...

$$\text{Pour tout } t \in [0 ; 1], |1 - e^{(a-x)t}| = |e^{(a-x)t} - 1| \quad \text{Pour tout réel } y, |y| = |-y|$$

Nous aimerions nous servir de l'inégalité établie en 1), mais pour cela, il faut que $(a-x)t$ appartienne à $[-1 ; 1]$. La bonne nouvelle, c'est que nous nous intéressons à ce qui se passe lorsque x tend vers a . Nous pouvons donc prendre x proche de a si cela nous arrange...

Pour tout $x \in [a-1 ; a+1]$, pour tout $t \in [0 ; 1]$: $|a-x| \leq 1$ et $|t| \leq 1$ donc $|(a-x)t| \leq 1$.

Autrement dit, $(a-x)t \in [-1 ; 1]$. D'après 1), nous savons donc : $|e^{(a-x)t} - 1| \leq |(a-x)t| e$

Ou encore : $|1 - e^{(a-x)t}| \leq |(a-x)t| e$

Et comme $|f_a(t)| \geq 0$, nous obtenons : $|f_x(t) - f_a(t)| \leq |f_a(t)| \times |(a-x)t| e = |f_a(t)| \times |a-x| \times t e$

Par croissance de l'intégrale :

$$\forall x \in [a-1 ; a+1], 0 \leq |F(x) - F(a)| \leq \int_0^1 |f_a(t)| |a-x| \times t e dt \quad (*)$$

$$\text{Or, par linéarité : } \int_0^1 |f_a(t)| |a-x| \times t e dt = |a-x| \times \int_0^1 |f_a(t)| t e dt$$

Nous avons sorti $|a-x|$, constante multiplicative vis-à-vis de t .

Et alors ? Certes, nous avons pas mal écrit, mais n'en oublions pas l'objectif : montrer que $F(x) - F(a)$ tend vers 0 quand x tend vers a (ce qui revient à montrer que $|F(x) - F(a)|$ tend vers 0 quand x tend vers a)

Or : $\lim_{x \rightarrow a} |a - x| = 0$, et l'intégrale $\int_0^1 |f_a(t)| t e \, dt$ est un réel ne dépendant pas de x .

Donc : $\lim_{x \rightarrow a} |a - x| \times \int_0^1 |f_a(t)| t e \, dt = 0$.

À partir de l'encadrement (*), le théorème des gendarmes fournit : $\lim_{x \rightarrow a} |F(x) - F(a)| = 0$

Autrement dit : $\lim_{x \rightarrow a} F(x) - F(a) = 0$, et enfin : $\lim_{x \rightarrow a} F(x) = F(a)$.

F est donc continue en a .

Cela étant vrai pour tout réel a , nous avons en fait montré que F est continue sur $\mathbb{R}$.

Exercice 10

*L'on te dénie le droit de prendre un air trop fier
tant que t'emplit d'effroi cette partie entière*

Énoncé : (temps conseillé : 25 min) (***) d'après Mines-Ponts 2016 PC Maths 1

On notera $\lfloor x \rfloor$ la partie entière (inférieure) d'un réel x .

On considère une suite $(a_n)_{n \in \mathbb{N}}$ décroissante de réels positifs et, pour tout entier naturel n , on pose $S_n = \sum_{k=0}^n a_k$. On suppose que : $\lim_{n \rightarrow +\infty} \frac{S_n}{\sqrt{n}} = 2$

1) Soient α et β deux réels vérifiant : $0 < \alpha < 1 < \beta$. Pour tout entier naturel n tel que $n - \lfloor \alpha n \rfloor$ et $n - \lfloor \beta n \rfloor$ soient non nuls, justifier l'encadrement : $\frac{S_{\lfloor \beta n \rfloor} - S_n}{\beta n - n} \leq a_n \leq \frac{S_n - S_{\lfloor \alpha n \rfloor}}{n - \alpha n}$

2) Soit γ un réel strictement positif. Déterminer les limites des suites de termes généraux $\frac{n}{\lfloor \gamma n \rfloor}$ et $\frac{S_{\lfloor \gamma n \rfloor}}{\sqrt{n}}$

Remarques sur l'énoncé :

Par défaut, lorsqu'on dit « partie entière de x » sans préciser inférieure ou supérieure, on parle de la partie entière inférieure $\lfloor x \rfloor$.

Rappelons que pour tout réel x , $\lfloor x \rfloor$ est le plus grand entier relatif inférieur ou égal à x . C'est l'unique entier relatif vérifiant : $\lfloor x \rfloor \leq x < \lfloor x \rfloor + 1$. De manière équivalente, c'est aussi l'unique entier relatif vérifiant : $x - 1 < \lfloor x \rfloor \leq x$. En particulier, si $x \in \mathbb{Z}$, $\lfloor x \rfloor = x$

On pourra utiliser sans démonstration le fait (qui découle simplement de sa définition) que la fonction partie entière est croissante sur $\mathbb{R}$.

Sinon, pas de rappels sur les sommes cette fois... Revenez ici si vous en sentez le besoin.

Enfin, l'an prochain, plutôt que d'écrire : $\lim_{n \rightarrow +\infty} \frac{S_n}{\sqrt{n}} = 2$, vous écririez aussi : $S_n \underset{n \rightarrow +\infty}{\sim} 2\sqrt{n}$
Mais ne vous embêtons pas avec les équivalents pour le moment...

Correction de l'exercice 10 :

1) Soit un entier naturel n tel que $n - \lfloor \alpha n \rfloor$ et $n - \lfloor \beta n \rfloor$ soient non nuls.

En particulier, $n \in \mathbb{N}^*$ (si n était nul, on aurait eu : $n = \lfloor \alpha n \rfloor = \lfloor \beta n \rfloor = 0$) Puisque $\alpha < 1$, $\alpha n < n$ et donc, par croissance de la fonction partie entière sur $\mathbb{R}$, $\lfloor \alpha n \rfloor \leq \lfloor n \rfloor = n$. Et comme $n - \lfloor \alpha n \rfloor \neq 0$, nous pouvons en déduire : $\lfloor \alpha n \rfloor < n$

Attention, la croissance simple de $x \mapsto \lfloor x \rfloor$ ne suffit pas, à partir de $\alpha n < n$, pour établir $\lfloor \alpha n \rfloor < \lfloor n \rfloor$. La fonction partie entière a beau être croissante, elle n'est pas strictement croissante (loin de là, avec les plateaux qu'elle fait entre deux entiers consécutifs). C'est le fait, par ailleurs, de savoir que $n - \lfloor \alpha n \rfloor \neq 0$ qui nous a permis d'obtenir l'inégalité stricte. De même, à partir de l'inégalité $1 < \beta$, nous obtenons : $n < \lfloor \beta n \rfloor$

$\lfloor \alpha n \rfloor$ et $\lfloor \beta n \rfloor$ étant des entiers naturels, $S_{\lfloor \alpha n \rfloor}$ et $S_{\lfloor \beta n \rfloor}$ sont bien définis.

$$\text{Puis : } S_{\lfloor \beta n \rfloor} - S_n = \sum_{k=0}^{\lfloor \beta n \rfloor} a_k - \sum_{k=0}^n a_k = \sum_{k=n+1}^{\lfloor \beta n \rfloor} a_k \quad \text{Cette somme est non vide car } n+1 \leq \lfloor \beta n \rfloor$$

De plus, la suite (a_n) est décroissante. Donc : $\forall k \in \llbracket n+1; \lfloor \beta n \rfloor \rrbracket, a_k \leq a_{n+1} \leq a_n$

Pourquoi avoir pensé à cette dernière majoration par a_n ? Parce que j'ai les yeux rivés sur le résultat à obtenir, qui est un encadrement de a_n et pas a_{n+1} .

$$\text{D'où : } S_{\lfloor \beta n \rfloor} - S_n = \sum_{k=n+1}^{\lfloor \beta n \rfloor} a_k \leq \sum_{k=n+1}^{\lfloor \beta n \rfloor} a_n = (\lfloor \beta n \rfloor - n) a_n$$

Le terme général de la dernière somme est une constante vis-à-vis de l'indice de sommation k . Ecrire $\sum_{k=n+1}^{\lfloor \beta n \rfloor} a_n$ revient donc à écrire $a_n + a_n + \dots + a_n$, avec a_n répété autant de fois qu'il y a de termes dans la somme, à savoir $\lfloor \beta n \rfloor - (n+1) + 1 = \lfloor \beta n \rfloor - n$

Nous avons donc établi : $S_{\lfloor \beta n \rfloor} - S_n \leq (\lfloor \beta n \rfloor - n) a_n$

Puis, en divisant par $\lfloor \beta n \rfloor - n > 0$:

$$\frac{S_{\lfloor \beta n \rfloor} - S_n}{\lfloor \beta n \rfloor - n} \leq a_n$$

$$\text{D'autre part : } S_n - S_{\lfloor \alpha n \rfloor} = \sum_{k=0}^n a_k - \sum_{k=0}^{\lfloor \alpha n \rfloor} a_k = \sum_{k=\lfloor \alpha n \rfloor+1}^n a_k \geq \sum_{k=\lfloor \alpha n \rfloor+1}^n a_n$$

par décroissance de (a_n)

Donc : $S_n - S_{[\alpha n]} \geq (n - [\alpha n])a_n$. Puis, comme $n - [\alpha n] > 0$: $\frac{S_n - S_{[\alpha n]}}{n - [\alpha n]} \geq a_n$

Enfin, nous avons bien montré que pour tout entier naturel n tel que $n - [\alpha n]$ et $n - [\beta n]$ soient non nuls : $\frac{S_{[\beta n]} - S_n}{\beta n - n} \leq a_n \leq \frac{S_n - S_{[\alpha n]}}{n - \alpha n}$

2) Notons $u_n = \frac{n}{[\gamma n]}$ et $v_n = \frac{S_{[\gamma n]}}{\sqrt{n}}$ Mais pour quels n ? J'arrive, j'arrive...

u_n est bien défini pour tout $n \in \mathbb{N}^*$, et v_n est bien défini si et seulement si $[\gamma n] \neq 0$

Or : $[\gamma n] = 0 \iff 0 \leq \gamma n < 1 \iff 0 \leq n < \frac{1}{\gamma}$.

A contrario, v_n est bien défini pour tout entier naturel n vérifiant $n \geq \frac{1}{\gamma}$

Il existe bien un entier $n_0 \in \mathbb{N}^*$ tel que : $\forall n \geq n_0, n \geq \frac{1}{\gamma}$

Il n'est pas nécessaire d'expliciter un tel entier n_0 ici, mais si vous voulez, $n_0 = \lfloor \frac{1}{\gamma} \rfloor + 1$ convient.

Les suites (u_n) et (v_n) sont bien définies à partir du rang n_0 .

Ce qui permet déjà de se poser la question de leur limite.

Pour tout $n \geq n_0$, $\gamma n - 1 < [\gamma n] \leq \gamma n$. Donc : $\gamma - \frac{1}{n} < \frac{[\gamma n]}{n} \leq \gamma$

Or : $\lim_{n \rightarrow +\infty} \gamma - \frac{1}{n} = \gamma$. Donc, d'après le théorème des gendarmes : $\lim_{n \rightarrow +\infty} \frac{[\gamma n]}{n} = \gamma$ avec $\gamma \neq 0$

Enfin, par quotient de limites : $\lim_{n \rightarrow +\infty} u_n = \frac{1}{\gamma}$. Autrement dit : $\lim_{n \rightarrow +\infty} \frac{n}{[\gamma n]} = \frac{1}{\gamma}$

Il était plus commode pour moi de travailler avec $[\gamma n]$ au numérateur. Je ne m'en suis donc pas privé, quitte à passer à l'inverse après. Dans le cas présent, plutôt que « par quotient de limites », j'aurais aussi pu dire « par continuité de la fonction inverse en γ »

Intéressons-nous maintenant à la convergence de (v_n) . Nous savons : $\lim_{n \rightarrow +\infty} \frac{S_n}{\sqrt{n}} = 2$

Ce qui va nous permettre de justifier assez simplement que : $\lim_{n \rightarrow +\infty} \frac{S_{[\gamma n]}}{\sqrt{[\gamma n]}} = 2...$

Pour tout $n \geq n_0$, $\gamma n - 1 < [\gamma n]$. Or : $\lim_{n \rightarrow +\infty} \gamma n - 1 = +\infty$. Donc, par théorème de com-

paraison : $\lim_{n \rightarrow +\infty} [\gamma n] = +\infty$. Puis, par composée de limites : $\lim_{n \rightarrow +\infty} \frac{S_{[\gamma n]}}{\sqrt{[\gamma n]}} = 2$

Mais la limite que nous voulons, c'est celle de $\frac{S_{[\gamma n]}}{\sqrt{n}}$...

$$\text{Pour tout } n \geq n_0, v_n = \frac{S_{[\gamma n]}}{\sqrt{n}} = \frac{S_{[\gamma n]}}{\sqrt{[\gamma n]}} \times \frac{\sqrt{[\gamma n]}}{\sqrt{n}} = \frac{S_{[\gamma n]}}{\sqrt{[\gamma n]}} \times \sqrt{\frac{[\gamma n]}{n}}$$

Or : $\lim_{n \rightarrow +\infty} \frac{[\gamma n]}{n} = \gamma$ (limite établie précédemment)

Donc, par continuité de la fonction racine carrée en γ : $\lim_{n \rightarrow +\infty} \sqrt{\frac{[\gamma n]}{n}} = \sqrt{\gamma}$

Et, par hypothèse : $\lim_{n \rightarrow +\infty} \frac{S_{[\gamma n]}}{\sqrt{[\gamma n]}} = 2$.

Enfin, par produit de limites : $\lim_{n \rightarrow +\infty} \frac{S_{[\gamma n]}}{\sqrt{n}} = 2\sqrt{\gamma}$

Une fois n'est pas coutume, il n'y a pas de lien entre les deux questions de cet exercice (elles servent plus tard dans l'énoncé originel). À vous de vous adapter à chaque situation : si, dans un exercice donné, les questions ont très souvent un lien entre elles, il ne faut pas non plus « forcer les choses » et se sentir obligé, coûte que coûte, d'utiliser le résultat d'une question précédente là où c'est manifestement inutile...

Pour être honnête, le reproche statistique à faire va plutôt dans l'autre sens : l'élève moyen est plus souvent coupable d'oublier les questions précédentes qui pourraient lui servir plutôt que de leur utilisation « forcée ».

Exercice 11

*Je le présente ici car il m'a fait tiquer :
c'est un sujet de l'X que j'ai décortiqué.*

Énoncé : (temps conseillé : 1 h) (***) *d'après X-ENS-ESPCI 2022 PC*

Soient α et β deux réels tels que $\alpha > \beta$. Soit $I = \left] \frac{\alpha + \beta}{2}, +\infty \right[$, et soit f la fonction définie sur l'intervalle I par : $f(x) = (x - \alpha)(x - \beta)$.
Soit h la fonction définie sur $\mathbb{R} \setminus \{\beta\}$ par $h(x) = \frac{x - \alpha}{x - \beta}$

1) Montrer que : $\forall x \in \mathbb{R} \setminus \{\beta\} : |h(x)| < 1 \iff x \in I$

2) Soit $x_0 \in I$. On pose : $\forall n \in \mathbb{N}, x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}$

a) Montrer que la suite (x_n) est bien définie, et que pour tout entier naturel n , $x_n \in I$

b) On pose : $\forall n \in \mathbb{N}, u_n = h(x_n)$. Justifier que la suite (u_n) est bien définie, puis expliciter la relation de récurrence satisfaite par cette suite (u_n) .

c) Montrer que la suite (u_n) tend vers 0 et en déduire que (x_n) tend vers α .

Remarques sur l'énoncé :

La question 2)b) vous demande d'explicitier la relation de récurrence satisfaite par (u_n) . Il s'agit de réussir à exprimer, pour tout entier naturel n , u_{n+1} en fonction de u_n . Ce serait une récurrence d'ordre 1 dans cette situation, mais vous pourrez, l'an prochain, être confrontés à des ordres supérieurs. Par exemple, une relation de récurrence d'ordre 2 serait l'expression de u_{n+2} en fonction de u_n et u_{n+1}

Correction de l'exercice 11 :

1) Nous pourrions étudier la fonction h , et, à partir de ses variations, établir l'équivalence demandée. Mais faisons plus simple. Résolvons directement l'inéquation. D'aucuns pourraient penser à distinguer des cas pour se débarrasser de la valeur absolue, mais faisons quelque chose de plus rigolo ici...

$$\begin{aligned}\forall x \in \mathbb{R} \setminus \{\beta\} : |h(x)| < 1 &\iff \frac{|x - \alpha|}{|x - \beta|} < 1 \iff |x - \alpha| < |x - \beta| \text{ (par stricte positivité de } |x - \beta|) \\ &\iff |x - \alpha|^2 < |x - \beta|^2 \text{ par stricte croissance de la fonction carré sur } \mathbb{R}_+ \\ &\iff (x - \alpha)^2 < (x - \beta)^2 \quad \text{En effet, pour tout réel } y, |y|^2 = y^2 = y^2 \\ &\iff x^2 - 2\alpha x + \alpha^2 < x^2 - 2\beta x + \beta^2 \\ &\iff x^2 - 2\alpha x + \alpha^2 < x^2 - 2\beta x + \beta^2 \quad \text{Au revoir les } x^2 \dots \\ &\iff (2\beta - 2\alpha)x < \beta^2 - \alpha^2 \iff x > \frac{\beta^2 - \alpha^2}{2(\beta - \alpha)} \text{ car } \beta - \alpha < 0 \\ &\iff x > \frac{(\beta + \alpha)(\beta - \alpha)}{2(\beta - \alpha)} \iff x > \frac{\beta + \alpha}{2} \iff x \in I\end{aligned}$$

Nous avons bien établi : $\forall x \in \mathbb{R} \setminus \{\beta\} : |h(x)| < 1 \iff x \in I$

2)a) « Montrer que la suite (x_n) est bien définie » ? Comment ça ? Qu'est-ce qui pourrait mal se passer ? Si pour un entier naturel n donné, $f'(x_n)$ avait le mauvais goût d'être nul, ça poserait un gros souci pour x_{n+1} ...

f est une fonction polynomiale, donc dérivable sur I .

Pour tout réel $x \in I$, $f(x) = x^2 - (\alpha + \beta)x + \alpha\beta$ donc $f'(x) = 2x - (\alpha + \beta) = 2\left(x - \frac{\alpha + \beta}{2}\right)$

Remarquons : $\forall x \in I, f'(x) > 0$ (donc a fortiori : $\forall x \in I, f'(x) \neq 0$)

Bonne nouvelle ! Si pour un n donné, $x_n \in I$, alors x_{n+1} est bien défini. Reste à savoir s'il appartient aussi à I . Mais supposer que c'est vrai pour un certain rang n , et vouloir l'établir pour le rang suivant, ça s'inscrirait dans le cadre de quel raisonnement ? Vous me voyez probablement venir. Mais avant de me lancer dans une récurrence, il serait commode d'établir que l'intervalle I est stable par la fonction $g : x \mapsto x - \frac{f(x)}{f'(x)}$. Autrement dit, que : $\forall x \in I, g(x) \in I$. Etudions tout simplement les variations de g pour l'établir. Ensuite, notre récurrence, et en particulier, l'étape d'hérédité, deviendra un jeu d'enfant...

g est définie et dérivable sur I par quotient (dont le dénominateur ne s'annule pas sur I) et somme de fonctions dérivables. Pour tout réel x appartenant à I :

$$g'(x) = 1 - \frac{f'(x) \times f'(x) - f(x)f''(x)}{(f'(x))^2} = 1 - \frac{(f'(x))^2}{(f'(x))^2} + \frac{f(x)f''(x)}{(f'(x))^2} \quad \text{Attention au signe...}$$

Donc : $\forall x \in I, g'(x) = \frac{f(x)f''(x)}{(f'(x))^2}$ avec $f(x) = (x - \alpha)(x - \beta)$, $f'(x) = 2x - (\alpha + \beta)$ et $f''(x) = 2$

C'est le signe de $g'(x)$ qui nous intéresse.

Pour tout x appartenant à I , $f''(x) = 2 > 0$ et $(f'(x))^2 > 0$. $g'(x)$ est donc du signe de $f(x)$

f est une fonction polynomiale du second degré, admettant deux racines distinctes α et β (avec $\beta < \alpha$), et de coefficient dominant $1 > 0$ « Son a » dans l'écriture $ax^2 + bx + c$
Sur $\mathbb{R}$ tout entier, le signe de $(x - \alpha)(x - \beta)$ est le suivant :

x	$-\infty$	β	α	$+\infty$	
$(x-\alpha)$ $\times(x-\beta)$	+	0	-	0	+

J'ai dit : « sur $\mathbb{R}$ tout entier, le signe de $(x - \alpha)(x - \beta)$ » et pas celui de $f(x)$. En toute rigueur, l'énoncé n'a choisi de définir f que sur $\left] \frac{\alpha + \beta}{2}, +\infty \right[$.

Pour obtenir le signe de $f(x) = (x - \alpha)(x - \beta)$ (et donc de $g'(x)$ sur l'intervalle $\left] \frac{\alpha + \beta}{2}, +\infty \right[$, il suffit de « zoomer » comme suit.

Puisque $\beta < \alpha$, remarquons : $\beta < \frac{\alpha + \beta}{2} < \alpha$

La moyenne arithmétique de deux réels distincts est évidemment strictement supérieure au plus petit de ces deux réels, et strictement inférieure au plus grand. Si vous en voulez une preuve formelle : $\beta + \beta < \alpha + \beta < \alpha + \alpha$, puis on divise tout ce beau monde par 2.

Nous pouvons donc établir le tableau de signe de f (qui est aussi celui de g') sur I et, en conséquence, le tableau de variations de g sur I :

x	$\frac{\alpha + \beta}{2}$	α	$+\infty$	
$g'(x)$		-	0	+
g				

$g(\alpha) = \alpha - \frac{f(\alpha)}{f'(\alpha)} = \alpha$ car $f(\alpha) = 0$ Et si vous n'avez pas perdu notre but de vue, vous comprenez pourquoi je n'ai aucune raison de m'embêter à calculer les limites de g en $\frac{\alpha + \beta}{2}$ et en $+\infty$...

g admet un minimum en α , et ce minimum est $g(\alpha) = \alpha$. Donc : $\forall x \in I, g(x) \geq \alpha > \frac{\alpha + \beta}{2}$

D'où : $\forall x \in I, g(x) \in \left] \frac{\alpha + \beta}{2}, +\infty \right[= I$. Ouf! Résultat qui va nous être très utile.

Soit, pour tout $n \in \mathbb{N}$, la propriété P_n : « x_n est bien défini et $x_n \in I$ »

Montrons par récurrence que pour tout entier naturel n , P_n est vraie.

Initialisation : x_0 est bien défini et $x_0 \in I$ par hypothèse, donc P_0 est vraie.

Hérédité : Supposons que pour un certain $n \in \mathbb{N}$, P_n soit vraie, et montrons que P_{n+1} aussi est vraie.

Supposons donc que x_n soit bien défini, et que $x_n \in I$.

g est définie sur I , donc $x_{n+1} = g(x_n)$ est bien défini, et d'après ce qui précède, $x_{n+1} \in I$ P_{n+1} est donc vraie.

Conclusion : Le principe de raisonnement par récurrence nous permet de conclure que pour tout entier naturel n , P_n est vraie.

Autrement dit, (x_n) est bien définie et, pour tout entier naturel n , $x_n \in I$.

2)b) Pour tout entier naturel n , $x_n \in I$, donc $x_n > \frac{\alpha + \beta}{2} > \beta$. D'où : $x_n \in \mathbb{R} \setminus \{\beta\}$, et donc $u_n = h(x_n)$ est bien défini. La suite (u_n) est donc bien définie.

$$\text{De plus : } \forall n \in \mathbb{N}, u_{n+1} = h(x_{n+1}) = \frac{x_{n+1} - \alpha}{x_{n+1} - \beta} = \frac{g(x_n) - \alpha}{g(x_n) - \beta}$$

Que faire à partir de ça ? Peut-être expliciter $g(x)$?

Pour tout x appartenant à I :

$$g(x) = x - \frac{f(x)}{f'(x)} = x - \frac{(x - \alpha)(x - \beta)}{2x - (\alpha + \beta)} = \frac{2x^2 - (\alpha + \beta)x - (x^2 - (\alpha + \beta)x + \alpha\beta)}{2x - (\alpha + \beta)} = \frac{x^2 - \alpha\beta}{2x - (\alpha + \beta)}$$

$$\text{Donc : } g(x_n) - \alpha = \frac{x_n^2 - \alpha\beta}{2x_n - (\alpha + \beta)} - \alpha = \frac{x_n^2 - \alpha\beta - 2x_n\alpha + \alpha(\alpha + \beta)}{2x_n - (\alpha + \beta)} = \frac{x_n^2 - 2x_n\alpha + \alpha^2}{2x_n - (\alpha + \beta)} \quad \text{Oh..}$$

$$\text{Puis } g(x_n) - \alpha = \frac{(x_n - \alpha)^2}{2x_n - (\alpha + \beta)}. \text{ On montre de même : } g(x_n) - \beta = \frac{(x_n - \beta)^2}{2x_n - (\alpha + \beta)}$$

$$\text{Par quotient : } \forall n \in \mathbb{N}, u_{n+1} = \frac{(x_n - \alpha)^2}{(x_n - \beta)^2} = \left(\frac{x_n - \alpha}{x_n - \beta} \right)^2. \text{ Autrement dit : } \forall n \in \mathbb{N}, u_{n+1} = u_n^2$$

2)c) Intéressant : c'est la convergence de (u_n) qui est censée nous aider à établir celle de (v_n) , et pas l'inverse... u_n , c'est $h(x_n)$, et nous avons obtenu une info intéressante sur la valeur absolue de $h(x)$ en 1)...

Pour tout entier naturel n : $|u_{n+1}| = |u_n^2| = |u_n|^2 = |u_n| \times |u_n| = |h(x_n)| \times |u_n|$

Quelle idée d'avoir remplacé l'un des $|u_n|$ par $|h(x_n)|$ et pas l'autre ! Mais pour quoi faire ? Ben je sais pas moi, descends... Ah non pardon, là je pars sur carrément autre chose. Plus sérieusement, cela m'a permis d'exprimer $|u_{n+1}|$ en fonction de $|u_n|$ et d'une quantité dont je sais qu'elle est inférieure à 1...

Pour tout entier naturel n , $x_n \in I$, donc d'après l'équivalence établie en 1), $|h(x_n)| < 1$. Puis (comme $|u_n| \geq 0$), $|h(x_n)| \times |u_n| \leq |u_n|$. Autrement dit : $\forall n \in \mathbb{N}$, $|u_{n+1}| \leq |u_n|$. La suite $(|u_n|)_{n \in \mathbb{N}}$ est donc décroissante.

Pourquoi ne pas être plutôt passé par le quotient $\frac{|u_{n+1}|}{|u_n|}$? Pour m'éviter une réflexion sur l'annulation éventuelle de $|u_n|$ au dénominateur.

De plus, cette suite $(|u_n|)_{n \in \mathbb{N}}$ est minorée par 0 (tous ses termes sont positifs). Le théorème de convergence monotone nous permet donc d'affirmer que $(|u_n|)_{n \in \mathbb{N}}$ converge. Notons $l \in \mathbb{R}$ sa limite.

Nous savons : $\forall n \in \mathbb{N}$, $|u_{n+1}| = |u_n|^2$. Par passage à la limite des deux côtés de cette égalité, nous obtenons : $l = l^2$

J'ai laissé votre théorème du point fixe (et notamment l'hypothèse de continuité) au repos pour cette fois. Le fait que $|u_n|$ tend vers l me donne, par un simple produit de limites, que $|u_n|^2$ tend vers l^2 . Si j'avais voulu utiliser ce théorème ici, j'aurais parlé de la continuité de la fonction $r : x \mapsto x^2$ sur $\mathbb{R}$. Les autres hypothèses sont bien au rendez-vous, $(|u_n|)$ converge et : $\forall n \in \mathbb{N}$, $|u_{n+1}| = r(|u_n|)$

Nous savons donc : $l^2 = l$, ce qui revient à dire $l^2 - l = 0$, ou encore $l(l - 1) = 0$. D'où : $l = 0$ ou $l = 1$.

Nous y sommes presque ! Il nous faut juste - et sainte Geneviève ne me contredirait pas - écarter ce 1 ...

$|u_0| = |h(x_0)| < 1$ et $(|u_n|)$ est décroissante, donc : $\forall n \in \mathbb{N}, |u_n| \leq |u_0|$

Puis (par passage à la limite) : $l \leq |u_0| < 1$. Donc $l \neq 1$, et enfin, $l = 0$.

$(|u_n|)$ converge donc vers 0, ce qui revient à dire que (u_n) converge vers 0

A nous deux, x_n . Attends que je t'exprime en fonction de u_n , et ton compte sera réglé...

Pour tout entier naturel n , $u_n = h(x_n) = \frac{x_n - \alpha}{x_n - \beta}$. Donc : $(x_n - \beta)u_n = x_n - \alpha$

Puis : $x_n u_n - \beta u_n = x_n - \alpha$, ou encore : $x_n(u_n - 1) = \beta u_n - \alpha$

J'ai bien envie de diviser par $u_n - 1$, mais il me faut d'abord expliquer pourquoi $u_n - 1 \neq 0$

La décroissance de $(|u_n|)$ nous a permis d'établir : $\forall n \in \mathbb{N}, |u_n| \leq |u_0| < 1$. Donc $|u_n| < 1$
 u_n ne peut donc pas être égal à 1.

Il s'ensuit : $\forall n \in \mathbb{N}, x_n = \frac{\beta u_n - \alpha}{u_n - 1}$. Et nous savons : $\lim_{n \rightarrow +\infty} u_n = 0$

Enfin, par opérations sur les limites (produit, somme, quotient) : $\lim_{n \rightarrow +\infty} x_n = \alpha$

Ici, nous étions gâtés niveau hypothèses. Dans d'autres situations (par exemple en l'absence d'informations sur la décroissance de $(|u_n|)$), plutôt que de justifier que pour tout entier naturel n , $u_n \neq 1$, nous aurions pu utiliser le fait que (u_n) converge vers 0 : en revenant à la définition de convergence (pour vous rafraîchir la mémoire, c'est [ici](#)), et en choisissant $\epsilon = \frac{1}{2}$, nous savons qu'il existe un entier naturel n_0 tel que : $\forall n \geq n_0, |u_n - 0| < \frac{1}{2}$.

Autrement dit : $\forall n \geq n_0, -\frac{1}{2} < u_n < \frac{1}{2}$. Et, en particulier : $\forall n \geq n_0, u_n \neq 1$.

Travailler non plus pour tout entier naturel n , mais à partir d'un certain rang n_0 n'aurait pas été dérangeant, puisqu'il s'agit d'obtenir une limite lorsque n tend vers $+\infty$.

Exercice 12

*Le complexe se boit en sortie de frigo
avec une cuillère à soupe de trigo*

Énoncé : (temps conseillé : 50 min) (***) d'après Mines-Ponts 2022 PC Maths 1

Les questions 1 et 2 sont indépendantes.

On rappelle que pour tous réels a et b : $\cos(a+b) = \cos(a)\cos(b) - \sin(a)\sin(b)$

Pour tout nombre complexe z , on notera respectivement $\operatorname{Re}(z)$ et $\operatorname{Im}(z)$ la partie réelle et la partie imaginaire de z .

1) Soit $x \in [0 ; 1[$ et θ un réel.

Montrer que $\frac{1}{1-x} - \operatorname{Re}\left(\frac{1}{1-xe^{i\theta}}\right) \geq \frac{x(1-\cos\theta)}{(1-x)((1-x)^2 + 2x(1-\cos\theta))}$

2) On cherche à montrer qu'il existe un réel $K > 0$ tel que : $\forall \theta \in [-\pi ; \pi], 1 - \cos\theta \geq K\theta^2$

a) Montrer que : $\forall t \in \left[0 ; \frac{\pi}{2}\right], \sin t \geq \frac{2t}{\pi}$, puis que : $\forall t \in \left[-\frac{\pi}{2} ; \frac{\pi}{2}\right], \sin^2(t) \geq \frac{4t^2}{\pi^2}$

b) Conclure.

Remarques sur l'énoncé :

Comme l'énoncé originel, lorsqu'il n'y a pas d'ambiguïté, je note nonchalamment $\cos\theta$ plutôt que $\cos(\theta)$ pour alléger la notation (vous avez probablement pris l'habitude l'an dernier avec $\ln x$). Il pourra tout de même m'arriver de mettre, par moments, les parenthèses dans la correction.

Juste au cas où, $\sin^2(\theta)$, c'est tout simplement $(\sin(\theta))^2$, ou encore $(\sin\theta)^2$

Correction de l'exercice 12 :

1) Tout d'abord, puisque x est un réel positif, $|xe^{i\theta}| = x$. Donc $|xe^{i\theta}| < 1$ et, a fortiori, $xe^{i\theta} \neq 1$. Donc $1 - xe^{i\theta} \neq 0$ Histoire de nous rassurer sur le dénominateur $1 - xe^{i\theta}$

$$\text{De plus : } \frac{1}{1 - xe^{i\theta}} = \frac{\overline{1 - xe^{i\theta}}}{|1 - xe^{i\theta}|^2} = \frac{\overline{1} - \overline{xe^{i\theta}}}{|1 - xe^{i\theta}|^2} = \frac{1 - xe^{-i\theta}}{|1 - xe^{i\theta}|^2}$$

D'accord pour multiplier notre fraction en haut et en bas par le conjugué de $1 - xe^{i\theta}$: c'est ce qu'il y a mieux à faire en général pour mettre une fraction complexe sous forme algébrique. Pas d'accord pour dire n'importe quoi sur ce conjugué de $1 - xe^{i\theta}$: par exemple, par hasard, inventer que ce serait $1 + xe^{i\theta}$, comme si nous avions une forme algébrique $a - ib$ (avec a et b réels, du coup) dont le conjugué est effectivement $a + ib$...

Profitions-en pour rappeler des propriétés de la conjugaison utilisées ci-haut :

$$\forall z_1, z_2 \in \mathbb{C}, \overline{z_1 + z_2} = \overline{z_1} + \overline{z_2} \text{ Et, de même, } \overline{z_1 - z_2} = \overline{z_1} - \overline{z_2}. \text{ En outre, } \overline{z_1 \times z_2} = \overline{z_1} \times \overline{z_2}$$

Nous nous en sommes servi de manière indirecte pour écrire : $\overline{xe^{i\theta}} = xe^{-i\theta}$, sachant que $\overline{\overline{x}} = x$ puisque x est réel.

Enfin, une égalité utile et sous-côtée, alliée de bien des transformations efficaces :

$$\forall z \in \mathbb{C}, z \times \overline{z} = |z|^2$$

$$\begin{aligned} \text{Or : } |1 - xe^{i\theta}|^2 &= (1 - xe^{i\theta})(\overline{1 - xe^{i\theta}}) = (1 - xe^{i\theta})(1 - xe^{-i\theta}) = 1 - xe^{-i\theta} - xe^{i\theta} + x^2 e^{i\theta} e^{-i\theta} \\ &= 1 - x(e^{i\theta} + e^{-i\theta}) + x^2 = 1 - 2x \cos \theta + x^2 \text{ d'après les formules d'Euler.} \end{aligned}$$

$$\text{Ces formules stipulent en effet : } \forall \theta \in \mathbb{R}, \cos(\theta) = \frac{e^{i\theta} + e^{-i\theta}}{2} \text{ et } \sin(\theta) = \frac{e^{i\theta} - e^{-i\theta}}{2i}$$

Elles s'obtiennent simplement à partir des égalités : $e^{i\theta} = \cos \theta + i \sin \theta$ et $e^{-i\theta} = \cos \theta - i \sin \theta$

Voici une autre façon (moins élégante à mon goût, mais tout à fait valable) d'obtenir

$$|1 - xe^{i\theta}|^2 = 1 - 2x \cos \theta + x^2 :$$

$$\begin{aligned} |1 - xe^{i\theta}|^2 &= |1 - x(\cos(\theta) + i \sin(\theta))|^2 = |1 - x \cos(\theta) - ix \sin(\theta)|^2 \\ &= (1 - x \cos \theta)^2 + (x \sin \theta)^2 \quad \text{en effet, si } z = a + ib \text{ avec } a \text{ et } b \text{ réels, } |z|^2 = a^2 + b^2 \\ &= 1 - 2x \cos \theta + x^2 \cos^2(\theta) + x^2 \sin^2(\theta) = 1 - 2x \cos \theta + x^2 (\cos^2(\theta) + \sin^2(\theta)) = 1 - 2x \cos \theta + x^2 \end{aligned}$$

Reprenons le fil...

$$\text{Donc } \frac{1}{1 - xe^{i\theta}} = \frac{1 - xe^{-i\theta}}{1 - 2x \cos \theta + x^2} = \frac{1 - x(\cos(\theta) + i \sin(-\theta))}{1 - 2x \cos \theta + x^2} = \frac{1 - x \cos(\theta) + ix \sin(\theta)}{1 - 2x \cos \theta + x^2}$$

Profitions-en pour rappeler que la fonction sinus est impaire : $\sin(-\theta) = -\sin(\theta)$

Puis : $\frac{1}{1-xe^{i\theta}} = \frac{1-x\cos(\theta)}{1-2x\cos\theta+x^2} + i\frac{x\sin(\theta)}{1-2x\cos\theta+x^2}$. D'où : $\operatorname{Re}\left(\frac{1}{1-xe^{i\theta}}\right) = \frac{1-x\cos(\theta)}{1-2x\cos\theta+x^2}$

Donc : $\frac{1}{1-x} - \operatorname{Re}\left(\frac{1}{1-xe^{i\theta}}\right) = \frac{1}{1-x} - \frac{1-x\cos\theta}{1-2x\cos\theta+x^2} = \frac{1-2x\cos\theta+x^2-(1-x)(1-x\cos\theta)}{(1-x)(1-2x\cos\theta+x^2)}$

Un coup d'œil à ce $\frac{x(1-\cos\theta)}{(1-x)((1-x)^2+2x(1-\cos\theta))}$ par lequel nous cherchons à minorer notre quantité... Est-ce que nous ne tenons pas déjà le bon dénominateur ?

$$(1-x)^2 + 2x(1-\cos\theta) = 1-2x+x^2+2x-2x\cos\theta = 1-2x\cos\theta+x^2$$

$$\begin{aligned} \text{D'où : } \frac{1}{1-x} - \operatorname{Re}\left(\frac{1}{1-xe^{i\theta}}\right) &= \frac{1-2x\cos\theta+x^2-1+x\cos\theta+x-x^2\cos\theta}{(1-x)((1-x)^2+2x(1-\cos\theta))} \\ &= \frac{-x\cos\theta+x^2+x-x^2\cos\theta}{(1-x)((1-x)^2+2x(1-\cos\theta))} = \frac{x(x-x\cos\theta+1-\cos\theta)}{(1-x)((1-x)^2+2x(1-\cos\theta))} \end{aligned}$$

J'ai factorisé par x au numérateur, mais il me semble que je peux faire mieux...

$$\text{Donc : } \frac{1}{1-x} - \operatorname{Re}\left(\frac{1}{1-xe^{i\theta}}\right) = \frac{x(x(1-\cos\theta)+1-\cos\theta)}{(1-x)((1-x)^2+2x(1-\cos\theta))} = \frac{x(1-\cos\theta)(x+1)}{(1-x)((1-x)^2+2x(1-\cos\theta))}$$

Génial !

x étant un réel positif, nous savons : $x+1 \geq 1$. De plus, $\frac{x(1-\cos\theta)}{(1-x)((1-x)^2+2x(1-\cos\theta))} \geq 0$

En effet : $x \geq 0$, $1-\cos\theta \geq 0$, $1-x > 0$ (car $x < 1$)

Et $(1-x)^2 + 2x(1-\cos\theta) \geq 0$ par somme et produit de termes positifs.

$$\text{Donc : } (x+1) \times \frac{x(1-\cos\theta)}{(1-x)(1-2x\cos\theta+x^2)} \geq 1 \times \frac{x(1-\cos\theta)}{(1-x)(1-2x\cos\theta+x^2)}$$

$$\text{Autrement dit : } \boxed{\frac{1}{1-x} - \operatorname{Re}\left(\frac{1}{1-xe^{i\theta}}\right) \geq \frac{x(1-\cos\theta)}{(1-x)((1-x)^2+2x(1-\cos\theta))}}$$

2)a) Comme dans l'exercice précédent, je vais esquiver une étude de fonction. J'aurais pu étudier, sur $\left[0 ; \frac{\pi}{2}\right]$, les variations puis le signe de la fonction $f : t \mapsto \sin(t) - \frac{2}{\pi}t$. Cette fois-ci, à une telle étude de fonction, je vais préférer un argument de convexité (en l'occurrence ici de concavité), pour vous convaincre (et renforcer par là-même ma propre conviction, moi qui n'en ai pas toujours été friand) de l'intérêt de cette notion.

La fonction sinus est deux fois dérivable sur $\mathbb{R}$, et : $\forall x \in \mathbb{R}, \sin''(t) = \cos'(t) = -\sin(t)$

En particulier : $\forall t \in \left[0 ; \frac{\pi}{2}\right]$, $\sin(t) \geq 0$ donc $\sin''(t) \leq 0$

La fonction sinus est donc concave sur $\left[0 ; \frac{\pi}{2}\right]$. Sur cet intervalle, la courbe de la fonction sinus est donc au-dessus de ses cordes (ou sécantes si vous les appelez comme ça).

En particulier, si l'on se place dans un repère, sur cet intervalle $\left[0 ; \frac{\pi}{2}\right]$, la courbe de la fonction sinus est au-dessus du segment reliant les point $A(0; \sin(0))$ et $B(\frac{\pi}{2}; \sin(\frac{\pi}{2}))$ (autrement dit, $A(0;0)$ et $B(\frac{\pi}{2};1)$) de la courbe de la fonction sinus.

Ce segment $[AB]$ est inclus dans la droite (AB) d'équation $y = \frac{2}{\pi}x$

Hein ? D'où sort cette équation ? Cette droite passe par l'origine du repère (son ordonnée à l'origine est donc nulle), et son coefficient directeur est $\frac{y_B - y_A}{x_B - x_A} = \frac{1 - 0}{\frac{\pi}{2} - 0} = \frac{2}{\pi}$. D'où l'équation de (AB) : $y = \frac{2}{\pi}x + 0$

Enfin : $\forall t \in \left[0 ; \frac{\pi}{2}\right]$, $\sin t \geq \frac{2}{\pi}t$

Sans indication de l'énoncé, il est assez peu raisonnable d'attendre d'un élève au sortir de la Terminale qu'il pense seul à la concavité de $\sin$ sur cet intervalle, à la position de sa courbe par rapport à ses cordes, et qu'il choisisse la corde qui, comme par hasard, a la bonne équation. Dans le supérieur, il acquerra le réflexe de penser un peu plus souvent à la convexité/concavité pour établir certaines inégalités. Parfois, l'habitude parlera plus que l'éclair de génie : l'inégalité que nous venons d'établir est en effet un classique de première année.

Pour quiconque la découvre, ce $\frac{2}{\pi}$ pouvait mettre la puce à l'oreille. Comment faire apparaître une division par $\frac{\pi}{2}$...

Un mot, tout de même, sur l'autre méthode évoquée, à savoir l'étude de fonction, que ces considérations de convexité m'ont permis d'éviter. On définit la fonction h sur $\left[0 ; \frac{\pi}{2}\right]$

par $h(t) = \sin(t) - \frac{2}{\pi}t$.

h est bien dérivable sur $\left[0 ; \frac{\pi}{2}\right]$ par somme de telles fonctions.

$\forall t \in \left[0 ; \frac{\pi}{2}\right]$, $h'(t) = \cos(t) - \frac{2}{\pi}$

Le signe de h' sur $\left[0 ; \frac{\pi}{2}\right]$ n'est pas immédiat. Remarquons que h' est strictement décrois-

sante sur $\left[0 ; \frac{\pi}{2}\right]$ (car $\cos$ l'est). $h'(0) = 1 - \frac{2}{\pi} > 0$ et $h'(1) = -\frac{2}{\pi} < 0$. h' étant continue sur $\left[0 ; \frac{\pi}{2}\right]$, toutes les hypothèses sont réunies pour appliquer le corollaire du théorème des valeurs intermédiaires et établir l'existence d'un unique réel $\alpha \in \left[0 ; \frac{\pi}{2}\right]$ tel que $h'(\alpha) = 0$. La stricte décroissance de h' sur $\left[0 ; \frac{\pi}{2}\right]$ nous permet d'en déduire le signe de h' , puis les variations de h sur $\left[0 ; \frac{\pi}{2}\right]$

t	0	α	$\frac{\pi}{2}$
$h'(t)$	+	0	-
h	0		0

Dois-je m'inquiéter de ne pas connaître α , et donc pas non plus $h(\alpha)$? Pas nécessairement : $h(0) = \sin(0) - \frac{2}{\pi} \times 0 = 0$ et $h\left(\frac{\pi}{2}\right) = \sin\left(\frac{\pi}{2}\right) - \frac{2}{\pi} \times \frac{\pi}{2} = 1 - 1 = 0$ (que je rajoute à mon tableau, du coup). Les variations de h nous permettent d'établir que h admet 0 pour minimum sur $\left[0 ; \frac{\pi}{2}\right]$ (minimum atteint en 0 et en $\frac{\pi}{2}$). Donc : $\forall t \in \left[0 ; \frac{\pi}{2}\right], h(t) \geq 0$. D'où l'inégalité souhaitée.

Bon, finalement, c'était plus qu'« un mot » ... Je me suis senti obligé de détailler un minimum parce que la situation (avec ce α inconnu notamment) pouvait sembler plus inquiétante, par exemple, qu'à la question 3 de l'exercice 2.

Maintenant, comment obtenir la seconde inégalité demandée - qui semble correspondre à l'élévation au carré de la première - sur un intervalle plus grand, à savoir $\left[-\frac{\pi}{2} ; \frac{\pi}{2}\right]$?

Avant d'élever au carré (prudemment), il faut déjà se demander ce qui se passe sur $\left[-\frac{\pi}{2} ; 0\right]$

Pour tout $t \in \left[-\frac{\pi}{2} ; 0\right]$, $-t \in \left[0 ; \frac{\pi}{2}\right]$. D'après ce qui précède : $\sin(-t) \geq \frac{2}{\pi} \times (-t)$

Autrement dit, par imparité de $\sin$: $-\sin(t) \geq -\frac{2}{\pi} \times t$. Puis : $\sin(t) \leq \frac{2}{\pi} \times t$

Par décroissance de la fonction carré sur $\mathbb{R}_-$ (intervalle dans lequel se situent $\sin(t)$ et $\frac{2}{\pi} \times t$, le plus grand des deux étant $\frac{2}{\pi} \times t$, évidemment négatif lorsque $t \in \left[-\frac{\pi}{2} ; 0\right]$) :

$$\forall t \in \left[-\frac{\pi}{2} ; 0\right], \sin^2(t) \geq \frac{4t^2}{\pi^2}.$$

D'autre part, nous savions : $\forall t \in \left[0 ; \frac{\pi}{2}\right], \sin t \geq \frac{2}{\pi} t$. Par croissance de la fonction carré

sur $\mathbb{R}_+ : \forall t \in \left[0 ; \frac{\pi}{2}\right], \sin^2(t) \geq \frac{4t^2}{\pi^2}$.

Nous avons donc bien montré : $\forall t \in \left[-\frac{\pi}{2} ; \frac{\pi}{2}\right], \sin^2(t) \geq \frac{4t^2}{\pi^2}$

2)b) Il nous faut établir, à partir de là, qu'il existe un réel $K > 0$ tel que :

$$\forall \theta \in [-\pi ; \pi], 1 - \cos \theta \geq K\theta^2$$

L'inégalité obtenue précédemment est valable pour $t \in \left[-\frac{\pi}{2} ; \frac{\pi}{2}\right]$, et elle concerne non pas un cos mais un $\sin^2$...

Pour tout réel $\theta \in [-\pi ; \pi]$, $\frac{\theta}{2} \in \left[-\frac{\pi}{2} ; \frac{\pi}{2}\right]$.

$$\begin{aligned} \text{Et : } \cos(\theta) &= \cos\left(2 \times \frac{\theta}{2}\right) = \cos\left(\frac{\theta}{2} + \frac{\theta}{2}\right) = \cos\left(\frac{\theta}{2}\right)\cos\left(\frac{\theta}{2}\right) - \sin\left(\frac{\theta}{2}\right)\sin\left(\frac{\theta}{2}\right) = \cos^2\left(\frac{\theta}{2}\right) - \sin^2\left(\frac{\theta}{2}\right) \\ &= 1 - \sin^2\left(\frac{\theta}{2}\right) - \sin^2\left(\frac{\theta}{2}\right) = 1 - 2\sin^2\left(\frac{\theta}{2}\right) \end{aligned}$$

$$\text{Donc : } 1 - \cos(\theta) = 1 - \left[1 - 2\sin^2\left(\frac{\theta}{2}\right)\right] = 2\sin^2\left(\frac{\theta}{2}\right)$$

Certains d'entre vous (rares à ce stade, j'imagine) connaissaient peut-être déjà la formule : $\cos(2x) = 1 - 2\sin^2(x)$, que nous avons retrouvé ici (avec $x = \frac{\theta}{2}$), à partir de la formule $\cos(a+b) = \cos(a)\cos(b) - \sin(a)\sin(b)$ rappelée par l'énoncé. Formule rappelée par sécurité : il me semblait qu'elle était vue en général en trigonométrie de Première ou Terminale, mais certains Terminale m'ont récemment assuré du contraire...

Comme $\frac{\theta}{2} \in \left[-\frac{\pi}{2} ; \frac{\pi}{2}\right]$, nous savons, d'après 2)a) : $\sin^2\left(\frac{\theta}{2}\right) \geq \frac{4}{\pi^2} \times \left(\frac{\theta}{2}\right)^2 = \frac{\theta^2}{\pi^2}$

Donc : $2\sin^2\left(\frac{\theta}{2}\right) \geq \frac{2\theta^2}{\pi^2}$. Autrement dit : $1 - \cos(\theta) \geq \frac{2\theta^2}{\pi^2}$

Enfin, en posant $K = \frac{2}{\pi^2} > 0$, nous avons bien établi : $\forall \theta \in [-\pi ; \pi], 1 - \cos \theta \geq K\theta^2$

Une pensée émue pour la notation Im, introduite en début d'énoncé, mais qui ne nous aura servi à rien...

Exercice 13

*Les inégalités, les tangentes, les cordes
ne sauraient effrayer ma plume monocorde.*

Énoncé : (temps conseillé : 1 h 10 min) (****) d'après Centrale 2017 PC Maths 2

1) Soit I un intervalle de $\mathbb{R}$, et soit f une fonction convexe sur I . Soient a et b appartenant à I , avec $a < b$, et soit $\lambda \in [0; 1]$.

a) Justifier que $\lambda a + (1 - \lambda)b \in [a; b]$.

b) Montrer que $f(\lambda a + (1 - \lambda)b) \leq \lambda f(a) + (1 - \lambda)f(b)$

On pourra utiliser la position de la courbe de f par rapport à ses sécantes (ou cordes).

c) Justifier que l'inégalité précédente reste valable pour tous a et b appartenant à I .

2) Soient $a, b \in \mathbb{R}$. Montrer que : $\forall \lambda \in [0; 1], e^{\lambda a + (1 - \lambda)b} \leq \lambda e^a + (1 - \lambda)e^b$

3) On se place dans un univers Ω . On admet que toutes les variables aléatoires mises en jeu ici possèdent une espérance. Soit c un réel strictement positif, et soit X une variable aléatoire réelle telle que $X(\Omega)$ est fini avec $X(\Omega) \subset [-c; c]$, et telle que $E(X) = 0$. On considère la variable aléatoire réelle Y définie par $Y = \frac{1}{2} - \frac{X}{2c}$.

a) Montrer que $e^X \leq Y e^{-c} + (1 - Y)e^c$

b) Montrer que $E(e^X) \leq \frac{e^c + e^{-c}}{2}$

Remarques sur l'énoncé :

A partir de l'inégalité de convexité qu'on nous demande d'établir en 1)b) et 1)c), il est possible d'obtenir cette inégalité de convexité généralisée (ou inégalité de Jensen) :

pour tout entier $n \geq 2$, on a : pour tous $x_1, x_2, \dots, x_n \in I$, pour tous $\lambda_1, \lambda_2, \dots, \lambda_n \geq 0$ tels que $\sum_{k=1}^n \lambda_k = 1$, $f\left(\sum_{k=1}^n \lambda_k x_k\right) \leq \sum_{k=1}^n \lambda_k f(x_k)$ C'est l'objet de cet exercice (difficile).

Sans rentrer dans trop de détails théoriques que vous aurez tout le loisir de découvrir l'an prochain, une variable aléatoire réelle X est une fonction définie sur un certain univers Ω et à valeurs dans $\mathbb{R}$. $X(\Omega)$ est l'ensemble des valeurs prises par X . Pour tout $\omega \in \Omega$, $X(\omega)$ est la valeur que prend X lorsque l'issue ω se réalise.

$Z = e^X$ est la variable aléatoire définie par : $\forall \omega \in \Omega, Z(\omega) = e^{X(\omega)}$

Parlons d'inégalités entre variables aléatoires. Pour deux variables aléatoires X et Y définies sur un univers Ω , affirmer $X \leq Y$ revient à dire : $\forall \omega \in \Omega, X(\omega) \leq Y(\omega)$. Autrement dit, quelle que soit l'issue, la valeur prise par X est inférieure ou égale à la valeur prise par Y . Il ne faut pas confondre le fait de dire $X \leq Y$ avec le fait de considérer l'événement $[X \leq Y]$, qui peut se réaliser ou pas. Événement dont on peut éventuellement calculer la probabilité, que l'on note $P(X \leq Y)$ par un abus de notation largement admis.

De même, affirmer $X = Y$ revient à dire : $\forall \omega \in \Omega, X(\omega) = Y(\omega)$. A ne pas confondre avec l'événement $[X = Y]$.

Dernière déclinaison (celle qui, d'expérience, induit le plus souvent en erreur) : affirmer $X = 3$ revient à dire que X est la variable aléatoire constante égale à 3. A ne pas confondre avec l'événement $[X = 3]$, qui peut se réaliser ou pas.

A priori, en Terminale, on ne vous embêtait pas trop sur l'existence de l'espérance des variables aléatoires que vous manipuliez. Vous verrez dans le supérieur que toute variable aléatoire n'admet pas forcément une espérance. Mais dans le cas particulier (qui était souvent le vôtre, et qui est celui de cet énoncé) où l'on manipule des variables aléatoires à valeurs dans un ensemble fini, l'existence de cette espérance ne pose pas de problème. En particulier, il n'est pas attendu de votre part ici de justifier l'existence de $E(e^X)$. L'an prochain, dans ce genre de cas simple, vous vous fendrez d'une justification du style « l'ensemble $X(\Omega)$ est fini, donc X admet une espérance ».

On rappelle enfin la croissance de l'espérance : si deux variables aléatoires X et Y admettent toutes deux une espérance, et si $X \leq Y$, alors $E(X) \leq E(Y)$

Correction de l'exercice 13 :

1)a) Une première question en apparence relativement gentille...

$a \leq b$ et $\lambda \in [0; 1]$, donc $\lambda \geq 0$ et $1 - \lambda \geq 0$.

Donc $\lambda a + (1 - \lambda)a \leq \lambda a + (1 - \lambda)b \leq \lambda b + (1 - \lambda)b$. Autrement dit : $a \leq \lambda a + (1 - \lambda)b \leq b$

Nous avons bien montré : $\lambda a + (1 - \lambda)b \in [a; b]$

En toute rigueur, la justification ci-haut suffit à établir le résultat demandé. On multiplie des inégalités par λ et $1 - \lambda$: il était donc important de préciser leur positivité. Si c'est allé trop vite à votre goût, voici une rédaction plus détaillée, en aboutissant séparément aux deux inégalités qui constituent l'encadrement demandé. Partant de $a \leq b$:

D'une part, $\lambda \geq 0$ donc $\lambda a \leq \lambda b$. Puis : $\lambda a + (1 - \lambda)b \leq \lambda b + (1 - \lambda)b$. Donc : $\lambda a + (1 - \lambda)b \leq b$

D'autre part, $\lambda \leq 1$, donc $1 - \lambda \geq 0$. D'où : $(1 - \lambda)b \geq (1 - \lambda)a$. Puis : $\lambda a + (1 - \lambda)b \geq \lambda a + (1 - \lambda)a$.

Donc : $\lambda a + (1 - \lambda)b \geq a$. Les deux inégalités soulignées donnent l'encadrement demandé.

1)b) f est convexe sur I , et a et b sont deux réels de cet intervalle. La corde (ou sécante) (d) joignant les points d'abscisses a et b de la courbe $\mathcal{C}_f$ de f est donc au-dessus de $\mathcal{C}$ sur le segment $[a; b]$ (ici, on a bien $a < b$)

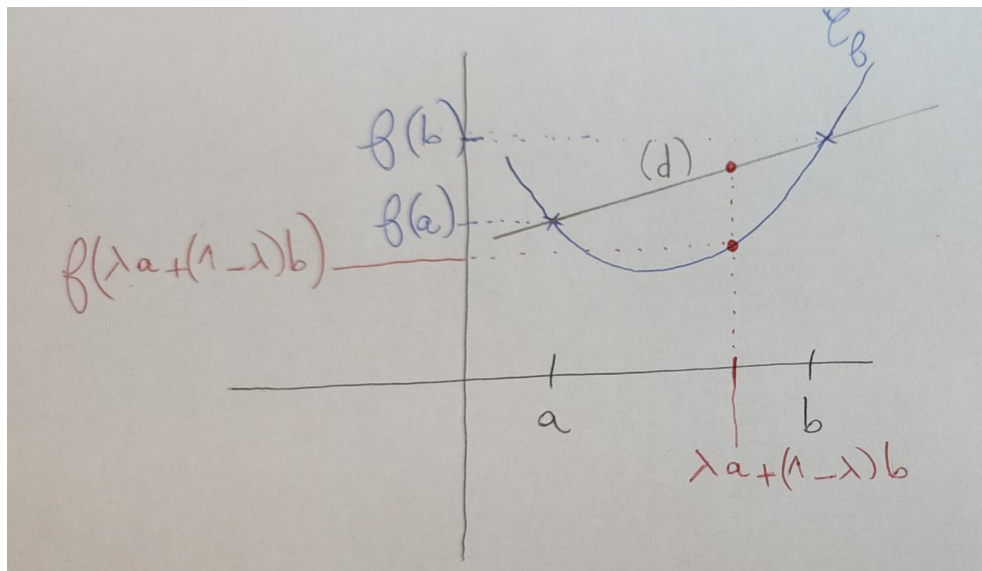


Schéma à main levée (encore...) représentant la situation.

La sécante (d) est au-dessus de $\mathcal{C}_f$ sur $[a; b]$

De plus, d'après 1)a) : $\lambda a + (1 - \lambda)b \in [a; b]$.

Et $f(\lambda a + (1 - \lambda)b)$ est l'ordonnée du point de la courbe de $\mathcal{C}_f$ d'abscisse $\lambda a + (1 - \lambda)b$

Point que nous avons placé sur la courbe. Il était intéressant de voir à quoi correspond géométriquement ce fameux $f(\lambda a + (1 - \lambda)b)$, membre de gauche de l'inégalité demandée.

Ce point est situé en-dessous du point de même abscisse situé sur la droite (d).

On aimerait bien avoir l'ordonnée de ce dernier point d'abscisse $\lambda a + (1 - \lambda)b$ et situé sur la droite (d) : on pourrait alors affirmer que $f(\lambda a + (1 - \lambda)b)$ est inférieur ou égal à cette ordonnée. Et avec un peu de chance, cette ordonnée serait...

Notons M le point d'abscisse $\lambda a + (1 - \lambda)b$ et situé sur la droite (d). $M(x_M, y_M)$ avec $x_M = \lambda a + (1 - \lambda)b$ et y_M à déterminer. De même, soient les points $A(a, f(a))$ et $B(b, f(b))$. La droite (d) est donc en fait la droite (AB).

Les points A , M et B étant alignés, les vecteurs $\overrightarrow{AM} \begin{pmatrix} x_M - a \\ y_M - f(a) \end{pmatrix}$ et $\overrightarrow{AB} \begin{pmatrix} b - a \\ f(b) - f(a) \end{pmatrix}$ sont colinéaires.

Donc $(x_M - a)(f(b) - f(a)) - (b - a)(y_M - f(a)) = 0$ Oui, un déterminant si vous voulez... Dans cette égalité, isolons notre inconnue y_M (et remplaçons x_M par son expression connue).

$$\text{D'où : } (b - a)(y_M - f(a)) = (\lambda a + (1 - \lambda)b - a)(f(b) - f(a))$$

$$\text{Autrement dit : } (b - a)(y_M - f(a)) = ((\lambda - 1)a + (1 - \lambda)b)(f(b) - f(a))$$

Il y a moyen de factoriser par $\lambda - 1$ à droite... Ben oui, $\lambda - 1 = -(1 - \lambda)$

$$\text{Donc : } (b - a)(y_M - f(a)) = (1 - \lambda)(b - a)(f(b) - f(a))$$

Ce serait sympa de simplifier par $b - a$ des deux côtés... Mais il faut avoir conscience de l'opération mathématique qui se cache derrière le mot « simplifier », pour pouvoir, le cas échéant, justifier son application si nécessaire. Ici en l'occurrence, il s'agit de diviser par $b - a$.

$$\text{Comme } b - a \neq 0, \text{ nous pouvons en déduire : } y_M - f(a) = (1 - \lambda)(f(b) - f(a))$$

$$\text{Puis : } y_M = (1 - \lambda)(f(b) - f(a)) + f(a) = (1 - \lambda)f(b) - (1 - \lambda)f(a) + f(a)$$

Remarquez que dans la dernière ligne, je ne me suis pas senti obligé de développer $(1 - \lambda)(f(b) - f(a))$ en quatre termes. Je n'avais pas d'intérêt à casser le $1 - \lambda$. Par contre, j'ai séparé le $f(a)$ et le $f(b)$ parce qu'ils sont séparés dans le résultat que j'escompte.

Donc : $y_M = (1 - \lambda)f(b) + (1 - 1 + \lambda)f(a) = \lambda f(a) + (1 - \lambda)f(b)$. Ouf!

Et rappelons : $f(\lambda a + (1 - \lambda)b) \leq y_M$. Finalement : $f(\lambda a + (1 - \lambda)b) \leq \lambda f(a) + (1 - \lambda)f(b)$

Il était possible d'aboutir à l'expression de y_M d'une autre manière, moins efficace (un brin plus calculatoire) mais peut-être plus intuitive pour certains : à partir des coordonnées des deux points $A(a, f(a))$ et $B(b, f(b))$, on pouvait déterminer l'équation de la droite (d). Par définition, son coefficient directeur est $\frac{f(b) - f(a)}{b - a}$. Quant à son ordonnée à l'origine (appelons-la β), on peut l'obtenir en se servant par exemple du fait que les coordonnées de A vérifient l'équation de (d). Autrement dit : $y_A = \frac{f(b) - f(a)}{b - a}x_A + \beta$

$$\text{Puis : } \beta = y_A - \frac{f(b) - f(a)}{b - a}x_A = f(a) - \frac{f(b) - f(a)}{b - a} \times a$$

$$(d) \text{ a donc pour équation : } y = \frac{f(b) - f(a)}{b - a}x + f(a) - \frac{f(b) - f(a)}{b - a} \times a \text{ Pas très joli...}$$

$$\text{Sous une forme plus condensée, (d) : } y = \frac{f(b) - f(a)}{b - a}(x - a) + f(a)$$

Pour obtenir y_M , il ne reste plus qu'à remplacer x par l'expression de x_M dans l'équation précédente et, après calculs, vous obtiendrez : $y_M = \lambda f(a) + (1 - \lambda)f(b)$

1)c) Que nous veut cette question ? L'inégalité établie dans la question précédente n'est-elle déjà pas valable pour tous a et b appartenant à I ? Non : pour l'instant, elle l'est uniquement dans le cas $a < b$.

Soient a et b appartenant à I et soit $\lambda \in [0; 1]$.

- Si $a < b$, l'inégalité demandée est établie d'après 1)b).
- Si $a = b$: d'une part, $f(\lambda a + (1 - \lambda)b) = f(\lambda a + (1 - \lambda)a) = f(a)$ et d'autre part : $\lambda f(a) + (1 - \lambda)f(b) = \lambda f(a) + (1 - \lambda)f(a) = f(a)$. L'inégalité demandée est donc établie (au sens large, c'est une égalité ici).
- Si $a > b$, en posant $\lambda' = 1 - \lambda$, $a' = b$ et $b' = a$, nous avons :
 $f(\lambda a + (1 - \lambda)b) = f((1 - \lambda')b' + \lambda'a') = f(\lambda'a' + (1 - \lambda')b')$ avec $\lambda' \in [0; 1]$ et $a' < b'$.
D'après 1)b) : $f(\lambda'a' + (1 - \lambda')b') \leq \lambda' f(a') + (1 - \lambda')f(b')$

1)b) est en effet valable pour tous a et b de I tels que $a < b$ (nos a' et b' conviennent donc) et pour tout $\lambda \in [0; 1]$ (notre λ' convient donc).

Autrement dit : $f(\lambda a + (1 - \lambda)b) \leq \lambda f(a) + (1 - \lambda)f(b)$

Nous avons bien établi : $\forall \lambda \in [0; 1], \forall a, b \in I, f(\lambda a + (1 - \lambda)b) \leq \lambda f(a) + (1 - \lambda)f(b)$

2) Le fait qu'il faille appliquer l'inégalité établie précédemment semble assez évident... La fonction exponentielle est convexe sur $\mathbb{R}$ (dérivable deux fois, de dérivée seconde elle-même, donc strictement positive sur $\mathbb{R}$).

D'après 1)c) : $\forall \lambda \in [0; 1], e^{\lambda a + (1 - \lambda)b} \leq \lambda e^a + (1 - \lambda)e^b$

*Le calme après la tempête.
Et avant la suivante...*

3)a) Commençons par exprimer X en fonction de Y .

$$Y = \frac{1}{2} - \frac{X}{2c} \text{ donc } \frac{X}{2c} = \frac{1}{2} - Y \text{ puis } X = c - 2cY = (1 - 2Y)c$$

Ces histoires de Y et $1 - Y$ que l'on me demande de faire apparaître... Mais moi, je vois du $1 - 2Y$...

$$\text{D'où : } X = (1 - Y - Y)c = (1 - Y)c - Yc = Y \times (-c) + (1 - Y) \times c$$

Une transformation astucieuse à laquelle il n'est pas évident de penser et qui, je le confesse, était indiquée dans l'énoncé originel par le biais d'une question intermédiaire.

A ce stade, il est tentant de passer à l'exponentielle et d'appliquer l'inégalité obtenue en 2). Mais attention, cette inégalité est valable (avec les bonnes hypothèses) pour des réels, pas a priori pour des variables aléatoires... Que faire ? Revenir simplement à la définition de $A \leq B$ pour deux variables aléatoires A et B : $\forall \omega \in \Omega, A(\omega) \leq B(\omega)$ (ça tombe bien, $A(\omega)$ et $B(\omega)$ sont, eux, des réels)

L'égalité de variables aléatoires précédente nous permet d'affirmer :

$$\forall \omega \in \Omega, X(\omega) = Y(\omega) \times (-c) + (1 - Y(\omega)) \times c$$

$$\text{Puis : } \forall \omega \in \Omega, e^{X(\omega)} = \exp(Y(\omega) \times (-c) + (1 - Y(\omega)) \times c)$$

Ce serait bien que $Y(\omega)$ appartienne à $[0; 1]$...

Par ailleurs, X est à valeurs dans $[-c; c]$. Autrement dit : $\forall \omega \in \Omega, -c \leq X(\omega) \leq c$

$$\text{Et : } Y(\omega) = \frac{1}{2} - \frac{X(\omega)}{2c}, \text{ avec } -\frac{1}{2} \leq -\frac{X(\omega)}{2c} \leq \frac{1}{2} \text{ (car } -2c < 0) \text{ puis } \frac{1}{2} - \frac{1}{2} \leq \frac{1}{2} - \frac{X(\omega)}{2c} \leq \frac{1}{2} + \frac{1}{2}$$

C'est-à-dire : $\forall \omega \in \Omega, 0 \leq Y(\omega) \leq 1$. L'inégalité obtenue en 2) nous permet alors d'affirmer : $\forall \omega \in \Omega, e^{X(\omega)} = \exp(Y(\omega) \times (-c) + (1 - Y(\omega)) \times c) \leq Y(\omega)e^{-c} + (1 - Y(\omega))e^c$

Autrement dit (en revenant aux variables aléatoires) : $e^X \leq Y e^{-c} + (1 - Y)e^c$

3)b) A partir de l'inégalité 3a), la croissance de l'espérance fournit :

$$E(e^X) \leq E(Y e^{-c} + (1 - Y)e^c).$$

Puis, e^c et e^{-c} étant des constantes, la linéarité de l'espérance nous permet d'obtenir :

$$E(e^X) \leq e^{-c} E(Y) + e^c E(1 - Y) = e^{-c} E(Y) + e^c (1 - E(Y))$$

Or, $Y = \frac{1}{2} - \frac{X}{2c}$ donc (encore par linéarité) :

$$E(Y) = \frac{1}{2} - \frac{1}{2c} E(X) = \frac{1}{2} \text{ car } E(X) = 0 \text{ par hypothèse.}$$

Par suite : $E(e^X) \leq e^{-c} \frac{1}{2} + e^c \times (1 - \frac{1}{2})$, et enfin : $E(e^X) \leq \frac{e^c + e^{-c}}{2}$

Entre autres, ce que j'ai trouvé rigolo dans cet exercice - largement remanié - c'est le fait, à partir de l'expression $\exp(Y(\omega) \times (-c) + (1 - Y(\omega)) \times c)$, de d'abord « sortir » $Y(\omega)$ et $1 - Y(\omega)$ de l'exponentielle par l'inégalité de convexité, puis, en passant à l'espérance, de « sortir » cette fois e^c et e^{-c} par linéarité. A chaque contexte, à chaque situation ses prisonniers à libérer...*

**La longue question 1), notamment, est un pur ajout, visant ensuite à pouvoir vous faire profiter de cette inégalité de convexité pour la suite de l'exercice. Cette question 1) est un Pokémon que cette version PCSI a en commun avec un exercice de la version MPSI ...*

Exercice 14

*Souvenir d'un été parsemé de dimanches,
la boule retirée de l'urne est rouge ou blanche*

Énoncé : (temps conseillé : 1h 10 min) (**)

d'après CCINP 2021 PC

Commençons par présenter le principe de raisonnement par récurrence forte, que l'on pourra être amené à utiliser. Pour montrer que, pour tout $n \in \mathbb{N}$, la propriété P_n est vraie, la récurrence simple se faisait ainsi :

Initialisation : On montre que P_0 est vraie.

Hérédité : On montre que si, pour un certain $n \in \mathbb{N}$, P_n est vraie, alors P_{n+1} est vraie.

Conclusion : Cela nous permet de conclure que pour tout entier naturel n , P_n est vraie.

Le principe de raisonnement par récurrence forte est le suivant :

Initialisation : On montre que P_0 et P_1 sont vraies.

Hérédité : On montre que si, pour un certain $n \in \mathbb{N}$, pour tout $k \in \llbracket 0; n \rrbracket$, P_k est vraie, alors P_{n+1} aussi est vraie.

Conclusion : Cela nous permet de conclure que pour tout entier naturel n , P_n est vraie.

On fixe un couple d'entiers $(b, r) \in \mathbb{N}^* \times \mathbb{N}^*$. On dispose d'un stock illimité de boules blanches et de boules rouges et on considère une urne contenant initialement b boules blanches et r boules rouges indiscernables au toucher. On procède à des tirages successifs dans cette urne en respectant à chaque fois le protocole suivant :

1. Si la boule tirée est de couleur blanche, on la replace dans l'urne et on ajoute une boule blanche supplémentaire.
2. Si la boule tirée est de couleur rouge, on la replace dans l'urne et on ajoute une boule rouge supplémentaire.

Pour tout $n \in \mathbb{N}^*$, on désigne par X_n la variable aléatoire prenant la valeur 1 si la boule tirée au n -ième tirage est blanche, et la valeur 0 si cette boule est rouge. On considère également la suite de variables aléatoires réelles $(S_n)_{n \in \mathbb{N}}$ définie par :

$$S_0 = b \text{ et } \forall n \in \mathbb{N}^*, S_n = b + \sum_{k=1}^n X_k$$

- 1) Déterminer la loi de X_1 , puis la loi de X_2 .

2) Soit $n \in \mathbb{N}$. Que représente la variable aléatoire S_n ? Quel est l'ensemble des valeurs prises par la variable aléatoire S_n ?

3) Soit $n \in \mathbb{N}^*$. Pour tout $k \in \llbracket b; n+b \rrbracket$, déterminer $P_{[S_n=k]}(X_{n+1} = 1)$.

4) Montrer que, pour tout $n \in \mathbb{N}^*$: $P(X_{n+1} = 1) = \frac{E(S_n)}{b+r+n}$

5) Montrer que pour tout $n \in \mathbb{N}^*$, X_n suit la loi de Bernoulli de paramètre $\frac{b}{b+r}$.

Remarques sur l'énoncé :

Le principe de récurrence forte présenté par l'énoncé diffère de la récurrence simple en ce sens : à l'étape d'hérédité, on ne suppose pas juste que P_n est vraie pour aboutir à P_{n+1} , mais on suppose plutôt que $P_0, P_1, P_2, \dots, P_n$ sont vraies pour aboutir à P_{n+1} . Cela reviendrait à montrer, par récurrence simple, que pour tout $n \in \mathbb{N}$, la propriété $Q_n : \ll \forall k \in \llbracket 0; n \rrbracket, P_k \text{ est vraie} \gg$ est vraie

Profitons-en pour rappeler (ou faire découvrir) le sens de la notation $E \times F$, lorsque E et F sont deux ensembles. $E \times F$, appelé produit cartésien de E et F , est l'ensemble des couples (a, b) avec $a \in E$ et $b \in F$.

« $(b, r) \in \mathbb{N}^* \times \mathbb{N}^*$ » veut donc dire la même chose que « $b \in \mathbb{N}^*$ et $r \in \mathbb{N}^*$ »

Cas particulier : si E est un ensemble, E^2 est l'ensemble des couples (a, b) avec $a \in E$ et $b \in E$. On note aussi $E^2 = E \times E$. Dans notre cas, il était donc possible d'écrire plutôt $(b, r) \in (\mathbb{N}^*)^2$ (les parenthèses autour de $\mathbb{N}^*$ étant là juste pour faire plus « propre »)

Dans la question 4), $E(S_n)$ désigne l'espérance de S_n .

Correction de l'exercice 14 :

1) $X_1(\Omega) \subset \{0; 1\}$ donc X_1 suit une loi de Bernoulli de paramètre $p = P(X_1 = 1)$. Reste à déterminer ce paramètre.

$X_1(\Omega)$ est simplement l'ensemble des valeurs prises par la variable aléatoire X_1 .

Sans rentrer dans trop de détails théoriques, Ω est l'univers considéré.

L'événement $[X_1 = 1]$ est réalisé si et seulement si la première boule tirée est blanche. Or, l'urne contient initialement b boules blanches et r rouges, et il y a équiprobabilité de tirer chacune de ces boules, indiscernables au toucher. Donc $P(X_1 = 1) = \frac{b}{b+r}$

En conclusion, X_1 suit une loi de Bernoulli de paramètre $p = \frac{b}{b+r}$

Comme j'ai lu l'énoncé en entier jusqu'à la question 5) (donnez-moi un cookie), je suis rassuré.

X_2 suit aussi une loi de Bernoulli, de paramètre $p' = P(X_2 = 1)$

Mais comment calculer cette probabilité ? La couleur de la boule tirée au deuxième tirage dépend du résultat du premier tirage, non ? Si, parfaitement. D'où l'idée suivante :

Les événements $[X_1 = 0]$ et $[X_1 = 1]$ forment un système complet d'événements.

Dans le supérieur, et en particulier en prépa, vous parlerez plus souvent de « système complet d'événements » que de « partition de l'univers », qui désignent la même chose (même s'il pourra vous arriver de recroiser la seconde formulation, souvenir heureux de vos années lycée).

D'après la formule des probabilités totales :

$$P(X_2 = 1) = P([X_1 = 0] \cap [X_2 = 1]) + P([X_1 = 1] \cap [X_2 = 1])$$

$$\text{Puis : } P(X_2 = 1) = P(X_1 = 0) \times P_{[X_1=0]}(X_2 = 1) + P(X_1 = 1) \times P_{[X_1=1]}(X_2 = 1)$$

Rappelons que pour deux événements A et B , si $P(A) \neq 0$, $P_A(B) = \frac{P(A \cap B)}{P(A)}$. D'où $P(A \cap B) = P(A) \times P_A(B)$

Ici, on a bien $P(X_1 = 0) \neq 0$ et $P(X_1 = 1) \neq 0$, car $P(X_1 = 1) = \frac{b}{b+r}$ et $P(X_1 = 0) = \frac{r}{b+r}$ avec $b, r \in \mathbb{N}^*$

Or, si l'événement $[X_1 = 0]$ est réalisé (autrement dit, la première boule tirée est rouge),

la composition de l'urne juste avant le deuxième tirage est la suivante : toujours b boules blanches, mais $r + 1$ boules rouges (puisqu'on en a rajouté une).

$$\text{Donc } P_{[X_1=0]}(X_2 = 1) = \frac{b}{b+r+1}.$$

De même, si l'événement $[X_1 = 1]$ est réalisé (autrement dit, la première boule tirée est blanche), la composition de l'urne juste avant le deuxième tirage est la suivante : $b + 1$ boules blanches, et toujours r boules rouges.

$$\text{Donc } P_{[X_1=1]}(X_2 = 1) = \frac{b+1}{b+r+1}$$

$$\text{D'où : } P(X_2 = 1) = \frac{r}{b+r} \times \frac{b}{b+r+1} + \frac{b}{b+r} \times \frac{b+1}{b+r+1} = \frac{rb + b(b+1)}{(b+r)(b+r+1)} = \frac{b(r+b+1)}{(b+r)(b+r+1)}$$

$$\text{Puis : } P(X_2 = 1) = \frac{b}{b+r}$$

X_2 suit donc aussi une loi de Bernoulli de paramètre $p = \frac{b}{b+r}$

Je l'avais appelé p' au départ parce que je n'avais pas encore montré que c'est le même. On a bien $p' = p$, comme attendu en lisant la question 5)

$$2) S_0 = b \text{ et } : \forall n \in \mathbb{N}^*, S_n = b + \sum_{k=1}^n X_k$$

Avant les tirages, il y a b boules blanches dans l'urne. Chaque fois qu'une boule blanche est tirée (c'est-à-dire chaque fois qu'un événement $[X_k = 1]$ - avec $k \in \mathbb{N}^*$ - est réalisé), l'urne gagne une boule blanche.

Pour tout $n \in \mathbb{N}$, S_n représente donc le nombre de boules blanches dans l'urne après le n -ième tirage.

A chaque tirage, il est toujours possible d'obtenir une boule blanche ou une boule rouge (aucune de ces deux couleurs ne peut disparaître de l'urne).

L'ensemble des valeurs prises par S_n est donc $S_n(\Omega) = \llbracket b ; b+n \rrbracket$

Autrement dit, l'ensemble des entiers naturels de b à $b+n$. Aux b boules blanches initialement présentes peuvent s'être ajoutées, au bout de n tirages : $0, 1, \dots, n$ autres boules blanches.

3) Pour tout $k \in \llbracket b, b+n \rrbracket$, l'événement $[S_n = k]$ (de probabilité non nulle) correspond à une urne composée, juste avant le $n+1$ -ème tirage, de k boules blanches et $b+r+n$

boules au total. En effet, après chaque tirage, l'urne gagne 1 boule, qu'elle soit blanche ou rouge : au bout de n tirages, elle en a donc gagné n .

Le nombre de boules rouges ne m'intéresse pas particulièrement, mais si j'en avais besoin, j'aurais tout simplement dit : $b + r + n - k$

La probabilité de tirer une boule blanche d'une urne avec une telle composition est donc : $\frac{k}{b+r+n}$

Autrement dit : $\forall k \in [b, n+b], P_{[S_n=k]}(X_{n+1}=1) = \frac{k}{b+r+n}$

4) Soit $n \in \mathbb{N}^*$.

On nous a fait calculer $P_{[S_n=k]}(X_{n+1}=1) = \frac{k}{b+r+n}$ pour tout $k \in [b, n+b]$

Comment accéder, à partir de ces informations, à $P(X_{n+1}=1)$? La formule des probabilités totales est notre amie, comme elle le sera souvent dans ce genre de situation.

$([S_n = k])_{k \in [b, n+b]}$ est un système complet d'événements.

D'après la formule des probabilités totales : $P(X_{n+1}=1) = \sum_{k=b}^{n+b} P([S_n = k] \cap [X_{n+1}=1])$

Sans le signe somme, ça donne tout simplement : $P(X_{n+1}=1)$

$= P([S_n = b] \cap [X_{n+1}=1]) + P([S_n = b+1] \cap [X_{n+1}=1]) + \dots + P([S_n = b+n] \cap [X_{n+1}=1])$

Donc : $P(X_{n+1}=1) = \sum_{k=b}^{n+b} P(S_n = k) \times P_{[S_n=k]}(X_{n+1}=1)$. Puis, d'après 3) :

$P(X_{n+1}=1) = \sum_{k=b}^{n+b} P(S_n = k) \times \frac{k}{b+r+n}$ Et $\frac{1}{b+r+n}$ est constant vis-à-vis de k ...

Par suite : $P(X_{n+1}=1) = \frac{1}{b+r+n} \sum_{k=b}^{n+b} k \times P(S_n = k)$ Oh, mais cette dernière somme...

Enfin : $P(X_{n+1}=1) = \frac{1}{b+r+n} \times E(S_n)$

La dernière somme portait bien sur l'ensemble des valeurs k prises par S_n .

Nous avons bien montré : $\forall n \in \mathbb{N}^*, P(X_{n+1}=1) = \frac{E(S_n)}{b+r+n}$

5) Enfin cette question 5)... Bon, l'énoncé nous a infligé une longue tartine sur la

réurrence forte. C'est peut-être le moment de l'utiliser...

L'énoncé nous a présenté ce raisonnement pour montrer que P_n est vraie pour tout $n \in \mathbb{N}$.

Ici, c'est plutôt : pour tout $n \in \mathbb{N}^$. Pas de panique ! Un très léger effort d'adaptation suffira.*

Soit, pour tout $n \in \mathbb{N}^*$, la propriété P_n :

« X_n suit la loi de Bernoulli de paramètre $\frac{b}{b+r}$ »

Montrons par récurrence forte que pour tout entier naturel $n \in \mathbb{N}^*$, P_n est vraie.

Initialisation : d'après la question 1), X_n suit la loi de Bernoulli de paramètre $\frac{b}{b+r}$ donc P_1 est vraie.

Hérédité : Supposons que pour un certain $n \in \mathbb{N}^*$: pour tout $k \in \llbracket 1; n \rrbracket$, P_k est vraie, et montrons que P_{n+1} aussi est vraie.

Supposons donc que : pour tout $k \in \llbracket 1; n \rrbracket$, X_k suit la loi de Bernoulli de paramètre $\frac{b}{b+r}$.

Comment à partir de ça, accéder à la loi de X_{n+1} ?

$X_{n+1}(\Omega) \subset \{0; 1\}$ donc X_{n+1} suit une loi de Bernoulli, de paramètre $P(X_{n+1} = 1)$.

Or, d'après 4) : $P(X_{n+1} = 1) = \frac{E(S_n)}{b+r+n}$

Et $E(S_n) = E(b + \sum_{k=1}^n X_k) = b + E(\sum_{k=1}^n X_k) = b + \sum_{k=1}^n E(X_k)$ par linéarité de l'espérance.

Or, par hypothèse de récurrence, pour tout k appartenant à $\llbracket 1; n \rrbracket$, X_k suit la loi de Bernoulli de paramètre $\frac{b}{b+r}$, donc $E(X_k) = \frac{b}{b+r}$

C'est là que nous avons eu besoin de l'information sur tous les $X_1, \dots, X_n$ précédant X_{n+1} (et pas seulement le prédécesseur immédiat X_n comme dans le cas d'une récurrence simple). D'où l'intérêt de la récurrence forte.

D'où : $E(S_n) = b + \sum_{k=1}^n \frac{b}{b+r} = b + n \times \frac{b}{b+r} = \frac{b(b+r) + nb}{b+r} = \frac{b(b+r+n)}{b+r}$

Puis : $P(X_{n+1} = 1) = \frac{E(S_n)}{b+r+n} = \frac{b(b+r+n)}{b+r} \times \frac{1}{b+r+n} = \frac{b}{b+r}$

Enfin, X_{n+1} suit bien la loi de Bernoulli de paramètre $\frac{b}{b+r}$ et P_{n+1} est vraie.

Conclusion : Le principe de raisonnement par récurrence forte nous permet de conclure que pour tout entier naturel $n \in \mathbb{N}^*$, P_n est vraie.

Autrement dit, pour tout $n \in \mathbb{N}^*$, X_n suit la loi de Bernoulli de paramètre $\frac{b}{b+r}$.

Exercice 15

*Aspirants députés, voici de quoi vous plaire :
les cas d'égalité dans la triangulaire.*

Énoncé : (temps conseillé : 50 min) (****) d'après Centrale 2023 PC Maths 1

On rappelle l'inégalité triangulaire : $\forall z, z' \in \mathbb{C}, |z + z'| \leq |z| + |z'|$.

- 1) Soit z un nombre complexe tel que $|1 + z| = 1 + |z|$. Montrer que $z \in \mathbb{R}_+$.
- 2) En déduire que, si z et z' sont deux nombres complexes vérifiant $|z + z'| = |z| + |z'|$ et $z \neq 0$, alors : $\exists \alpha \in \mathbb{R}_+, z' = \alpha z$.
- 3) Soit n un entier supérieur ou égal à 2 et $z_1, \dots, z_n$ des nombres complexes non tous nuls tels que $\left| \sum_{j=1}^n z_j \right| = \sum_{j=1}^n |z_j|$.

Montrer qu'il existe un réel θ tel que : $\forall k \in \llbracket 1; n \rrbracket, z_k = e^{i\theta} |z_k|$.

Dans le cas où $z_1 \neq 0$, on pourra appliquer le résultat de la question précédente aux couples (z_1, z_k) pour $k \in \llbracket 2; n \rrbracket$.

Remarques sur l'énoncé :

Oui, l'inégalité triangulaire déjà rappelée en début d'énoncé de l'exercice 8 est plus généralement valable sur $\mathbb{C}$ (en parlant cette fois de module plutôt que de valeur absolue). Là encore, elle se généralise à une somme de plusieurs termes : $\left| \sum_{j=1}^N z_k \right| \leq \sum_{j=1}^N |z_k|$. Cet exercice a pour but de vous faire réfléchir sur le cas d'égalité de cette inégalité triangulaire.

Dans la question 3), l'énoncé dit que $z_1, \dots, z_n$ sont « non tous nuls ». Cela veut tout simplement dire qu'ils ne sont pas tous nuls. Autrement dit, qu'au moins l'un d'entre eux est non nul. Mais ça n'interdit pas à certains d'entre eux d'être éventuellement non nuls. A ne pas confondre avec « tous non nuls » qui aurait voulu dire qu'aucun d'entre eux n'est nul.

« Tous non nuls » implique « non tous nuls » mais « non tous nuls » n'implique pas nécessairement « tous non nuls »

Correction de l'exercice 15 :

1) Le réflexe de beaucoup ici sera d'écrire z sous forme algébrique $z = a + ib$ (avec a et b réels) puis d'exprimer le module (quitte à élever au carré) en fonction de a et b . Il y a légèrement plus élégant (et moins calculatoire), en se souvenant que $z\bar{z} = |z|^2$...

En élevant au carré l'égalité vérifiée par z , nous obtenons : $|1 + z|^2 = 1 + 2|z| + |z|^2$

Autrement dit : $(1 + z)(\overline{1 + z}) = 1 + 2|z| + z\bar{z}$. Ou encore : $(1 + z)(1 + \bar{z}) = 1 + 2|z| + z\bar{z}$

Un petit rappel sur les propriétés de la conjugaison ? C'est par ici.

Puis : $1 + \bar{z} + z + z\bar{z} = 1 + 2|z| + z\bar{z}$, ce qui donne $z + \bar{z} = 2|z|$. Or, $z + \bar{z} = 2 \operatorname{Re}(z)$

Si vous l'aviez oublié, passez à la forme algébrique, vous verrez les parties imaginaires de z et $\bar{z}$, opposées l'une de l'autre, se simplifier, et leurs parties réelles, égales, s'additionner. On a de même : $z - \bar{z} = 2i \operatorname{Im}(z)$

Donc : $2 \operatorname{Re}(z) = 2|z|$, c'est-à-dire $\operatorname{Re}(z) = |z|$. D'où : $\operatorname{Re}^2(z) = |z|^2 = \operatorname{Re}^2(z) + \operatorname{Im}^2(z)$

Bon, j'avoue, cela revient à parler de forme algébrique de z . Mais au moins, on ne s'est pas coltiné, dès le début, du $z = a + ib$ qui aurait légèrement alourdi nos calculs.

$\operatorname{Re}^2(z) = \operatorname{Re}^2(z) + \operatorname{Im}^2(z)$ donc $\operatorname{Im}^2(z) = 0$ puis $\operatorname{Im}(z) = 0$. z est donc un nombre réel.

De plus : $\operatorname{Re}(z) = |z| \geq 0$. z étant réel, nous avons : $z = \operatorname{Re}(z) \geq 0$. Enfin, $z \in \mathbb{R}_+$.

2) Soient z et z' sont deux nombres complexes vérifiant $|z + z'| = |z| + |z'|$ et $z \neq 0$
« En déduire ». Quel lien avec ce qui précède ? $z \neq 0$, ça me permet de faire quelque chose...

z étant non nul, nous pouvons écrire : $\left|z\left(1 + \frac{z'}{z}\right)\right| = |z|\left(1 + \frac{|z'|}{|z|}\right)$.

Autrement dit : $|z| \times \left|1 + \frac{z'}{z}\right| = |z|\left(1 + \frac{|z'|}{|z|}\right)$. Puis ($|z|$ étant non nul) : $\left|1 + \frac{z'}{z}\right| = \left(1 + \frac{|z'|}{|z|}\right)$

Ou encore : $\left|1 + \frac{z'}{z}\right| = 1 + \left|\frac{z'}{z}\right|$ Oh mais dis-donc...

De 1), nous pouvons conclure : $\frac{z'}{z} \in \mathbb{R}_+$. Il existe donc un réel positif α tel que $\frac{z'}{z} = \alpha$

Ben oui, en 1), nous avons établi que tout nombre complexe z vérifiant $|1 + z| = 1 + |z|$ est en fait un réel positif. Peu importe, donc, qu'il s'appelle ici non pas z mais $\frac{z'}{z}$.

En conclusion, si z et z' sont deux nombres complexes vérifiant $|z + z'| = |z| + |z'|$ et $z \neq 0$, alors : $\exists \alpha \in \mathbb{R}_+, z' = \alpha z$.

3) Attention : le θ dont il faut montrer l'existence ici doit être le même pour tous les z_k .

$z_1, \dots, z_n$ sont des nombres complexes non tous nuls. L'un au moins d'entre eux est donc non nul. Sans perte de généralité, supposons $z_1 \neq 0$.

Comment ça, « sans perte de généralité » ? Qu'est-ce que c'est que cette expression douteuse ? Pas d'inquiétude ici, il n'y a aucune arnaque. Les $z_1, \dots, z_n$ étant non tous nuls, on peut toujours, quitte à les réordonner, se débrouiller pour que z_1 soit non nul. Et leur ordre n'intervient pas dans le résultat demandé.

Ce « sans perte de généralité » promet au lecteur que malgré l'ajout d'une hypothèse (ici $z_1 \neq 0$), la portée générale du résultat n'est en rien remise en cause.

A utiliser avec énormément de modération, uniquement dans les cas où vous avez l'assurance - comme moi ici - que l'hypothèse que vous rajoutez ne fait perdre aucune généralité au résultat. Au moindre doute, ne le faites pas, quitte à opter pour un raisonnement plus détaillé - et peut-être plus fastidieux. A une colle ou un oral, si vous utilisez un tel raccourci verbal, vous devez être en mesure de justifier, sur demande, qu'il n'y a effectivement pas de perte de généralité.

En quoi aurait consisté un raisonnement plus détaillé ici ? En le fait, par exemple, d'expliquer la réindexation : si $z_1 \neq 0$, pas besoin de réindexer. Sinon, $z_1, \dots, z_n$ étant non tous nuls (et z_1 étant nul), il existe $k \in \llbracket 2; n \rrbracket$ tel que $z_k \neq 0$. Définissons alors les complexes $z'_1, z'_2, \dots, z'_n$ ainsi : $z'_1 = z_k$, $z'_k = z_1$, et pour tout $j \in \llbracket 2; n \rrbracket$, si $j \notin \{1; k\}$, $z'_j = z_j$ (autrement dit, on permute z_1 et z_k et on laisse les autres tranquilles). La permutation effectuée ne modifiant pas la valeur des deux sommes, on a toujours $\left| \sum_{j=1}^n z'_j \right| = \sum_{j=1}^n |z'_j|$.

Reprenons le fil de notre rédaction.

Pour tout $k \in \llbracket 2; n \rrbracket$, $\left| \sum_{j=1}^n z_j \right| = \left| z_1 + z_k + \sum_{\substack{j=2 \\ j \neq k}}^n z_j \right|$

$\sum_{\substack{j=2 \\ j \neq k}}^n z_j$, c'est tout simplement la somme $\sum_{j=2}^n z_j$ à laquelle on a retiré le terme z_k .

Si $n = 2$, il n'y a aucun entier j différent de k entre 2 et n , et la somme $\sum_{\substack{j=2 \\ j \neq k}}^n z_j$ est nulle par convention. D'accord, mais quel intérêt d'avoir sorti z_1 et z_k de la somme ? Nous mettre en situation d'utiliser l'indication fournie par l'énoncé.

Pour tout $k \in \llbracket 2 ; n \rrbracket$, $\left| \sum_{j=1}^n z_j \right| = \left| z_1 + z_k + \sum_{\substack{j=2 \\ j \neq k}}^n z_j \right| \leq |z_1 + z_k| + \left| \sum_{\substack{j=2 \\ j \neq k}}^n z_j \right|$ d'après l'inégalité triangulaire. La « simple », celle pour une somme de deux termes, ici en l'occurrence $z_1 + z_k$ d'une part et $\sum_{\substack{j=2 \\ j \neq k}}^n z_j$ d'autre part.

De plus : $|z_1 + z_k| \leq |z_1| + |z_k|$. Donc : $\left| \sum_{j=1}^n z_j \right| \leq |z_1 + z_k| + \left| \sum_{\substack{j=2 \\ j \neq k}}^n z_j \right| \leq |z_1| + |z_k| + \left| \sum_{\substack{j=2 \\ j \neq k}}^n z_j \right|$

Puis, toujours par l'inégalité triangulaire (sa généralisation à une somme, cette fois) :

$\left| \sum_{\substack{j=2 \\ j \neq k}}^n z_j \right| \leq \sum_{\substack{j=2 \\ j \neq k}}^n |z_j|$. Donc, si l'on récapitule les inégalités successives obtenues :

$$\left| \sum_{j=1}^n z_j \right| \leq |z_1 + z_k| + \left| \sum_{\substack{j=2 \\ j \neq k}}^n z_j \right| \leq |z_1| + |z_k| + \left| \sum_{\substack{j=2 \\ j \neq k}}^n z_j \right| \leq |z_1| + |z_k| + \sum_{\substack{j=2 \\ j \neq k}}^n |z_j|$$

Autrement dit : $\left| \sum_{j=1}^n z_j \right| \leq |z_1 + z_k| + \left| \sum_{\substack{j=2 \\ j \neq k}}^n z_j \right| \leq |z_1| + |z_k| + \left| \sum_{\substack{j=2 \\ j \neq k}}^n z_j \right| \leq \sum_{j=1}^n |z_j|$ (*)

C'est bien joli tout ça, mais pourquoi avoir conservé la trace de toutes ces majorations successives ? Parce qu'un phénomène rigolo va se produire : rappelons que par hypothèse, le membre tout à gauche et le membre tout à droite sont égaux...

Par hypothèse : $\left| \sum_{j=1}^n z_j \right| = \sum_{j=1}^n |z_j|$. Donc les quatre membres de (*) sont égaux.

En particulier : $|z_1 + z_k| + \left| \sum_{\substack{j=2 \\ j \neq k}}^n z_j \right| = |z_1| + |z_k| + \left| \sum_{\substack{j=2 \\ j \neq k}}^n z_j \right|$. Et en simplifiant : $|z_1 + z_k| = |z_1| + |z_k|$

z_1 et z_k vérifient donc les hypothèses de la question 2) (z_1 étant bien non nul).

Pour tout $k \in \llbracket 2 ; n \rrbracket$, il existe donc un réel positif α_k tel que $z_k = \alpha_k z_1$.

C'est aussi vrai pour $k = 1$, en prenant tout simplement $\alpha_1 = 1 : z_1 = 1 \times z_1$

Attention : il n'y a aucune raison a priori que le α soit le même pour tous les z_k . D'où

la pertinence de l'appeler α_k . Rien de dérangentant ; nous sommes à deux doigts du résultat...

Notons θ un argument du nombre complexe non nul z_1 .

La mention « non nul » est utile : 0 n'a pas d'argument.

Sous forme exponentielle, z_1 s'écrit : $z_1 = |z_1|e^{i\theta}$. Donc : $\forall k \in \llbracket 1 ; n \rrbracket$, $z_k = \alpha_k |z_1| e^{i\theta}$

Et $|z_k| = \alpha_k |z_1|$ (car $|e^{i\theta}| = 1$). Enfin : $\forall k \in \llbracket 1 ; n \rrbracket$, $z_k = |z_k| e^{i\theta}$

Nous avons bien établi l'existence d'un réel θ tel que : $\forall k \in \llbracket 1 ; n \rrbracket$, $z_k = e^{i\theta} |z_k|$.

Autrement dit, tous les z_k non nuls ont un argument commun (ce fameux θ). Géométriquement, cela se traduit sur le plan complexe par le fait que les points M_k d'affixes les z_k sont sur une même demi-droite partant de l'origine du repère.

Exercice 16

- Quel est cet énoncé ? C'est un peu court jeune homme !

- Il fait dire déjà bien des choses aux sommes...

Énoncé : (temps conseillé : 25 min) (****) d'après Centrale 2022 PC Maths 2

On pourra librement utiliser la formule du binôme de Newton :

$$\forall a, b \in \mathbb{R}, \forall n \in \mathbb{N}, (a + b)^n = \sum_{k=0}^n \binom{n}{k} a^k b^{n-k} = \sum_{k=0}^n \binom{n}{k} a^{n-k} b^k$$

Soit $n \in \mathbb{N}^*$. En développant $(1 + x)^n$ pour deux réels x bien choisis, montrer que

$$\sum_{p=0}^{\lfloor n/2 \rfloor} \binom{n}{2p} = 2^{n-1}.$$

Remarques sur l'énoncé :

Une fois n'est pas coutume, et comme relevé par Cyrano, un énoncé particulièrement bref. La formule du binôme de Newton rappelée par l'énoncé fait figurer deux sommes. On peut vérifier qu'elles sont égales par changement d'indice $j = n - k$ (plus de détails dans [cette vidéo](#) de la playlist sur le signe $\sum$) en rappelant $\binom{n}{n-j} = \binom{n}{j}$. Par ailleurs, il est tout à fait normal que l'on puisse intervertir a et b comme c'est le cas entre les deux sommes, puisque $(a + b)^n = (b + a)^n$. Quant à n , constant dans la somme, il n'est évidemment pas affecté par le changement d'indice.

Cette formule du binôme de Newton se démontre par récurrence sur n , à coups de manipulations sur les sommes (changements d'indice, expulsions de termes) et les coefficients binomiaux (formule du triangle de Pascal). Elle est plus généralement valable pour a et b complexes.

Revenez [ici](#) si vous estimez qu'un rappel sur la partie entière est nécessaire. Au cas où (ce n'est peut-être pas évident en termes de lisibilité), la borne du haut de la somme à calculer est $\left\lfloor \frac{n}{2} \right\rfloor$

Correction de l'exercice 16 :

La somme à calculer, avec du $2p$ en bas du coefficient binomial, me fait peur... Qui sont les $2p$ lorsque p se balade de 0 à $\left\lfloor \frac{n}{2} \right\rfloor$? Ce sont « tout simplement » les entiers pairs de 0 à n , comme nous le montrerons plus loin. Chaque chose en son temps : commençons par développer $(1+x)^n$ comme le préconise l'énoncé.

D'après la formule du binôme de Newton, pour tout réel x :

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

« Pour deux réels x bien choisis », nous précise l'énoncé. Lesquels? Déjà, pour $x = 1$, on se retrouve avec une somme dont le terme général est un coefficient binomial seul...

En particulier, pour $x = 1$: $(1+1)^n = \sum_{k=0}^n \binom{n}{k} 1^k$. Autrement dit : $\sum_{k=0}^n \binom{n}{k} = 2^n$

D'accord, mais je ne vois toujours pas ce qui pourrait bien m'amener à cette histoire de nombre pairs. Peut-être une alternance de signes, que pourrait par exemple m'offrir $x = -1$...

Et pour $x = -1$: $(1-1)^n = \sum_{k=0}^n \binom{n}{k} (-1)^k$. Autrement dit : $\sum_{k=0}^n \binom{n}{k} (-1)^k = 0^n = 0$ (car $n \in \mathbb{N}^*$)

Précision utile car $0^0 = 1$ par convention...

Que faire avec les deux sommes obtenues? Comment ne garder que des coefficients binomiaux dont le terme du bas est pair? En sommant ces deux sommes!

$$\text{Or : } \sum_{k=0}^n \binom{n}{k} + \sum_{k=0}^n \binom{n}{k} (-1)^k = \sum_{k=0}^n \left[\binom{n}{k} + \binom{n}{k} (-1)^k \right] = \sum_{k=0}^n \binom{n}{k} (1 + (-1)^k)$$

Quel intérêt? Je vous réponds par une question : que pouvez-vous me dire sur $1 + (-1)^k$?

Pour tout entier naturel k : $(-1)^k = 1$ si k est pair, et $(-1)^k = -1$ si k est impair.

Donc : $1 + (-1)^k = 2$ si k est pair, et $1 + (-1)^k = 0$ si k est impair.

« Cassons » donc la somme $\sum_{k=0}^n \binom{n}{k} (1 + (-1)^k)$ en deux : une somme qui ne contiendra que les termes d'indice k pair, et une somme qui ne contiendra que les termes d'indice k impair.

$$\begin{aligned}\sum_{k=0}^n \binom{n}{k} (1 + (-1)^k) &= \sum_{\substack{k=0 \\ k \text{ pair}}}^n \binom{n}{k} (1 + (-1)^k) + \sum_{\substack{k=0 \\ k \text{ impair}}}^n \binom{n}{k} (1 + (-1)^k) \\ &= \sum_{\substack{k=0 \\ k \text{ pair}}}^n \binom{n}{k} \times 2 + \sum_{\substack{k=0 \\ k \text{ impair}}}^n \binom{n}{k} \times 0 = 2 \sum_{\substack{k=0 \\ k \text{ pair}}}^n \binom{n}{k}\end{aligned}$$

Or, pour tout $k \in \llbracket 0 ; n \rrbracket$, k est pair si et seulement si il existe un entier p tel que $k = 2p$, avec $0 \leq 2p \leq n$, c'est-à-dire $0 \leq p \leq \frac{n}{2}$, c'est-à-dire (puisque p est entier) $0 \leq p \leq \left\lfloor \frac{n}{2} \right\rfloor$

Rappelons que pour tout réel y , $\lfloor y \rfloor$ est le plus grand entier inférieur ou égal à y . Dire d'un entier p qu'il est inférieur ou égal à y revient donc à dire qu'il est inférieur ou égal à $\lfloor y \rfloor$.

$$\text{Donc : } \sum_{\substack{k=0 \\ k \text{ pair}}}^n \binom{n}{k} = \sum_{p=0}^{\lfloor n/2 \rfloor} \binom{n}{2p}$$

Malgré toutes mes explications, j'ai conscience que cette égalité semblera la moins évidente pour un élève au sortir de la Terminale. J'ai fait, sans prononcer le mot, un changement d'indice $k = 2p$. Attention ! Je me le suis permis car la somme de gauche ne porte que sur des k pairs, et que j'ai montré au préalable que les k pairs vérifiant $0 \leq k \leq n$ sont exactement les entiers qui s'écrivent $2p$ avec p vérifiant $0 \leq p \leq \left\lfloor \frac{n}{2} \right\rfloor$

$$\text{D'où : } \sum_{k=0}^n \binom{n}{k} (1 + (-1)^k) = 2 \times \sum_{p=0}^{\lfloor n/2 \rfloor} \binom{n}{2p}.$$

$$\text{Puis : } \sum_{p=0}^{\lfloor n/2 \rfloor} \binom{n}{2p} = \frac{1}{2} \sum_{k=0}^n \binom{n}{k} (1 + (-1)^k) = \frac{1}{2} \left[\sum_{k=0}^n \binom{n}{k} + \sum_{k=0}^n \binom{n}{k} (-1)^k \right] = \frac{1}{2} \times (2^n + 0)$$

$$\text{Enfin : } \sum_{p=0}^{\lfloor n/2 \rfloor} \binom{n}{2p} = 2^{n-1}$$

Quelques rappels de calcul matriciel

Ces rappels **ne sont pas exhaustifs** et se concentreront notamment sur les notions nécessaires pour aborder les exercices 17, 18, et 20. Vous pouvez aussi consulter [cette playlist de vidéos courtes](#), série introductive sur le calcul matriciel qui reprend en grande partie les rappels ci-après.

Enfin, si vous estimez ne pas avoir besoin de ces rappels, vous pouvez attaquer directement [les exercices ici](#).

Tailles et positions

Soient n et p deux entiers naturels non nuls. Une matrice réelle de taille $n \times p$ est un tableau de nombres qui comporte n lignes et p colonnes. Ces nombres sont appelés coefficients de la matrice.

Le coefficient situé à la i -ème ligne et j -ième colonne d'une matrice M sera noté $m_{i,j}$, ou parfois $[M]_{i,j}$.

Exemples : $M = \begin{pmatrix} 1 & 0 & 1 \\ 0 & 1 & 0 \\ 0 & 0 & 2 \end{pmatrix}$ $N = \begin{pmatrix} 3 & -1 \\ 1 & 1 \\ 3 & 0 \end{pmatrix}$ $P = \begin{pmatrix} 0 \\ 1 \\ \frac{1}{2} \end{pmatrix}$ $Q = \begin{pmatrix} 8 & -2 & 1 & 1 \end{pmatrix}$

M est une matrice de taille 3×3 . N est une matrice de taille 3×2 .

P est une matrice de taille 2×1 . Q est une matrice de taille 1×4 .

On note aussi : $M \in \mathcal{M}_{3,3}(\mathbb{R})$, $N \in \mathcal{M}_{3,2}(\mathbb{R})$, $P \in \mathcal{M}_{2,1}(\mathbb{R})$ et $Q \in \mathcal{M}_{1,4}(\mathbb{R})$

Par ailleurs, on peut aussi dire que P est une matrice colonne de taille 2, et que Q est une matrice ligne de taille 4.

Quelques coefficients : $m_{3,3} = 2$, $n_{1,2} = -1$, $p_{2,1} = \frac{1}{2}$ et $q_{1,4} = 1$

Lorsqu'une matrice a autant de lignes que de colonnes, on dit qu'elle est carrée (de taille n si elle a n lignes). Plus simplement que $\mathcal{M}_{n,n}(\mathbb{R})$, l'ensemble des matrices carrées de taille n à coefficients réels se note aussi $\mathcal{M}_n(\mathbb{R})$. Dans les exemples ci-dessus, M est une matrice carrée de taille 3. On peut noter : $M \in \mathcal{M}_3(\mathbb{R})$

Egalité de deux matrices

Deux matrices sont égales si et seulement si elles ont la même taille et que leurs coefficients (aux mêmes emplacements) sont deux à deux égaux.

Matrice identité

La matrice identité de taille n est la matrice dont les coefficients diagonaux (diagonale qui va d'en haut à gauche jusqu'en bas à droite) sont égaux à 1, les autres étant égaux à 0. On la note souvent I_n .

Par exemple : $I_2 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$, $I_3 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \dots$

Matrice nulle

La matrice nulle $O_{n,p}$ est la matrice de taille $n \times p$ dont tous les coefficients sont égaux à 0. Par commodité, $O_{n,n}$ est notée O_n .

Par exemple : $O_2 = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}$, $O_{2,3} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \dots$

Somme de matrices

On peut sommer deux matrices A et B de même taille.

Si A et B sont deux matrices de taille $n \times p$, la matrice $A + B$ est la matrice de taille $n \times p$ obtenue en sommant deux à deux les coefficients aux mêmes emplacements.

Exemple : si $M = \begin{pmatrix} 1 & 0 & 1 \\ 0 & 1 & 0 \\ 0 & 0 & 2 \end{pmatrix}$, $N = \begin{pmatrix} 1 & 2 & 0 \\ -1 & 3 & 4 \\ 5 & 0 & 0 \end{pmatrix}$, et $P = \begin{pmatrix} 4 & 7 \\ 1 & 2 \end{pmatrix}$

$$M + N = \begin{pmatrix} 1+1 & 0+2 & 1+0 \\ 0-1 & 1+3 & 0+4 \\ 0+5 & 0+0 & 2+0 \end{pmatrix} = \begin{pmatrix} 2 & 2 & 1 \\ -1 & 4 & 4 \\ 5 & 0 & 2 \end{pmatrix}.$$

Remarquez que l'addition matricielle est commutative : lorsque l'addition est possible, on a toujours : $M + N = N + M$

Ici, on ne peut pas sommer M et P , ou N et P (pas les mêmes tailles)

Par ailleurs, pour toute matrice A de taille $n \times p$, $A + O_{n,p} = A$. On dit que $O_{n,p}$ est l'élément neutre pour l'addition dans $\mathcal{M}_{n,p}(\mathbb{R})$

Produit d'une matrice par un réel

On peut multiplier n'importe quelle matrice A par un réel k . La matrice kA est la matrice de même taille que A , obtenue en multipliant chaque coefficient de A par k .

Exemple : si $M = \begin{pmatrix} 1 & 0 & 1 \\ 0 & 1 & 0 \\ 0 & 0 & 2 \end{pmatrix}$, $-5M = \begin{pmatrix} -5 & 0 & -5 \\ 0 & -5 & 0 \\ 0 & 0 & -10 \end{pmatrix}$

Produit de deux matrices

C'est plus compliqué que la somme... Il ne s'agit pas d'un simple produit terme à terme. Si A est une matrice de taille $n \times p$ et si B est une matrice de taille $p \times q$, on définit la matrice $A \times B$ (notée aussi AB) ainsi : $C = A \times B$ est la matrice de taille $n \times q$ telle que, pour tous entiers i et j vérifiant $1 \leq i \leq n$ et $1 \leq j \leq q$: $c_{i,j} = \sum_{k=1}^p a_{i,k} b_{k,j}$

Autrement dit : $c_{i,j} = a_{i,1}b_{1,j} + a_{i,2}b_{2,j} + \dots + a_{i,p}b_{p,j}$

Pour obtenir le coefficient $c_{i,j}$, on prend la i -ème ligne de A et la j -ième colonne de B . On effectue alors, pour tout entier k entre 1 et p , le produit du k -ième terme de cette i -ème ligne de A avec le k -ième terme de cette j -ième ligne de B , et on somme ces p produits. Un exemple nous permettra d'y voir plus clair :

Exemple : si $A = \begin{pmatrix} 3 & 4 \\ 1 & 0 \\ 2 & 2 \end{pmatrix}$ (de taille 3×2) et $B = \begin{pmatrix} 1 & 8 & 0 & 7 \\ 1 & -1 & 0 & 5 \end{pmatrix}$ (de taille 2×4),

$C = A \times B$ est une matrice de taille 3×4 ,

$$\text{et } C = \begin{pmatrix} 3 \times 1 + 4 \times 1 & 3 \times 8 + 4 \times (-1) & 3 \times 0 + 4 \times 0 & 3 \times 7 + 4 \times 5 \\ 1 \times 1 + 0 \times 1 & 1 \times 8 + 0 \times (-1) & 1 \times 0 + 0 \times 0 & 1 \times 7 + 0 \times 5 \\ 2 \times 1 + 2 \times 1 & 2 \times 8 + 2 \times (-1) & 2 \times 0 + 2 \times 0 & 2 \times 7 + 2 \times 5 \end{pmatrix} = \begin{pmatrix} 7 & 20 & 0 & 41 \\ 1 & 8 & 0 & 7 \\ 4 & 14 & 0 & 24 \end{pmatrix}$$

Handwritten diagram illustrating the calculation of the element $c_{2,4}$ in the product matrix C . It shows matrix A with rows 1, 2, 3 and columns 1, 2. Matrix B has rows 1, 2 and columns 1, 2, 3, 4. The calculation for $c_{2,4}$ is shown as $a_{2,1}b_{1,4} + a_{2,2}b_{2,4} = 1 \times 7 + 0 \times 5 = 7$. The final matrix C is shown with the value 7 in the second row, fourth column.

Disposition plus simple pour le calcul, et détail du calcul de $c_{2,4}$

Attention : le produit matriciel $A \times B$ n'est possible que lorsque le nombre de colonnes de A est égal aux lignes de B .

Cela implique notamment que le produit matriciel n'est pas commutatif : on n'a pas en général $A \times B = B \times A$ (puisque en général, l'un peut être défini sans que l'autre ne le soit).

Dans le cas où A et B sont deux matrices carrées de même taille, les produits $A \times B$ et $B \times A$ sont bien définis, mais pas nécessairement égaux. Lorsque $A \times B = B \times A$, on dit que les matrices A et B commutent.

Exemple : si $A = \begin{pmatrix} 3 & 4 \\ 1 & 0 \end{pmatrix}$ et $B = \begin{pmatrix} 1 & 8 \\ 1 & -1 \end{pmatrix}$:

$$A \times B = \begin{pmatrix} 3 & 4 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} 1 & 8 \\ 1 & -1 \end{pmatrix} = \begin{pmatrix} 3 \times 1 + 4 \times 1 & 3 \times 8 + 4 \times (-1) \\ 1 \times 1 + 0 \times 1 & 1 \times 8 + 0 \times (-1) \end{pmatrix} = \begin{pmatrix} 7 & 20 \\ 1 & 8 \end{pmatrix}$$

$$B \times A = \begin{pmatrix} 1 & 8 \\ 1 & -1 \end{pmatrix} \begin{pmatrix} 3 & 4 \\ 1 & 0 \end{pmatrix} = \begin{pmatrix} 1 \times 3 + 8 \times 1 & 1 \times 4 + 8 \times 0 \\ 1 \times 3 + (-1) \times 1 & 1 \times 4 + (-1) \times 0 \end{pmatrix} = \begin{pmatrix} 11 & 4 \\ 2 & 4 \end{pmatrix}$$

On remarque : $A \times B \neq B \times A$

Pour toute matrice carrée A de taille n , on a $A \times I_n = I_n \times A = A$

On dit que I_n est l'élément neutre pour la multiplication dans $\mathcal{M}_n(\mathbb{R})$.

Par ailleurs, pour toute matrice A de taille $n \times p$, $A \times O_{p,q} = O_{n,q}$ et $O_{m,n} A = O_{m,p}$

Cas particulier du produit d'une matrice carrée par une matrice colonne

$$\text{Si } A = \begin{pmatrix} 3 & 4 & 1 \\ 1 & 0 & 2 \\ -1 & -2 & 6 \end{pmatrix} \text{ et } X = \begin{pmatrix} -3 \\ 5 \\ 0 \end{pmatrix}, \quad AX = \begin{pmatrix} 3 & 4 & 1 \\ 1 & 0 & 2 \\ -1 & -2 & 6 \end{pmatrix} \begin{pmatrix} -3 \\ 5 \\ 0 \end{pmatrix} = \begin{pmatrix} 3 \times (-3) + 4 \times 5 + 1 \times 0 \\ 1 \times (-3) + 0 \times 5 + 2 \times 0 \\ -1 \times (-3) - 2 \times 5 + 6 \times 0 \end{pmatrix}$$

$$\text{Donc } AX = \begin{pmatrix} 11 \\ -3 \\ -7 \end{pmatrix}$$

Quelques propriétés calculatoires

Pour tous réels a et b , pour toutes matrices M et N de même taille :

$$(a + b)M = aM + bM \quad a(bM) = (ab)M = abM \quad a(M + N) = aM + aN$$

Pour toutes matrices A, B et C tels que les produits matriciels soient possibles :

- $A \times (B \times C) = (A \times B) \times C$ (qu'on peut donc écrire sans ambiguïté $A \times B \times C$ ou ABC)
C'est ce qu'on appelle l'associativité du produit matriciel.
- $A \times (B + C) = A \times B + A \times C$ et $(A + B) \times C = A \times C + B \times C$
C'est ce qu'on appelle la distributivité du produit sur l'addition.

L'intégrité tombe à l'eau

Notez qu'un produit de deux matrices peut être nul sans qu'aucune des deux matrices ne soit nulle. Voyez plutôt :

$$\text{si } A = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \text{ et } B = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix} : A \times B = \begin{pmatrix} 1 \times 0 + 0 \times 0 & 1 \times 0 + 0 \times 1 \\ 0 \times 0 + 0 \times 0 & 0 \times 0 + 0 \times 1 \end{pmatrix} = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix} = O_2$$

Et pourtant, ni A ni B n'est nulle...

La jolie propriété « un produit de facteurs est nul si et seulement si l'un au moins des facteurs est nul » que vous avez sur $\mathbb{R}$ (et aussi sur $\mathbb{C}$ si vous avez vu les nombres complexes), qui vous accompagne depuis le collège, et qui fait de $\mathbb{R}$ ce que l'on appelle - frimons un peu - un anneau intègre, tombe à l'eau pour les matrices...

Par contre, il est tout à fait exact d'affirmer que le produit λM entre un réel λ et une matrice M de taille $n \times p$ est nul si et seulement si $\lambda = 0$ ou $M = O_{n,p}$

Puissances d'une matrice carrée

Soit $n \in \mathbb{N}^*$. Soit A une matrice carrée de taille n .

Soit $k \in \mathbb{N}^*$. La puissance k -ième de A est $A^k = A \times A \times \dots \times A$ (où A apparaît k fois).

Et, par convention : $A^0 = I_n$. (de même que pour tout réel a , $a^0 = 1$)

Nous avons donc : $\forall k \in \mathbb{N}$, $A^{k+1} = A \times A^k = A^k \times A$.

Remarquez en particulier que pour tout $k \in \mathbb{N}$, $I_n^k = I_n$

Inversibilité et inverse d'une matrice carrée

Soit $n \in \mathbb{N}^*$, et soit M une matrice carrée de taille n .

M est dite inversible lorsqu'il existe une matrice N carrée de taille n telle que :

$$M \times N = N \times M = I_n.$$

N est alors appelée l'inverse de M , et on note $N = M^{-1}$

(En fait, l'une des deux égalités $M \times N = I_n$ ou $N \times M = I_n$ suffit et implique l'autre)

Exemple : si $A = \begin{pmatrix} 3 & -1 \\ 1 & 0 \end{pmatrix}$ et $B = \begin{pmatrix} 0 & 1 \\ -1 & 3 \end{pmatrix}$:

$$A \times B = \begin{pmatrix} 3 & -1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ -1 & 3 \end{pmatrix} = \begin{pmatrix} 3 \times 0 - 1 \times (-1) & 3 \times 1 - 1 \times 3 \\ 1 \times 0 + 0 \times (-1) & 1 \times 1 + 0 \times 3 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = I_2$$

A est donc inversible et $A^{-1} = B$ (de même, B est inversible et $B^{-1} = A$)

Les matrices carrées de taille n ne sont pas toutes inversibles.

Lorsqu'une matrice carrée A est inversible, son inverse A^{-1} est unique.

Lorsqu'une matrice carrée A est inversible d'inverse A^{-1} , pour tout $k \in \mathbb{N}$, A^{-k} est la matrice définie ainsi : $A^{-k} = (A^{-1})^k$

Lorsque A est inversible, pour tout $k \in \mathbb{N}$, A^k l'est aussi, et $(A^k)^{-1} = (A^{-1})^k = A^{-k}$

Transposée d'une matrice

Si A est une matrice de taille $n \times p$, la transposée de A , que l'on note A^T (ou tA), est la matrice de taille $p \times n$ telle que : $\forall i \in \llbracket 1; p \rrbracket, \forall j \in \llbracket 1; n \rrbracket, [A^T]_{i,j} = [A]_{j,i} = a_{j,i}$.

Autrement dit, le coefficient situé à la i -ème ligne et j -ième colonne de la matrice A^T est égal au coefficient situé à la j -ème ligne et i -ième colonne de la matrice A .

Exemple : si $M = \begin{pmatrix} 2 & 3 & 1 \\ -1 & 0 & 4 \end{pmatrix}$, $M^T = \begin{pmatrix} 2 & -1 \\ 3 & 0 \\ 1 & 4 \end{pmatrix}$. Et si $A = \begin{pmatrix} 1 & 2 & 0 \\ -1 & 3 & 4 \\ 5 & 0 & 7 \end{pmatrix}$, $A^T = \begin{pmatrix} 1 & -1 & 5 \\ 2 & 3 & 0 \\ 0 & 4 & 7 \end{pmatrix}$

Remarquez que dans le cas de la matrice A carrée, A^T s'obtient à partir de A en effectuant, sur ses coefficients, une sorte de symétrie axiale d'axe la diagonale de A .

Une matrice carrée A est dite symétrique lorsqu'elle est égale à sa transposée, autrement dit lorsque $A^T = A$.

Une matrice carrée A est dite antisymétrique lorsque $A^T = -A$

Exercice 17

*Je vous suis à la trace et vos pas me lessivent.
Mais vous ne fuirez pas, puissances successives !*

Énoncé : (temps conseillé : 1 h) (***)

d'après X-ESPCI 2008 PC

On appelle trace d'une matrice carrée A la somme de ses coefficients diagonaux (diagonale qui va d'en haut à gauche jusqu'en bas à droite) et on la note $\text{Tr}(A)$.

Autrement dit : $\forall n \in \mathbb{N}^*, \forall A \in \mathcal{M}_n(\mathbb{R}), \text{Tr}(A) = \sum_{i=1}^n a_{i,i}$

On considère la matrice $M = \begin{pmatrix} 0 & 0 & -1 \\ 0 & 1 & 1 \\ 1 & 1 & -1 \end{pmatrix}$, et on note I la matrice identité de $\mathcal{M}_3(\mathbb{R})$.

1) Calculer M^3 et l'exprimer comme combinaison linéaire de M et I .

2) Pour tout $n \in \mathbb{N}$, exprimer M^{n+3} en fonction de M^{n+1} et M^n .

3) Pour tout entier naturel n , on pose $u_n = \text{Tr}(M^n)$ et $v_n = \cos(\pi u_n)$

On admet que pour tout entier naturel n , u_n est un entier.

a) Pour $0 \leq n \leq 10$, calculer u_n et v_n

b) Montrer que la suite (v_n) est périodique, et en déterminer une période.

c) Montrer que la suite (w_n) définie par $w_n = \sum_{k=0}^n v_k$ n'est pas bornée.

Remarques sur l'énoncé :

La question 1) demande d'exprimer M^3 comme combinaison linéaire de I et M , c'est-à-dire d'arriver à une écriture $M^3 = \alpha M + \beta I$, avec α et β deux réels.

Correction de l'exercice 17 :

$$1) M^2 = \begin{pmatrix} 0 & 0 & -1 \\ 0 & 1 & 1 \\ 1 & 1 & -1 \end{pmatrix} \begin{pmatrix} 0 & 0 & -1 \\ 0 & 1 & 1 \\ 1 & 1 & -1 \end{pmatrix} = \begin{pmatrix} -1 & -1 & 1 \\ 1 & 2 & 0 \\ -1 & 0 & 1 \end{pmatrix}$$

Pour le détail du calcul :

$$M^2 = \begin{pmatrix} 0 \times 0 + 0 \times 0 + (-1) \times 1 & 0 \times 0 + 0 \times 1 + (-1) \times 1 & 0 \times (-1) + 0 \times 1 + (-1) \times (-1) \\ 0 \times 0 + 1 \times 0 + 1 \times 1 & 0 \times 0 + 1 \times 1 + 1 \times 1 & 0 \times (-1) + 1 \times 1 + 1 \times (-1) \\ 1 \times 0 + 1 \times 0 + (-1) \times 1 & 1 \times 0 + 1 \times 1 + (-1) \times 1 & 1 \times (-1) + 1 \times 1 + (-1) \times (-1) \end{pmatrix}$$

Vous n'êtes évidemment pas tenus de faire figurer un tel calcul sur vos copies. Je le fais (cette fois) juste pour que vous puissiez vérifier votre calcul matriciel. De rien.

$$\text{Puis : } M^3 = M^2 \times M = \begin{pmatrix} -1 & -1 & 1 \\ 1 & 2 & 0 \\ -1 & 0 & 1 \end{pmatrix} \times \begin{pmatrix} 0 & 0 & -1 \\ 0 & 1 & 1 \\ 1 & 1 & -1 \end{pmatrix} = \begin{pmatrix} 1 & 0 & -1 \\ 0 & 2 & 1 \\ 1 & 1 & 0 \end{pmatrix}$$

$$\text{On remarque que } M^3 = \begin{pmatrix} 0 & 0 & -1 \\ 0 & 1 & 1 \\ 1 & 1 & -1 \end{pmatrix} + \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}. \text{ Autrement dit : } \boxed{M^3 = M + I}$$

2) Nous savons : $M^3 = M + I$. Donc, pour tout entier naturel n : $M^3 \times M^n = (M + I) \times M^n$.
D'où (par distributivité) : $\forall n \in \mathbb{N}, M^{n+3} = M \times M^n + I \times M^n$.

$$\text{Enfin : } \boxed{\forall n \in \mathbb{N}, M^{n+3} = M^{n+1} + M^n}$$

3)a) Qu'est-ce que c'est que ce délire ? Vingt-deux calculs ?? C'est l'énoncé d'origine qui le demande, promis... Mais ce ne sera peut-être pas aussi fastidieux que ça en a l'air. Exploitions déjà l'égalité obtenue précédemment.

Pour tout entier naturel n , $u_{n+3} = \text{Tr}(M^{n+3}) = \text{Tr}(M^{n+1} + M^n)$

Au vu de la définition de la trace, il semble relativement intuitif que la trace d'une somme de deux matrices est égale à la somme de leurs traces. Mais n'ayant pas eu cette information, il nous faut la démontrer si nous voulons nous en servir.

Pour tout $p \in \mathbb{N}^*$, pour toutes matrices A et B de $\mathcal{M}_p(\mathbb{R})$, notons $C = A + B$. Nous avons alors : $\forall i, j \in \llbracket 1 ; p \rrbracket, c_{i,j} = a_{i,j} + b_{i,j}$. Puis : $\text{Tr}(C) = \sum_{i=1}^p c_{i,i} = \sum_{i=1}^p a_{i,i} + b_{i,i} = \sum_{i=1}^p a_{i,i} + \sum_{i=1}^p b_{i,i}$

Donc : $\text{Tr}(C) = \text{Tr}(A) + \text{Tr}(B)$. Autrement dit : $\text{Tr}(A + B) = \text{Tr}(A) + \text{Tr}(B)$.

J'ai utilisé la lettre p pour la taille des matrices carrées A et B , et surtout pas n , qui désigne déjà les exposants des puissances de M mises en jeu précédemment. J'aurais même pu me contenter de rester sur $\mathcal{M}_3(\mathbb{R})$ plutôt qu'un $\mathcal{M}_p(\mathbb{R})$ plus général, mais c'eût été dommage de me priver de la portée générale du résultat.

L'an prochain, vous verrez et pourrez utiliser à votre guise la linéarité de la trace :

$$\forall A, B \in \mathcal{M}_p(\mathbb{R}), \forall \alpha, \beta \in \mathbb{R}, \text{Tr}(\alpha A + \beta B) = \alpha \text{Tr}(A) + \beta \text{Tr}(B)$$

Revenons aux matrices spécifiques à notre exercice.

Nous savons donc : $\forall n \in \mathbb{N}, u_{n+3} = \text{Tr}(M^{n+1}) + \text{Tr}(M^n)$.

D'où : $\forall n \in \mathbb{N}, u_{n+3} = u_{n+1} + u_n$

Voilà qui va pas mal nous faciliter le calcul des termes de u_0 à u_{10} ...

$$u_0 = \text{Tr}(M^0) = \text{Tr}(I) = 1 + 1 + 1 = 3 \quad u_1 = \text{Tr}(M^1) = \text{Tr}(M) = 0 + 1 - 1 = 0$$

$$u_2 = \text{Tr}(M^2) = -1 + 2 + 1 = 2 \quad u_3 = u_1 + u_0 = 0 + 3 = 3$$

$$u_4 = u_2 + u_1 = 2 + 0 = 2$$

...

De plus : $\forall n \in \mathbb{N}, v_n = \cos(\pi u_n)$

Rappelons que pour tout entier naturel k , $\cos(k\pi) = 1$ si k est pair, et 0 si k est impair.

Autrement dit : $\cos(k\pi) = (-1)^k$.

Dès lors, $v_0, v_1, \dots, v_{10}$ s'obtiennent simplement à partir de $u_0, u_1, \dots, u_{10}$. Voici, sous forme de tableau, les valeurs demandées :

n	0	1	2	3	4	5	6	7	8	9	10
u_n	3	0	2	3	2	5	5	7	10	12	17
v_n	-1	1	1	-1	1	-1	-1	-1	1	1	-1

3)b) Au vu des valeurs obtenues en 3)a), (v_n) semble périodique de période 7 (on voit : $v_7 = v_0, v_8 = v_1, v_9 = v_2, v_{10} = v_3$)

Pour l'instant, cette périodicité n'est qu'une conjecture, à démontrer proprement.

En revanche, il est déjà certain que pour tout entier T tel que $1 \leq T \leq 6$, (v_n) ne peut être périodique de période T .

En effet, (v_n) est périodique de période T si et seulement si : $\forall n \in \mathbb{N}, v_{n+T} = v_n$

Et $v_1 \neq v_0$ (ce qui écarte une période 1, qui correspondrait d'ailleurs à une suite constante),
 $v_2 \neq v_0$ (ce qui écarte une période 2), $v_5 \neq v_2$, $v_4 \neq v_0$, $v_6 \neq v_1$, $v_7 \neq v_1$

Montrons que (v_n) est périodique de période 7. Pour tout $n \in \mathbb{N}$, $v_{n+7} = \cos(\pi u_{n+7})$

Or : $u_{n+7} = u_{n+4+3} = u_{n+4+1} + u_{n+4} = u_{n+5} + u_{n+4}$

D'une part : $u_{n+5} = u_{n+3} + u_{n+2}$ et d'autre part, $u_{n+4} = u_{n+2} + u_{n+1}$

Donc : $u_{n+7} = u_{n+3} + u_{n+2} + u_{n+2} + u_{n+1} = u_{n+3} + 2u_{n+2} + u_{n+1}$

Allez, un dernier coup de transformation...

Et $u_{n+3} = u_{n+1} + u_n$. D'où : $u_{n+7} = u_{n+1} + u_n + 2u_{n+2} + u_{n+1}$

Autrement dit : $\forall n \in \mathbb{N}$, $u_{n+7} = 2(u_{n+2} + u_{n+1}) + u_n$

u_{n+2} et u_{n+1} étant des entiers (admis par l'énoncé*), $u_{n+2} + u_{n+1}$ aussi.

Donc $2(u_{n+2} + u_{n+1})$ est un entier pair.

Il s'ensuit que u_{n+7} est un entier de même parité que u_n .

D'où : $\cos(\pi u_{n+7}) = \cos(\pi u_n)$. Autrement dit : $v_{n+7} = v_n$ (et ce, pour tout entier naturel n)

En conclusion, la suite (v_n) est périodique de période 7.

**Pourquoi l'énoncé nous le fait-il explicitement admettre ? N'était-ce pas évident en soi, au vu de la relation de récurrence vérifiée par la suite (u_n) , et au vu de ses premiers termes ? C'est assez trivial en effet : comme les trois premiers termes de (u_n) sont entiers et que la relation de récurrence vérifiée par cette suite nous assure du fait que, pour tout $n \in \mathbb{N}$, si u_n et u_{n+1} sont entiers, alors u_{n+3} aussi, nous avons la certitude que tous les termes de cette suite sont entiers. La justification rigoureuse qui se cache derrière, c'est une récurrence. Immédiate, certes, mais tout de même une récurrence d'ordre 3, nouveauté pour vous.*

3)c) Voyons déjà ce que donne la somme des termes de (v_n) sur une période...

$$w_6 = \sum_{k=0}^6 v_k = v_0 + v_1 + \dots + v_6 = -1$$

Ah, et donc si l'on sommait sur plusieurs périodes... La somme sur la période suivante serait $v_7 + v_8 + \dots + v_{13} = -1$ (par périodicité), celle d'après : $v_{14} + v_{15} + \dots + v_{20} = -1$...

(v_n) étant périodique de période 7, nous savons : $\forall N \in \mathbb{N}$, $v_{0+7i} + v_{1+7i} + \dots + v_{6+7i} = -1$

Autrement dit : $\forall i \in \mathbb{N}$, $\sum_{k=0}^6 v_{7i+k} = -1$

En sommant sur plusieurs périodes successives, nous obtenons :

$$\forall N \in \mathbb{N}, \sum_{k=0}^6 v_{7 \times 0 + k} + \sum_{k=0}^6 v_{7 \times 1 + k} + \sum_{k=0}^6 v_{7 \times 2 + k} + \dots + \sum_{k=0}^6 v_{7 \times N + k} = -1 + (-1) + (-1) + \dots + (-1)$$

(-1) étant répété autant de fois qu'il y a d'entiers naturels entre 0 et N compris, c'est-à-dire $N + 1$. Cette somme s'écrit, de manière explicite :

$$(v_0 + v_1 + \dots + v_6) + (v_7 + v_8 + \dots + v_{13}) + (v_{14} + v_{15} + \dots + v_{20}) + \dots + (v_{7N} + v_{7N+1} + \dots + v_{7N+6})$$

C'est en fait : $\sum_{k=0}^{7N+6} v_k$, c'est-à-dire w_{7N+6} .

Nous avons donc établi : $\forall N \in \mathbb{N}, w_{7N+6} = -(N + 1)$ *C'est difficilement minoré, ça...*

Par l'absurde : si (w_n) était bornée, elle serait en particulier minorée. Il existerait donc une constante réelle m telle que : $\forall n \in \mathbb{N}, m \leq w_n$. En particulier : $\forall N \in \mathbb{N}, m \leq w_{7N+6}$. C'est-à-dire : $\forall N \in \mathbb{N}, m \leq -(N + 1)$.

Or : $\lim_{N \rightarrow +\infty} -(N + 1) = -\infty$. Donc, par théorème de comparaison : $\lim_{N \rightarrow +\infty} m = -\infty$

C'est absurde car m est une constante.

Nous pouvons donc en déduire que (w_n) n'est pas bornée.

Les deux étoiles (difficulté moyenne) que j'ai attribuées à cet exercice ne tiennent pas compte de la dernière question, particulièrement difficile. J'ai hésité à mettre cette dernière question, que le Terminale moyen devrait trouver assez inaccessible en termes de raisonnement (même s'il peut avoir l'idée générale de sommer les valeurs de (v_n) sur différentes périodes).

Pourquoi l'ai-je tout de même gardée ? Parce qu'elle introduit un principe intéressant : le principe de sommation par paquets. Principe que j'utilise sans le dire, même si je mets bien, bout à bout, et de manière explicite, des paquets de termes consécutifs de (v_n)

De manière plus formelle, j'aurais pu écrire :

$$\forall i \in \mathbb{N}, \sum_{k=0}^6 v_{7i+k} = -1. \text{ Puis : } \forall N \in \mathbb{N}, \sum_{i=0}^N \sum_{k=0}^6 v_{7i+k} = \sum_{i=0}^N -1 = -(N + 1)$$

$$\text{Et (c'est là que la sommation par paquets intervient) : } \sum_{i=0}^N \sum_{k=0}^6 v_{7i+k} = \sum_{k=0}^{7N+6} v_k$$

Pardon d'avoir voulu vous épargner la frayeur de voir apparaître des sommes doubles...

Exercice 18

*Les muscles bien gainés sous son affreux survêt',
entre trois points donnés, ce pion fait la navette.*

Énoncé : (temps conseillé : 1 heure) (**)

d'après CCINP 2019 PC

On considère trois points distincts du plan nommés A , B et C . Nous allons étudier le déplacement aléatoire d'un pion se déplaçant sur ces trois points. A l'étape $n = 0$, on suppose que le pion se trouve sur le point A . Ensuite, le mouvement aléatoire du pion respecte les deux règles suivantes :

1. le mouvement du pion de l'étape n à l'étape $n + 1$ ne dépend que de la position du pion à l'étape n ; plus précisément, il ne dépend pas des positions occupées aux autres étapes précédentes.
2. pour passer de l'étape n à l'étape $n + 1$, on suppose que le pion a une chance sur deux de rester sur place, sinon il se déplace de manière équiprobable vers l'un des deux autres points.

Pour tout $n \in \mathbb{N}$, on note A_n l'évènement "le pion se trouve en A à l'étape n ", B_n l'évènement "le pion se trouve en B à l'étape n " et C_n l'évènement "le pion se trouve en C à l'étape n ". On note également :

$$\forall n \in \mathbb{N}, p_n = P(A_n), \quad q_n = P(B_n), \quad r_n = P(C_n), \quad V_n = \begin{pmatrix} p_n \\ q_n \\ r_n \end{pmatrix},$$

et on considère la matrice : $M = \frac{1}{4} \begin{pmatrix} 2 & 1 & 1 \\ 1 & 2 & 1 \\ 1 & 1 & 2 \end{pmatrix} \in \mathcal{M}_3(\mathbb{R})$.

Dans l'exercice, on pourra utiliser **sans le démontrer** le résultat suivant :

$$\forall n \in \mathbb{N}, M^n = \frac{1}{3 \times 4^n} \begin{pmatrix} 4^n + 2 & 4^n - 1 & 4^n - 1 \\ 4^n - 1 & 4^n + 2 & 4^n - 1 \\ 4^n - 1 & 4^n - 1 & 4^n + 2 \end{pmatrix}.$$

On rappelle que si E et F sont deux évènements avec $P(F) > 0$, on définit la probabilité conditionnelle de E sachant F (notée $P_F(E)$) par :

$$P_F(E) = \frac{P(E \cap F)}{P(F)}.$$

-
- 1) Calculer les nombres p_n , q_n et r_n pour $n = 0$ et $n = 1$.
 - 2) Démontrer que pour tout $n \in \mathbb{N}$, on a la relation $V_{n+1} = MV_n$.
 - 3) En déduire que $V_n = M^n V_0$, puis une expression de p_n , q_n et r_n pour tout $n \in \mathbb{N}$.
 - 4) Déterminer les limites respectives des suites $(p_n)_{n \in \mathbb{N}}$, $(q_n)_{n \in \mathbb{N}}$ et $(r_n)_{n \in \mathbb{N}}$. Interpréter le résultat.

Pour $n \in \mathbb{N}^*$, on note a_n le nombre moyen de passages du pion en A entre l'étape 1 et l'étape n et on définit la variable aléatoire :

$$X_n = \begin{cases} 1 & \text{si } A_n \text{ est réalisé} \\ 0 & \text{si } \overline{A_n} \text{ est réalisé} \end{cases}.$$

- 5) Interpréter la variable aléatoire $X_1 + \dots + X_n$ et le nombre $E(X_1 + \dots + X_n)$.
- 6) Déterminer, pour tout $n \in \mathbb{N}^*$, l'expression de a_n en fonction de n .

Remarques sur l'énoncé :

J'ai assez peu remanié l'énoncé originel. Un bon petit morceau de sujet faisable juste avec les connaissances de Terminale, quelle aubaine pour moi qui arrive à l'exercice 18 et commence à fatiguer... Tiens, je me suis même permis de supprimer une question intermédiaire à un moment donné, en me disant que ça vous rendait la tâche trop simple !

Correction de l'exercice 18 :

1) On sait que le pion se trouve sur le point A à l'étape 0.

Donc : $p_0 = P(A_0) = 1$, $q_0 = P(B_0) = 0$, et $r_0 = P(C_0) = 0$

Etant données les règles de déplacement du pion, il a une chance sur deux de rester sur le point A à l'étape 2, et une chance sur deux de rejoindre, de manière équiprobable, le point B ou le point C . D'où : $p_1 = P(A_1) = \frac{1}{2}$, $q_1 = P(B_1) = \frac{1}{4}$, et $r_1 = P(C_1) = \frac{1}{4}$

2) On cherche à exprimer V_{n+1} en fonction de V_n . Autrement dit, à exprimer p_{n+1} , q_{n+1} et r_{n+1} en fonction de p_n , q_n et r_n .

Soit $n \in \mathbb{N}$. Les événements A_n , B_n et C_n forment un système complet d'événements.

En effet, à l'étape n , le pion se situe sur un et un seul des points A , B , ou C .

D'après la formule des probabilités totales :

$$\begin{aligned} p_{n+1} &= P(A_{n+1}) = P(A_n \cap A_{n+1}) + P(B_n \cap A_{n+1}) + P(C_n \cap A_{n+1}) \\ &= P(A_n) \times P_{A_n}(A_{n+1}) + P(B_n) \times P_{B_n}(A_{n+1}) + P(C_n) \times P_{C_n}(A_{n+1}) \\ &= p_n \times \frac{1}{2} + q_n \times \frac{1}{4} + r_n \times \frac{1}{4} \end{aligned}$$

En effet, $P_{A_n}(A_{n+1})$ est la probabilité que le pion reste sur A à l'étape $n+1$ sachant qu'il y était déjà à l'étape n : cette probabilité vaut $\frac{1}{2}$ d'après l'énoncé. Et, toujours d'après l'énoncé, $P_{B_n}(A_{n+1}) = P_{C_n}(A_{n+1}) = \frac{1}{4}$

$$\text{Donc : } p_{n+1} = \frac{1}{2}p_n + \frac{1}{4}q_n + \frac{1}{4}r_n.$$

$$\text{De même : } q_{n+1} = \frac{1}{4}p_n + \frac{1}{2}q_n + \frac{1}{4}r_n, \text{ et } r_{n+1} = \frac{1}{4}p_n + \frac{1}{4}q_n + \frac{1}{2}r_n$$

Oui, je peux me permettre un petit « de même » sans avoir l'impression d'exagérer. Ce serait la même formule des probabilités totales, les mêmes calculs, à ceci près que le $\frac{1}{2}$ change de place

$$\text{Nous avons donc : } \forall n \in \mathbb{N}, \begin{cases} p_{n+1} &= \frac{1}{2}p_n + \frac{1}{4}q_n + \frac{1}{4}r_n \\ q_{n+1} &= \frac{1}{4}p_n + \frac{1}{2}q_n + \frac{1}{4}r_n \\ r_{n+1} &= \frac{1}{4}p_n + \frac{1}{4}q_n + \frac{1}{2}r_n \end{cases} \quad \text{Donc } V_{n+1} = \begin{pmatrix} \frac{1}{2}p_n + \frac{1}{4}q_n + \frac{1}{4}r_n \\ \frac{1}{4}p_n + \frac{1}{2}q_n + \frac{1}{4}r_n \\ \frac{1}{4}p_n + \frac{1}{4}q_n + \frac{1}{2}r_n \end{pmatrix}$$

$$\text{D'autre part : } MV_n = \frac{1}{4} \begin{pmatrix} 2 & 1 & 1 \\ 1 & 2 & 1 \\ 1 & 1 & 2 \end{pmatrix} \begin{pmatrix} p_n \\ q_n \\ r_n \end{pmatrix} = \frac{1}{4} \times \begin{pmatrix} 2p_n + q_n + r_n \\ p_n + 2q_n + r_n \\ p_n + q_n + 2r_n \end{pmatrix} = \begin{pmatrix} \frac{1}{2}p_n + \frac{1}{4}q_n + \frac{1}{4}r_n \\ \frac{1}{4}p_n + \frac{1}{2}q_n + \frac{1}{4}r_n \\ \frac{1}{4}p_n + \frac{1}{4}q_n + \frac{1}{2}r_n \end{pmatrix}$$

Nous avons bien établi : $\boxed{\forall n \in \mathbb{N}, V_{n+1} = MV_n}$

3) Pour passer d'un terme au suivant de la suite (V_n) , on multiplie chaque fois à gauche par M ... (V_n) est donc en quelque sorte géométrique de « raison » M . On peut comprendre assez vite d'où vient le M^n (simplement, $M \times M \times M \dots$), mais il faut le prouver rigoureusement. Une simple récurrence !

Soit, pour tout $n \in \mathbb{N}$, la propriété P_n : « $V_n = M^n V_0$ »

Montrons par récurrence que pour tout entier naturel n , P_n est vraie.

Initialisation : $M^0 V_0 = I_3 V_0 = V_0$, donc $\underline{P_0}$ est vraie.

Hérédité : Supposons que pour un certain $n \in \mathbb{N}$, P_n soit vraie, et montrons que P_{n+1} aussi est vraie.

Supposons donc que $V_n = M^n V_0$

Nous savons : $V_{n+1} = MV_n$, donc $V_{n+1} = M \times M^n V_0 = M^{n+1} V_0$ $\underline{P_{n+1}}$ est donc vraie.

Encore une fois, pas l'hérédité la plus dure...

Conclusion : Le principe de raisonnement par récurrence nous permet de conclure que pour tout entier naturel n , P_n est vraie. Autrement dit : $\boxed{\forall n \in \mathbb{N}, V_n = M^n V_0}$

$$\text{D'où : } \forall n \in \mathbb{N}, \begin{pmatrix} p_n \\ q_n \\ r_n \end{pmatrix} = \frac{1}{3 \times 4^n} \begin{pmatrix} 4^n + 2 & 4^n - 1 & 4^n - 1 \\ 4^n - 1 & 4^n + 2 & 4^n - 1 \\ 4^n - 1 & 4^n - 1 & 4^n + 2 \end{pmatrix} \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} = \frac{1}{3 \times 4^n} \begin{pmatrix} 4^n + 2 \\ 4^n - 1 \\ 4^n - 1 \end{pmatrix}$$

$$\text{Enfin : } \boxed{\forall n \in \mathbb{N}, p_n = \frac{4^n + 2}{3 \times 4^n}, q_n = \frac{4^n - 1}{3 \times 4^n} \text{ et } r_n = \frac{4^n - 1}{3 \times 4^n}}$$

$$4) \text{ Pour tout entier naturel } n, p_n = \frac{4^n}{3 \times 4^n} + \frac{2}{3 \times 4^n} = \frac{1}{3} + \frac{2}{3 \times 4^n}$$

On a tendance à l'oublier, mais casser une somme au numérateur permet parfois de lever des indéterminations en pacotille...

Or, $4 > 1$, donc : $\lim_{n \rightarrow +\infty} 4^n = +\infty$. Puis, par produit, quotient et somme de limites :

$$\lim_{n \rightarrow +\infty} p_n = \frac{1}{3}. \text{ De même : } \forall n \in \mathbb{N}, q_n = r_n = \frac{1}{3} - \frac{1}{3 \times 4^n}. \text{ D'où : } \lim_{n \rightarrow +\infty} q_n = \frac{1}{3} \text{ et } \lim_{n \rightarrow +\infty} r_n = \frac{1}{3}$$

À long terme, les probabilités pour le pion de se situer en A, B ou C deviennent sensiblement les mêmes.

5) $X_1 + \dots + X_n$ renvoie le nombre de passages du pion sur le point A entre l'étape 1 et l'étape n (puisque, dans cette somme, chaque X_k contribue à hauteur de 1 ou 0, selon que A_k est réalisé ou pas, c-à-d selon que le pion passe ou pas sur A à l'étape k).

Dès lors, $E(X_1 + \dots + X_n)$ est le nombre moyen du passage du pion sur le point A entre l'étape 1 et l'étape n . Autrement dit : $E(X_1 + \dots + X_n) = a_n$

6) Pour tout $n \in \mathbb{N}^*$, $a_n = E(X_1 + \dots + X_n) = E(X_1) + \dots + E(X_n) = \sum_{k=1}^n E(X_k)$ par linéarité de l'espérance. Dès lors, il suffit de calculer $E(X_k)$ pour tout $k \in \llbracket 1 ; n \rrbracket$

Pour tout $k \in \llbracket 1 ; n \rrbracket$, X_k suit une loi de Bernoulli de paramètre $P(A_k) = p_k = \frac{1}{3} + \frac{2}{3 \times 4^k}$

Rappelons que l'espérance d'une variable aléatoire suivant une loi de Bernoulli est égale à son paramètre... Autrement dit : si $X \sim \mathcal{B}(p)$, alors $E(X) = p$

$$\text{Donc : } \forall k \in \llbracket 1 ; n \rrbracket, E(X_k) = \frac{1}{3} + \frac{2}{3 \times 4^k}.$$

$$\text{Puis : } a_n = \sum_{k=1}^n \left(\frac{1}{3} + \frac{2}{3 \times 4^k} \right) = \sum_{k=1}^n \frac{1}{3} + \sum_{k=1}^n \frac{2}{3 \times 4^k}$$

$$\text{D'une part : } \sum_{k=1}^n \frac{1}{3} = n \times \frac{1}{3} = \frac{n}{3} \quad (\text{on somme } n \text{ fois la constante } \frac{1}{3})$$

$$\text{D'autre part : } \sum_{k=1}^n \frac{2}{3 \times 4^k} = \frac{2}{3} \times \sum_{k=1}^n \frac{1}{4^k} = \frac{2}{3} \times \sum_{k=1}^n \left(\frac{1}{4} \right)^k$$

Un petit rappel sur les manipulations de sommes ? Par ici les poissons rouges!

Et $\sum_{k=1}^n \left(\frac{1}{4} \right)^k$ est la somme des n premiers termes d'une suite géométrique de raison

$$\frac{1}{4} \neq 1 \text{ et de premier terme } \left(\frac{1}{4}\right)^1 = \frac{1}{4}. \text{ Donc : } \sum_{k=1}^n \left(\frac{1}{4}\right)^k = \frac{1}{4} \times \frac{1 - \left(\frac{1}{4}\right)^n}{1 - \frac{1}{4}} = \frac{1}{4} \times \left(1 - \left(\frac{1}{4}\right)^n\right) \times \frac{4}{3} \\ = \frac{1}{3} \times \left(1 - \left(\frac{1}{4}\right)^n\right). \text{ Puis : } \frac{2}{3} \times \sum_{k=1}^n \left(\frac{1}{4}\right)^k = \frac{2}{3} \times \frac{1}{3} \times \left(1 - \left(\frac{1}{4}\right)^n\right) = \frac{2}{9} \times \left(1 - \left(\frac{1}{4}\right)^n\right)$$

Il est bon de connaître en « français » la formule donnant la somme de termes consécutifs d'une suite géométrique, pour éviter de mettre par exemple, du $n + 1$ mécaniquement en exposant, sans comprendre d'où il vient...

En « français », la somme de termes consécutifs d'une suite géométrique de raison $q \neq 1$ est égale à : 1^{er} terme de la somme $\times \frac{1 - q^{nb \text{ de termes}}}{1 - q}$

Il y a bien n termes entre 1 et n compris. Le $n + 1$ que vous voyez souvent en exposant vient du fait que les sommes que vous considériez allaient souvent de 0 à n .

Enfin : $\forall n \in \mathbb{N}^*, a_n = \frac{n}{3} + \frac{2}{9} \times \left(1 - \frac{1}{4^n}\right)$

La question enlevée du sujet originel (cf remarques sur l'énoncé) précédait cette dernière demandant l'expression de a_n . Elle demandait d'exprimer $E(X_n)$ pour tout $n \in \mathbb{N}^$. Mais c'eût été trop vous prendre par la main ; j'estime qu'à ce stade du document et de votre poussée de croissance estivale, vous êtes assez grand pour abandonner les vélos à quatre roues...*

Exercice 19

*Rédacteur malicieux, je devine et prétends
que nombre d'entre vous perdront ici leur temps.*

Énoncé : (temps conseillé : 45 min) (***) d'après CCINP 2014 PC Maths 1

Pour $Q \in \mathbb{R}[X]$, on dit que Q est stable si toutes ses racines complexes ont une partie réelle strictement négative.

1) Soit $Q(X) = X^3 - 3X^2 + 4X - 2$. Q est-il stable ?

Soient a et b deux réels. On note $P(X) = X^2 + aX + b$ et $\Delta = a^2 - 4b$.

On veut montrer que P est stable si et seulement si $a > 0$ et $b > 0$

On note z_1 et z_2 deux nombres complexes tels que : $P(X) = (X - z_1)(X - z_2)$.

2) Montrer que $a = -(z_1 + z_2)$ et $b = z_1 z_2$.

3) On suppose dans cette question que $\Delta > 0$.

a) Vérifier que si P est stable, alors $a > 0$ et $b > 0$.

b) Montrer réciproquement que si $a > 0$ et $b > 0$, alors P est stable.

4) On suppose dans cette question que $\Delta \leq 0$. Montrer que P est stable si et seulement si $a > 0$ et $b > 0$.

Remarques sur l'énoncé :

Rappelons que $\mathbb{R}[X]$ est l'ensemble des polynômes à coefficients réels et d'indéterminée X .

Vous emploierez souvent l'adjectif « stable » l'an prochain, mais dans un tout autre sens : celui utilisé notamment dans la correction de l'exercice 11. Si $f : E \rightarrow F$ est une application et si $I \subset E$, alors I est stable par f si et seulement si $\forall x \in I, f(x) \in I$

Les coefficients du polynôme P ont des noms relativement inhabituels en termes de placement, ce qui donne une formule de discriminant peut-être déroutante... Il va falloir vous adapter à la situation, et ne pas vous attacher à la lettre au détriment de ce qu'elle représente.

Correction de l'exercice 19 :

1) Quand on cherche les racines d'un polynôme de degré 3 (ou quand on cherche à le factoriser), il est souvent judicieux de se demander s'il n'aurait pas une racine évidente, pour faciliter la factorisation.

$Q(1) = 1 - 3 + 4 - 2 = 0$, donc 1 est racine de Q

Et là, si je voulais vraiment m'embêter à trouver les autres racines de Q , voici ce que je ferais :

Il existe donc a, b et c réels tels que : $Q(X) = (aX^2 + bX + c)(X - 1)$

Q est de degré 3, d'où le degré 2 de $aX^2 + bX + c$ ($a \neq 0$, et même, $a = 1$ en regardant les coefficients dominants)

Autrement dit, il existe a, b et c réels tels que :

$$X^3 - 3X^2 + 4X - 2 = aX^3 - aX^2 + bX^2 - bX + cX - c, \text{ c'est-à-dire :}$$

$$X^3 - 3X^2 + 4X - 2 = aX^3 + (b - a)X^2 + (c - b)X - c$$

Or, deux polynômes sont égaux si et seulement si leurs coefficients sont deux à deux égaux.

$$D'où : \begin{cases} a &= 1 \\ b - a &= -3 \\ c - b &= 4 \\ -c &= -2 \end{cases} \quad \text{C'est-à-dire :} \begin{cases} a &= 1 \\ b &= -2 \\ c &= 2 \end{cases}$$

Donc : $Q(X) = (X^2 - 2X + 2)(X - 1)$. Restent à déterminer les racines de $X^2 - 2X + 2$. Le discriminant de ce polynôme du second degré est $\Delta = (-2)^2 - 4 \times 2 = -4 < 0$. Ce polynôme (à coefficients réels) admet donc deux racines complexes conjuguées $z_1 = \frac{2 - i\sqrt{4}}{2} = 1 - i$ et $z_2 = 1 + i$. Ces racines font partie de l'ensemble des racines de Q - dont l'ensemble des racines est en fait $\{1; 1 - i; 1 + i\}$. Or, $1 - i$ et $1 + i$ ont pour partie réelle 1, qui est positif (donc pas strictement négatif). Q n'est donc pas stable. Ouf!

Trop long pour rien. Si vous avez fait tout ça, vous méritez la punition de ne pas avoir considéré que 1 (oui oui, la toute-première racine trouvée) est bien un nombre complexe, de partie réelle 1 positive..

je me suis embêté à le faire tout de même car l'occasion était trop belle de vous rappeler la procédure habituelle pour déterminer les racines complexes d'un polynôme de degré 3.

1 est une racine complexe de Q , et $\text{Re}(1) = 1 \geq 0$. Q n'est donc pas stable.

2) D'une part, $P(X) = (X - z_1)(X - z_2) = X^2 - (z_1 + z_2)X + z_1z_2$ et d'autre part,

$P(X) = X^2 + aX + b$. Par identification des coefficients, $a = -(z_1 + z_2)$ et $b = z_1 z_2$

Relations coefficients-racines bonjour !

3)a) P est un polynôme du second degré à coefficients réels de discriminant $\Delta = a^2 - 4b$ strictement positif. Ses racines complexes z_1 et z_2 sont donc deux réels distincts. Donc $\operatorname{Re}(z_1) = z_1$ et $\operatorname{Re}(z_2) = z_2$.

Un lien a été fait avec a et b dans la question précédente...

Si P est stable, $\operatorname{Re}(z_1) = z_1 < 0$ et $\operatorname{Re}(z_2) = z_2 < 0$. D'après 2), d'une part, $b = z_1 z_2 > 0$. Et d'autre part, $z_1 + z_2 < 0$, puis $a = -(z_1 + z_2) > 0$

Nous avons bien établi, dans le cas $\Delta > 0$: si P est stable, alors $a > 0$ et $b > 0$.

3)b) Réciproquement (toujours dans le cas $\Delta > 0$, donc dans le cas z_1 et z_2 réels), si $a > 0$ et $b > 0$:

$z_1 z_2 > 0$ donc z_1 et z_2 sont de même signe (et non nuls). Puis, comme $-(z_1 + z_2) > 0$, $z_1 + z_2 < 0$, et donc z_1 et z_2 sont tous deux strictement négatifs.

Il était en effet plus judicieux de commencer par tirer une information de $z_1 z_2 > 0$, plutôt que de considérer immédiatement $-(z_1 + z_2) > 0$, qui ne donne aucune information sur les signes de z_1 et z_2 .

Puis : $\operatorname{Re}(z_1) = z_1 < 0$ et $\operatorname{Re}(z_2) = z_2 < 0$

Toutes les racines complexes (z_1 et z_2) de P ont donc une partie réelle strictement négative. Autrement dit : si $a > 0$ et $b > 0$, alors P est stable.

« Les racines complexes ? Tu viens pourtant de dire qu'elles étaient réelles ! », pourrait me lancer quelqu'un qui a mal digéré la 1)... $\mathbb{R}$ est inclus dans $\mathbb{C}$, l'ami.

4) $\Delta \leq 0$, donc P admet deux racines complexes conjuguées $z_1 = \frac{-a - i\sqrt{-\Delta}}{2}$ et $z_2 = \overline{z_1} = \frac{-a + i\sqrt{-\Delta}}{2}$

Mais ça, c'est valable uniquement lorsque Δ est strictement négatif, non ? Non, ça reste valable pour Δ nul. Dans ce cas, P admet une racine double $z = -\frac{a}{2}$ (réelle, donc égale à son conjugué). Il est donc légitime ici de traiter les cas $\Delta < 0$ et $\Delta = 0$ ensemble.

$$\operatorname{Re}(z_1) = \operatorname{Re}(z_2) = -\frac{a}{2}.$$

P est donc stable si et seulement si $-\frac{a}{2} < 0$, c'est-à-dire si et seulement si $a > 0$

Oulala, notre conclusion est manquée non ? Et « $b > 0$ » dans tout ça ? Pas de panique n'oublions pas dans quel cas nous sommes...

Or, dans notre cas $\Delta \leq 0$, $a^2 - 4b \leq 0$, c'est-à-dire : $4b \geq a^2$

Reprenons donc nos équivalences :

P est stable $\iff a > 0$

$\iff \begin{cases} a > 0 \\ 4b \geq a^2 \end{cases}$ équivalence vraie car on a toujours $4b \geq a^2$ dans cette question

$\iff \begin{cases} a > 0 \\ b > 0 \end{cases}$

Pour la dernière équivalence : le sens $\Leftarrow$ est immédiat car, une fois encore, l'hypothèse $4b \geq a^2$ est valable dans toute la question (*je peux donc l'ajouter quand je veux*). Quant au sens $\Rightarrow$: si $4b \geq a^2$ avec $a > 0$, $a^2 > 0$ car $a^2 \geq 0$ en tant que carré de réel, et $a^2 \neq 0$ car $a \neq 0$. D'où $4b > 0$ et donc $b > 0$

Si vous voulez jouer la prudence (de peur d'écrire une équivalence fausse à un moment donné), vous pouvez aussi procéder par double implication. Ici, c'était l'occasion pour moi d'aller un peu plus vite en raisonnant directement par équivalence, et en vous montrant par étape comment justifier les $\iff$ moins évidents.

Nous avons établi, dans le cas où Δ est négatif ou nul, l'équivalence suivante :

P est stable si et seulement si $a > 0$ et $b > 0$.

En réalité, comme en 3a) et 3b), nous avons aussi établi cette équivalence dans le cas $\Delta > 0$, cette équivalence est en fait vraie quel que soit le signe de Δ .

Exercice 20

*Produits, transpositions, réseaux d'équivalences
se donnent pour mission d'épicer vos vacances.*

Énoncé : (temps conseillé : 45 min) (****) *d'après X-ENS-ESPCI 2011 PC*

Pour toute matrice M , on notera M^T la matrice transposée de M .

Pour tout $n \in \mathbb{N}^*$, on note S_n l'ensemble des matrices symétriques de $\mathcal{M}_n(\mathbb{R})$. On note S_n^+ l'ensemble suivant : $S_n^+ = \{M \in S_n, \forall X \in \mathcal{M}_{n,1}(\mathbb{R}), X^T M X \geq 0\}$

Soient a, b et d trois réels. On note : $A = \begin{pmatrix} 1 & 1 & 0 \\ 1 & 2 & 1 \\ 0 & 1 & 1 \end{pmatrix}$ et $B = \begin{pmatrix} a & b \\ b & d \end{pmatrix}$

1) Montrer que $A \in S_3^+$

2) Montrer que : $B \in S_2^+$ si et seulement si ($a \geq 0$, $d \geq 0$ et $ad - b^2 \geq 0$)

Dans le cas où $a \neq 0$, on pourra procéder à une factorisation par a

Remarques sur l'énoncé :

Remarquons que pour toute matrice M de $\mathcal{M}_n(\mathbb{R})$ et pour toute matrice colonne X de $\mathcal{M}_{n,1}(\mathbb{R})$, le produit matriciel MX appartient à $\mathcal{M}_{n,1}(\mathbb{R})$. Et comme $X^T \in \mathcal{M}_{1,n}(\mathbb{R})$, le produit matriciel $X^T(MX)$ - c'est-à-dire, par associativité, $X^T M X$ - appartient à $\mathcal{M}_{1,1}(\mathbb{R})$. C'est donc une matrice carrée de taille 1, que l'on pourra simplement assimiler à un réel. C'est ce qui donne du sens à « $X^T M X \geq 0$ » dans la définition de S_n^+ .

Par définition, une matrice M de $\mathcal{M}_n(\mathbb{R})$ appartient à S_n^+ si et seulement si elle est symétrique et pour tout $X \in \mathcal{M}_{n,1}(\mathbb{R})$, le réel $X^T M X$ est positif ou nul.

Correction de l'exercice 20 :

1) A est une matrice carrée de taille 3, et $A^T = A$, donc A est symétrique.

Immédiat, mais à dire quand même. Le caractère symétrique est le prérequis de base ici.

Donc : $A \in S_3$.

Montrons maintenant que pour tout X appartenant à $\mathcal{M}_{3,1}(\mathbb{R})$, $X^T A X \geq 0$.

$$\text{Soit } X = \begin{pmatrix} x \\ y \\ z \end{pmatrix} \in \mathcal{M}_{3,1}(\mathbb{R}). \quad AX = \begin{pmatrix} 1 & 1 & 0 \\ 1 & 2 & 1 \\ 0 & 1 & 1 \end{pmatrix} \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} x+y \\ x+2y+z \\ y+z \end{pmatrix}$$

$$\text{Donc } X^T A X = \begin{pmatrix} x & y & z \end{pmatrix} \times \begin{pmatrix} x+y \\ x+2y+z \\ y+z \end{pmatrix}$$

$$= x(x+y) + y(x+2y+z) + z(y+z) = x^2 + xy + xy + 2y^2 + yz + yz + z^2 = x^2 + 2y^2 + z^2 + 2xy + 2yz$$

Oh, des identités remarquables cachées...

$$\text{D'où : } X^T A X = x^2 + 2xy + y^2 + y^2 + 2yz + z^2 = (x+y)^2 + (y+z)^2 \geq 0$$

C'est une somme de carrés de réels. Attention au jour où vous devrez manipuler des matrices à coefficients complexes...

Nous avons bien montré : $\forall X \in \mathcal{M}_{3,1}(\mathbb{R}), X^T A X \geq 0$

En conclusion, $A \in S_3^+$

2) Factoriser par a , factoriser par a ... On verra en temps voulu.

B est une matrice carrée de taille 2, et $B^T = B$, donc $B \in S_2^+$

Pour tout X appartenant à $\mathcal{M}_{2,1}(\mathbb{R})$, en posant $X = \begin{pmatrix} x \\ y \end{pmatrix}$:

$$BX = \begin{pmatrix} a & b \\ b & d \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} ax+by \\ bx+dy \end{pmatrix} \text{ donc } X^T B X = \begin{pmatrix} x & y \end{pmatrix} \times \begin{pmatrix} ax+by \\ bx+dy \end{pmatrix} = x(ax+by) + y(bx+dy) \\ = ax^2 + 2bxy + dy^2$$

Il nous faut établir l'équivalence suivante :

$$(\forall X \in \mathcal{M}_{2,1}(\mathbb{R}), X^T B X \geq 0) \iff (a \geq 0, d \geq 0 \text{ et } ad - b^2 \geq 0)$$

Ce qui revient à établir l'équivalence :

$$(\forall (x,y) \in \mathbb{R}^2, ax^2 + 2bxy + dy^2 \geq 0) \iff (a \geq 0, d \geq 0 \text{ et } ad - b^2 \geq 0)$$

Raisonnement directement par équivalence semble assez compliqué ici...

Procédons par double implication.

« $\Rightarrow$ » Supposons que pour tous réels x et y , $ax^2 + 2bxy + dy^2 \geq 0$

Peut-être est-il temps de distinguer les cas $a = 0$ et $a \neq 0$, pour pouvoir factoriser par a dans ce dernier, comme préconisé par l'énoncé.

Par disjonction de cas :

- Si $a = 0$, la conclusion $a \geq 0$ est immédiate. C'est déjà ça de pris...

Nous savons aussi : $\forall x, y \in \mathbb{R}, 2bxy + dy^2 \geq 0$

Et il nous faut aboutir à : $d \geq 0$, et $ad - b^2 \geq 0$, c'est-à-dire, dans notre cas : $-b^2 \geq 0$, c'est-à-dire, en fait : $b = 0$ (parce que, dans tous les cas $-b^2 \leq 0$, puisque b^2 est un carré de réel).

Profitions du fait que nous savons l'inégalité $2bxy + dy^2 \geq 0$ vraie pour tous réels x et y : nous pouvons prendre les x et y que l'on veut, elle reste vraie...

En particulier, en prenant $y = 1$ et $x = 0$: $d \geq 0$

Cool ! Maintenant, comment obtenir b égale zéro ?

En prenant $y = 1$, nous savons : $\forall x \in \mathbb{R}, 2bx + d \geq 0$

Oui, je peux aussi fixer l'un et garder l'autre quelconque...

La fonction affine $f : x \mapsto 2bx + d$ est positive sur $\mathbb{R}$, donc nécessairement constante.

Le coefficient directeur de la droite d'équation $y = 2bx + d$ est donc nul. Autrement dit, $2b = 0$, puis $b = 0$.

Eh oui... Une fonction affine, dont la courbe représentative est une droite, ne peut être de signe constant (autrement dit, sa droite représentative ne reste que d'un côté du plan par rapport à l'axe des abscisses) que si elle est constante.

Un raisonnement légèrement différent sur les limites était aussi possible : par l'absurde, si $b \neq 0$, alors, selon le signe de b , l'une des limites en $+\infty$ ou en $-\infty$ de $2bx + d$ est $-\infty$, ce qui est absurde vu l'inégalité : $\forall x \in \mathbb{R}, 2bx + d \geq 0$. Donc $b = 0$.

Nous avons bien établi, dans le cas $a = 0$: $a \geq 0$, $d \geq 0$, et $ad - b^2 = -b^2 \geq 0$

- Si $a \neq 0$: pour tous réels x et y , $a(x^2 + 2\frac{b}{a}xy + \frac{d}{a}y^2) \geq 0$

$$\text{Or : } a\left(x^2 + \frac{2b}{a}xy + \frac{d}{a}y^2\right) = a\left(x^2 + 2\frac{by}{a} \times x + \frac{b^2y^2}{a^2} + \frac{d}{a}y^2 - \frac{b^2y^2}{a^2}\right)$$

Souvenirs de forme canonique de Première... J'ai ajouté un $\frac{b^2y^2}{a^2}$ qui n'y était pas, pour former une identité remarquable. Bien sûr, j'ai aussi soustrait un $\frac{b^2y^2}{a^2}$ à la fin pour compenser. Une égalité, c'est sacré.

$$\text{Donc : } \forall x, y \in \mathbb{R}, a\left(x^2 + \frac{2b}{a}xy + \frac{d}{a}y^2\right) = a\left[\left(x + \frac{by}{a}\right)^2 + \frac{(ad - b^2)y^2}{a^2}\right] \geq 0 \quad (*)$$

En particulier : pour $x = 1 - \frac{by}{a}$ et $y = 0$:

$$a\left[\left(1 - \frac{by}{a} + \frac{by}{a}\right)^2 + 0\right] = a \times 1 = \underline{a \geq 0}. \text{ Et même ici : } a > 0 \text{ (puisque } a \neq 0)$$

$$\text{Donc (en divisant (*) par } a) : \forall x, y \in \mathbb{R}, \left(x + \frac{by}{a}\right)^2 + \frac{(ad - b^2)y^2}{a^2} \geq 0$$

$$\text{En particulier : pour } x = -\frac{by}{a} \text{ et } y = 1 : \frac{ad - b^2}{a^2} \geq 0 \text{ puis } \underline{ad - b^2 \geq 0}$$

Reste à prouver que $d \geq 0$... Quoi, encore un choix particulier de x et y ? Non non, c'est sous nos yeux.

D'où : $ad \geq b^2 \geq 0$, puis $ad \geq 0$. Et comme, $a > 0$, nous en déduisons : $\underline{d \geq 0}$.

Nous avons donc aussi établi dans le cas $a \neq 0$: $\underline{a \geq 0, d \geq 0, \text{ et } ad - b^2 \geq 0}$

L'implication suivante est démontrée :

$$(\forall (x, y) \in \mathbb{R}^2, ax^2 + 2bxy + dy^2 \geq 0) \implies (a \geq 0, d \geq 0 \text{ et } ad - b^2 \geq 0)$$

« $\Leftarrow$ » Supposons maintenant que $a \geq 0, d \geq 0$, et $ad - b^2 \geq 0$

• Si $a = 0$, $-b^2 \geq 0$ donc $b = 0$. D'où : $\forall x, y \in \mathbb{R} : ax^2 + 2bxy + dy^2 = dy^2 \geq 0$ car $d \geq 0$

• Si $a \neq 0$ (donc en fait $a > 0$), on peut écrire, comme en (*) :

$$\forall x, y \in \mathbb{R} : ax^2 + 2bxy + dy^2 = a\left[\left(x + \frac{by}{a}\right)^2 + \frac{(ad - b^2)y^2}{a^2}\right], \text{ avec } \left(x + \frac{by}{a}\right)^2 \geq 0, \frac{(ad - b^2)y^2}{a^2} \geq 0$$

(car $ad - b^2 \geq 0$) et, comme précisé plus haut, $a > 0$

$$\text{D'où : } \forall x, y \in \mathbb{R}, ax^2 + 2bxy + dy^2 \geq 0$$

Oui, cette implication-ci est plus immédiate, même si nous n'avons pas eu à réeffectuer la transformation ().*

Nous avons donc établi :

$$(a \geq 0, d \geq 0 \text{ et } ad - b^2 \geq 0) \implies (\forall (x, y) \in \mathbb{R}^2, ax^2 + 2bxy + dy^2 \geq 0)$$

$$\text{D'où l'équivalence : } (\forall (x, y) \in \mathbb{R}^2, ax^2 + 2bxy + dy^2 \geq 0) \iff (a \geq 0, d \geq 0 \text{ et } ad - b^2 \geq 0)$$

Autrement dit, nous avons montré :

$$(\forall X \in \mathcal{M}_{2,1}(\mathbb{R}), X^T B X \geq 0) \iff (a \geq 0, d \geq 0 \text{ et } ad - b^2 \geq 0)$$

$$\text{Autrement dit*}, \text{ nous avons établi : } B \in S_2^+ \text{ si et seulement si } (a \geq 0, d \geq 0 \text{ et } ad - b^2 \geq 0)$$

**Celui-là était fait exprès, pour finir en beauté ; c'est un tic irrésistible chez moi de mettre des « autrement dit » à tout va, j'aime bien... Tout comme cette manie insupportable de camoufler l'embarras d'une fin de phrase ou d'un adieu avec trois petits points. Bonnes vacances...*

