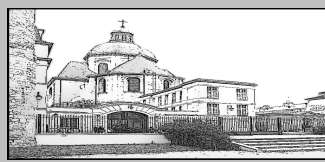




LYCEE CHARLEMAGNE

Révisions

M.P.S.I.2



2024

2025

Compétences

Le minimum vital pour aborder la seconde année...

○ Calculer module et argument de $1 + e^{i\theta}$. Calculer module et argument de $e^{i\alpha} + e^{i\beta}$.

Graphiquement, $1 + e^{i\theta}$ a pour argument $\theta/2$ et on mesure son demi module dans un triangle rectangle.

Mais proprement, on factorise : $e^{i\theta} + 1 = e^{i\theta/2} \cdot (e^{i\theta/2} + e^{-i\theta/2}) = 2 \cdot \cos(\theta/2) \cdot e^{i\theta/2}$.

Le module est donc bien $2 \cdot \cos(\theta/2)$ (positif si on a bien pris θ entre $-\pi$ et π).

L'argument est $\theta/2$.

On factorise de même : $e^{i\alpha} + e^{i\beta} = e^{i \cdot \frac{\alpha+\beta}{2}} \cdot (e^{i \cdot \frac{\alpha-\beta}{2}} + e^{i \cdot \frac{-\alpha+\beta}{2}}) = 2 \cdot \cos\left(\frac{\alpha-\beta}{2}\right) \cdot e^{i \cdot \frac{\alpha+\beta}{2}}$.

En passant aux parties réelles et imaginaires, on retrouve les formules classiques :

$\cos(\alpha) + \cos(\beta) = 2 \cdot \cos\left(\frac{\alpha-\beta}{2}\right) \cdot \cos\left(\frac{\alpha+\beta}{2}\right)$ et	$\sin(\alpha) + \sin(\beta) = 2 \cdot \cos\left(\frac{\alpha-\beta}{2}\right) \cdot \sin\left(\frac{\alpha+\beta}{2}\right)$
---	--

○ Calculer les deux racines carrées d'un complexe donné sous forme cartésienne. Résoudre une équation du second degré à coefficients dans $\mathbb{C}$.

On donne $Z = a + i \cdot b$ et on veut trouver z de la forme $x + i \cdot y$ vérifiant $z^2 = Z$ (avec évidemment, a , b , x et y réels, méfiez vous du piège). On développe $(x + i \cdot y)^2$ et surtout, on pense aussi aux modules

: $|z|^2 = |Z|$. On a trois équations :
$$\begin{cases} x^2 - y^2 = a \\ 2xy = b \\ x^2 + y^2 = \sqrt{a^2 + b^2} \end{cases}$$
 On utilise première et troisième

ligne, et la seconde ne sert qu'à choisir le signe de x et y : si b est positif, x et y sont de même signe ; si b est négatif, x et y sont de signes opposés.

On a alors deux racines : $\frac{a + \sqrt{a^2 + b^2}}{2} + \varepsilon \cdot i \cdot \frac{\sqrt{a^2 + b^2} - a}{2}$ et son opposé, avec $\varepsilon = \text{Signe}(b)$.

Il n'est pas demandé de connaître la formule par cœur, mais de savoir exploiter rapidement le module $x^2 + y^2 = \sqrt{a^2 + b^2}$. On pourra éventuellement retenir une formule toute prête avec $\Re(Z) + |z|$.

On peut prolonger l'histoire avec les formules en arc moitié.

Pour résoudre l'équation $a \cdot z^2 + b \cdot z + c = 0$ avec a , b et c complexes (et a non nul), on a juste besoin de calculer le discriminant et d'en extraire les racines. Les discussions sur le signe éventuel de Δ deviennent totalement oiseuses et ridicules.

On calcule : $\Delta = b^2 - 4 \cdot a \cdot c$.

On cherche δ vérifiant $\delta^2 = \Delta$ (ce qu'on appelle rechercher les racines carrées de Δ mais qu'on ne notera jamais par écrit $\delta = \sqrt{\Delta}$).

Les deux racines sont alors $\frac{-b + \delta}{2 \cdot a}$ et $\frac{-b - \delta}{2 \cdot a}$.

Comme il y a deux valeurs de δ possibles, opposées, il y a a bien un seul couple de solutions.

Dans le cas où a , b et c sont réels, et Δ négatif, on trouve pour δ un imaginaire pur, qui coïncide avec les formules lourdingues de Terminale.

○ Donner la définition d'un groupe. Montrer que le neutre est unique.

Un groupe est un ensemble G muni d'une loi $*$ et il est S.A.I.N. Mais on cite les propriétés dans l'ordre naturel.

I	la loi est interne	$\forall (a, b) \in G^2, a * b \in G$
A	la loi est associative	$\forall (a, b, c) \in G^3, (a * b) * c = a * (b * c)$
N	il y a un neutre	$\exists e \in G, \forall a \in G, a * e = e * a = a$
S	chaque élément a un symétrique	$\forall a \in G, \exists \alpha \in G, a * \alpha = \alpha * a = e$

Il faut quantifier les variables dans le bon ordre pour neutre et symétrique.

On ne parle pas du symétrique avant d'avoir parlé du neutre.

Le groupe est $(G, *)$, on cite la loi avec ; ainsi, $\mathbb{Z}$ n'est pas un groupe, c'est $(\mathbb{Z}, +)$ qui en est un.

En option, on a la commutativité : $\forall (a, b) \in G^2, a * b = b * a$.

Dans un groupe non commutatif, on a juste $\exists(a, b) \in G^2$, $a * b \neq b * a$ mais il reste des éléments qui sont obligés de commuter entre eux.

La preuve de l'unicité du neutre est une preuve directe, que l'on peut prendre pour une preuve par l'absurde. On se donne deux éléments qui tiennent le rôle de neutre, et on montre que c'est le même. On suppose donc que e et ε sont neutres, et on calcule :

$e * \varepsilon = e$ car ε est neutre (à droite)	et	$e * \varepsilon = \varepsilon$ car e est neutre (à gauche)
---	----	---

Par transitivité de l'égalité, $e = \varepsilon$.

C'est avec le même type de raisonnement qu'on prouve que si a un deux symétriques α et ∞ , alors c'est le même, il suffit d'écrire : $\alpha = \alpha * (a * \infty) = (\alpha * a) * \infty = \infty$.

○ Donner la définition et le domaine de définition des fonctions Arcsin , Arccos et Arctan . Calculer leurs dérivées. Savoir les intégrer par intégration par parties.

fonction	Arcsin	Arccos	Arctan
domaine	$x \in [-1, 1]$	$x \in [-1, 1]$	$x \in \mathbb{R}$
image	$[-\pi/2, \pi/2]$	$[0, \pi]$	$] - \pi/2, \pi/2[$
définition	$\theta = \text{Arcsin}(x)$ $\Leftrightarrow \left(\sin(\theta) = x \text{ et } \theta \leq \pi/2 \right)$	$\theta = \text{Arccos}(x)$ $\Leftrightarrow \left(\sin(\theta) = x \text{ et } 0 \leq \theta \leq \pi \right)$	$\theta = \text{Arctan}(x)$ $\Leftrightarrow \left(\sin(\theta) = x \text{ et } \theta < \pi/2 \right)$
dérivée	$\text{Arcsin}'(x) = \frac{1}{\sqrt{1-x^2}}$ sur $] -1, 1[$	$\text{Arccos}'(x) = \frac{-1}{\sqrt{1-x^2}}$ sur $] -1, 1[$	$\text{Arctan}'(x) = \frac{1}{1+x^2}$ sur $] -\infty, +\infty[$
une primitive	$x \cdot \text{Arcsin}(x) + \sqrt{1-x^2}$	$x \cdot \text{Arccos}(x) - \sqrt{1-x^2}$	$x \cdot \text{Arctan}(x) - \frac{\ln(1+x^2)}{2}$

Les primitives ne sont pas exigibles. Il est juste demandé de savoir les obtenir par parties : $\int_a^b \text{Arctan}(x).dx =$

$$\left[x \cdot \text{Arctan}(x) \right]_a^b - \int_a^b \frac{x}{1+x^2} \cdot dx$$

On peut aussi les obtenir par un argument géométrique de symétrie du graphe par rapport à la première bissectrice.

○ Démontrer la formule de Taylor avec reste intégrale pour une application de classe C^{n+1} sur un intervalle $[a, a+h]$ (h pas forcément positif).

$$f(a+h) = \sum_{k=0}^n \frac{f^{(k)}(a)}{k!} \cdot h^k + \frac{h^{n+1}}{n!} \cdot \int_{t=0}^1 (1-t)^n \cdot f^{(n+1)}(a+t.h).dt$$

On démontre cette formule par récurrence sur n .

$$\text{On initialise en calculant } \int_{t=0}^1 t \cdot f'(a+t.h).dt = \left[f(a+t.h) \right]_{t=0}^1 = f(a+h) - f(a).$$

On prouve l'hérédité en calculant $\frac{h^{n+1}}{n!} \cdot \int_{t=0}^1 (1-t)^n \cdot f^{(n+1)}(a+t.h).dt$ par parties

$h^{n+1} \cdot f^{(n+1)}(a+t.h)$	$\hookrightarrow$	$h^{n+2} \cdot f^{(n+2)}(a+t.h)$
$\frac{(1-t)^n}{n!}$	$\hookleftarrow$	$-\frac{(1-t)^{n+1}}{(n+1)!}$

Il se convertit en $\frac{h^{n+1}}{(n+1)!} \cdot f^{n+1}(a)$ qu'on intègre

à la somme et $\frac{h^{n+2}}{(n+1)!} \cdot \int_{t=0}^1 (1-t)^{n+1} \cdot f^{(n+2)}(a+t.h).dt$ qui devient le nouveau reste intégrale.

On peut aussi en une fois prendre la fonction $t \mapsto \sum_{k=0}^n \frac{((1-t).h)^k}{k!} \cdot f^{(k)}(a+t.h)$ (notée F), la dériver,

$$\text{la calculer en 0 et en 1 et écrire } F(1) - F(0) = \int_{t=0}^1 F'(t).dt.$$

○ Connaître le développement limité d'ordre n quand u tend vers 0 de $\ln(1+u)$ et $\ln(1-u)$. Déterminer le développement limité du logarithme en un point a .

On écrit avant tout $\frac{1}{1-x} - \frac{x^n}{1-x} = 1+x+\dots+x^{n-1}$ qui devient $\frac{1}{1-x} = 1+x+\dots+x^{n-1}+o(x^{n-1})_{x \rightarrow 0}$ (x plus petit que 1 en valeur absolue, ce qui ne pose pas de problème car il va tendre vers 0 par encadrement).

On intègre de 0 à u (qui va tendre vers 0) : $\ln(1-u) = -u - \frac{u^2}{2} - \frac{u^3}{3} - \dots - \frac{u^n}{n} + o(u^n)_{u \rightarrow 0}$
 Cette formule est la plus simple, elle ne contient que des signes moins, commence par l'équivalent $\ln(1-u) \sim -u$.
 Pour se souvenir des $\frac{1}{k}$, on se souvient que même si ça n'a aucun sens, pour u égal à 1, on a la divergence de la série harmonique.

On remplace u par $-u$: $\ln(1+u) = u - \frac{u^2}{2} + \frac{u^3}{3} - \dots + (-1)^{n-1} \cdot \frac{u^n}{n} + o(u^n)_{u \rightarrow 0}$.
 Dans cette formule, on commence par l'équivalent $\ln(1+u) \sim u$, les signes alternent, et on peut penser à la série harmonique alternée de somme $\ln(2)$.

En a , on écrit $\ln(a+h) = \ln(a) + \ln\left(1 + \frac{h}{a}\right) = \ln(a) + \sum_{k=1}^n (-1)^{k-1} \cdot \frac{h^k}{k \cdot a^k} + o(h^n)_{h \rightarrow 0}$.
 On peut d'ailleurs l'obtenir rapidement par la formule de Taylor si on sait dériver efficacement le logarithme à un ordre quelconque : $\ln^n = x \mapsto (-1) \cdot (n-1)! \cdot x^{-n}$ (n différent de 0 bien sûr).

On rappelle qu'il n'y a aucune formule pour le logarithme en 0, ni même d'équivalent autre que $\ln(x) \sim \ln(x)$.

○ Déterminer partie réelle et partie imaginaire d'une somme de n complexes donnés sous forme cartésienne. Donner module et argument d'un produit/quotient de n complexes donnés sous forme polaire.

C'est le programme officiel qui dit ça à un moment.

Les opérateurs partie réelle et partie imaginaire sont $\mathbb{R}$ -linéaires : $\Re\left(\sum_{k=1}^n z_k\right) = \sum_{k=1}^n \Re(z_k)$ et

$$\Im\left(\sum_{k=1}^n z_k\right) = \sum_{k=1}^n \Im(z_k).$$

On a même $\Re\left(\sum_{k=1}^n a_k \cdot z_k\right) = \sum_{k=1}^n a_k \cdot \Re(z_k)$ si les a_k sont des réels.

De même, le module du produit est le produit des modules.

Le module du quotient est le quotient des modules.

L'argument du produit n'est la somme des arguments que si on pense à réduire modulo 2π .

**○ Donner la liste des racines complexes de l'équation $z^n = 1$ d'inconnue complexe z . Calculer leur somme, leur produit.
 Donner la liste des n racines n^{iemes} d'un complexe a donné sous forme polaire $\rho \cdot e^{i \cdot \alpha}$.**

L'équation $z^n = 1$ conduit déjà par condition nécessaire à $|z| = 1$. En notant θ l'argument du complexe cherché, on a $n \cdot \theta = 0 [2\pi]$. On trouve la liste des n racines : $e^{2 \cdot i \cdot k \cdot \pi / n}$ avec k allant de 0 à $n-1$.

On trouve n racines réparties équitablement sur le cercle unité.

Pour k égal à 0, on a toujours la racine réelle 1.

Suivant la parité de n , on a ou non l'autre racine réelle -1 .

En écrivant sous forme factorisée $\prod_{k=0}^{n-1} (X - e^{2 \cdot i \cdot k \cdot \pi / n}) = X^n - 1$ et en utilisant les formules de Viète,

on trouve la somme des n racines de l'unité est nulle.

Sauf pour n égal à 1 ; il n'y a qu'une racine, de "somme" 1.

On retrouve le résultat en sommant : $\sum_{k=0}^{n-1} e^{2 \cdot i \cdot k \cdot \pi / n} = \frac{1 - e^{2 \cdot i \cdot n \cdot \pi / n}}{1 - e^{2 \cdot i \cdot \pi / n}} = 0$ (série géométrique).

Pour le produit, on fait encore usage des formules de Viète : $\prod_{k=0}^{n-1} e^{2 \cdot i \cdot k \cdot \pi / n} = (-1)^n$.

On la aussi avec $\prod_{k=0}^{n-1} e^{2 \cdot i \cdot k \cdot \pi / n} = (e^{2 \cdot i \cdot \pi / n})^{0+1+2+\dots+(n-1)}$.

On l'a enfin en regroupant les racines non réelles deux à deux par conjugaison. Il ne reste que la racine 1 (neutre) et la racine -1 si n est impair.

Si on résout $(r \cdot e^{i \cdot \theta})^n = \rho \cdot e^{i \cdot \alpha}$, on trouve le même module pour toutes : $r = \sqrt[n]{\rho}$ et des arguments en progression arithmétique : $\theta = \frac{\alpha}{n} + \frac{2 \cdot i \cdot k \cdot \pi}{n}$ avec k décrivant $\text{range}(n)$.

Si on dispose d'une des solutions (solution particulière), on trouve les autres en multipliant par les

racines n^{iemes} de l'unité : $z_0.e^{2.i.k.\pi/n}$ avec k de 0 à $n-1$.

Classiquement, l'équation $x^4 = -1$ a pour solution particulière $\frac{1+i}{\sqrt{2}}$ et on passe aux autres en multipliant par i , -1 et $-i$.

○ Donner la définition de la fonction indicatrice d'une partie A dans un ensemble E . Exprimer l'indicatrice de l'intersection de deux parties, de leur réunion, de leur différence symétrique et du complémentaire d'une partie A . Calculer $\sum_{e \in E} \chi_A(e)$.

On dispose de deux notations : χ_A et 1_A : $1_A = x \mapsto \begin{cases} 1 & \text{si } x \in A \\ 0 & \text{sinon} \end{cases}$. Les éléments de A répondent 1, les autres 0.

$1_{(A^c)} = 1 - 1_A$	$1_{A \cap B} = 1_A \cdot 1_B$	$1_{A \cup B} = 1_A + 1_B - 1_A \cdot 1_B$	$1_{A \Delta B} = 1_A + 1_B - 2 \cdot 1_A \cdot 1_B$
-----------------------	--------------------------------	--	--

Les deux dernières formules sont utiles pour les cardinaux. Mais on pourra préférer $1_{A \Delta B} = (1_A - 1_B)^2$ ou même $1_{A \Delta B} = |1_A - 1_B|$ qui justifie le nom.

Pour compter les éléments de A , on écrit $\sum_{e \in E} \chi_A(e) = \text{Card}(A)$ qui n'a d'intérêt que si A est un ensemble fini.

On pourra utiliser les notations $\text{Card}(A)$, $\#A$ ou même $|A|$ pour désigner le cardinal d'un ensemble. Pour les groupes, la coutume est même de parler d'"ordre".

○ Simplifier une somme télescopique $\sum_{k=0}^n (a_{k+1} - a_k)$.

La formule de base est $\sum_{k=0}^n (a_{k+1} - a_k) = a_{n+1} - a_0$.

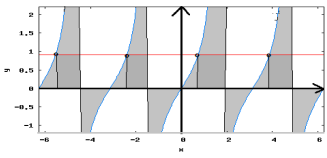
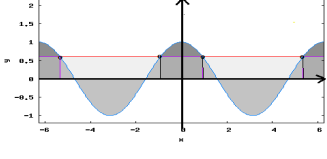
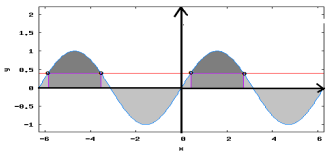
On peut la prouver par récurrence sur n .

On peut aussi décaler les indices $\sum_{k=0}^n (a_{k+1} - a_k) = \sum_{k=0}^n a_{k+1} - \sum_{k=0}^n a_k = \sum_{k=1}^{n+1} a_k - \sum_{k=0}^n a_k$.

On dispose de plusieurs formules : $\sum_{0 \leq k < n} (a_{k+1} - a_k) = a_n - a_0$ ou $\sum_{k=p}^q (a_{k+1} - a_k) = a_{q+1} - a_p$ si q est plus grand que p .

On peut aussi écrire des formules à l'ordre 2 : $\sum_{k=0}^n (a_{k+2} - 2.a_{k+1} + a_k)$ que je vous laisse compléter.

○ Résoudre et discuter $\cos(\theta) = a$ d'inconnue a et de paramètre θ . Compter le nombre de solutions sur un intervalle donné. Même question avec $\sin(\theta) = a$. Même question avec $\tan(\theta) = a$.

Équation	Condition	Solutions	Graphique
$\tan(\theta) = a$	aucune	$\{ \text{Arctan}(a) + k.\pi \mid k \in \mathbb{Z} \}$	
$\cos(\theta) = a$	$ a \leq 1$	$\{ \text{Arccos}(a) + 2.k.\pi \mid k \in \mathbb{Z} \}$ $\cup \{ -\text{Arccos}(a) + 2.k.\pi \mid k \in \mathbb{Z} \}$	
$\sin(\theta) = a$	$ a \leq 1$	$\{ \text{Arcsin}(a) + 2.k.\pi \mid k \in \mathbb{Z} \}$ $\cup \{ \pi - \text{Arcsin}(a) + 2.k.\pi \mid k \in \mathbb{Z} \}$	

Dans un intervalle de longueur $2.N.\pi$ l'équation $\cos(\theta) = a$ admet $2.N$ solutions.

Dans un intervalle de longueur $P.\pi$ l'équation $\tan(\theta) = a$ admet P solutions.

Quelques équations conduisent à des solutions de forme simple

$\cos(\theta) = 1$	$\{ 2.k.\pi \mid k \in \mathbb{Z} \}$	$\sin(\theta) = 1$	$\{ \frac{\pi}{2} + 2.k.\pi \mid k \in \mathbb{Z} \}$
$\cos(\theta) = 0$	$\{ \frac{\pi}{2} + k.\pi \mid k \in \mathbb{Z} \}$	$\sin(\theta) = 0$	$\{ k.\pi \mid k \in \mathbb{Z} \}$
$\cos(\theta) = -1$	$\{ (2.k+1).\pi \mid k \in \mathbb{Z} \}$	$\sin(\theta) = -1$	$\{ -\frac{\pi}{2} + 2.k.\pi \mid k \in \mathbb{Z} \}$

○ Passer dans $\mathbb{R}$ de “ $x \in [a, b]$ ” à “ $|x - \alpha| \leq r$ ” et vice versa.

On a évidemment $|x - \alpha| \leq r \Leftrightarrow x \in [\alpha - r, \alpha + r]$.

On a aussi $x \in [a, b] \Leftrightarrow \left| x - \frac{a+b}{2} \right| \leq \frac{|b-a|}{2}$ (centre et rayon).

On utilisera ces propriétés pour passer de $|x - a| \leq \varepsilon$ à $a\varepsilon \leq x \leq a + \varepsilon$ dans les définitions des limites.

On saura aussi utiliser que $|y - b| \leq |b|/2$ entraîne que y est non nul, du même signe que b .

○ Donner la définition de $(A, +, \times)$ est un anneau. Connaître la définition de “commutatif”, “intègre”.

Prouver que dans un anneau, on a toujours $a \times 0 = 0$ (où 0 est le neutre de la première loi).

groupe	anneau	espace vectoriel	algèbre
une loi $\oplus$	deux lois $\oplus, \otimes$	deux lois $\oplus, \odot$	trois lois $\oplus, \otimes, \odot$

Dans un anneau, on a donc deux lois internes, associatives, avec la seconde distributive sur la première. Les anneaux où chacune est distributive sur l'autre sont appelés anneaux de Boole et sont du domaine de la logique.

$\oplus$	
Interne	$\forall(a, b) \in A^2, a \oplus b \in A$
Associative	$\forall(a, b, c) \in A^3, (a \oplus b) \oplus c = a \oplus (b \oplus c)$
Neutre	$\exists 0 \in A, \forall a \in A, a \oplus 0 = 0 \oplus a = a$
Symétriques	$\forall a \in A, \exists -a \in A, a \oplus (-a) = (-a) \oplus a = 0$
Commutative	$\forall(a, b) \in A^2, a \oplus b = b \oplus a$
$\otimes$	
Interne	$\forall(a, b) \in A^2, a \otimes b \in A$
Associative	$\forall(a, b, c) \in A^3, (a \otimes b) \otimes c = a \otimes (b \otimes c)$
Neutre	$\exists 1 \in A, \forall a \in A, a \otimes 1 = 1 \otimes a = a$
Distributive sur $\oplus$	$\forall(a, b, c) \in A^3, a \otimes (b \oplus c) = (a \otimes b) \oplus (a \otimes c)$
	$\forall(a, b, c) \in A^3, (b \oplus c) \otimes a = (b \otimes a) \oplus (c \otimes a)$

Certaines conventions n'exigent pas que la seconde loi ait un neutre.

Les règles de priorités pousseront à écrire $b \times a + c \times a$ à la place de $(b \times a) + (c \times a)$.
 Un anneau est le domaine adapté pour écrire la formule du binôme. Surtout si cet anneau est commutatif.
 Avec la distributivité, on peut développer $(a + b).(c + d)$ par exemple en quatre termes.

L'anneau pourra être commutatif : $\forall(a, b) \in A^2, a \otimes b = b \otimes a$.

La loi "additive" est forcément commutative. L'adjectif ne porte que sur la seconde.

Si on note donc 0 le neutre de la première loi, on a $a \otimes (0 + 0) = a \otimes 0$ car 0 est neutre. On distribue : $a \otimes 0 + a \otimes 0 = a \otimes 0$. On simplifie par le symétrique de $a \otimes 0$ et il reste juste $a \otimes 0 = 0$.

Bref, le neutre additif est absorbant pour la multiplication.

Un produit de facteurs est nul si l'un des facteurs est nul.

L'intégrité est une "réciproque", c'est le "seulement si".

Elle admet plusieurs formulations, en jouant sur $(p \Rightarrow q) = (\bar{p} \text{ ou } q) = (\bar{q} \Rightarrow \bar{p})$.

$\forall(a, b) \in A^2, (a \otimes b = 0) \Rightarrow (a = 0 \text{ ou } b = 0)$	
$\forall(a, b) \in A^2, (a \otimes b = 0 \text{ et } a \neq 0) \Rightarrow (b = 0)$	
$\forall(a, b) \in A^2, (a \neq 0 \text{ et } b \neq 0) \Rightarrow (a \otimes b \neq 0)$	

○ **Factoriser une somme du type $\sum_{\substack{i \leq p \\ k \leq q}} a_i.b_k$.**

Il s'agit d'une somme double où les deux indices varient en toute indépendance. Elle se sépare en produit de deux sommes : $\sum_{\substack{i \leq p \\ k \leq q}} a_i.b_k = \left(\sum_{i \leq p} a_i \right) . \left(\sum_{k \leq q} b_k \right)$.

Il faut savoir utiliser cette formule dans les deux sens, et l'utiliser y compris pour développer :

$$\left(\sum_{i \leq n} a_i \right)^2 = \left(\sum_{i \leq n} b_i \right) . \left(\sum_{k \leq n} a_k \right) = \sum_{\substack{i \leq n \\ k \leq n}} a_i.a_k.$$

Si les variables sont liées, il faudra se méfier :

$$\sum_{0 \leq i \leq k \leq n} (a_i.b_k) = \sum_{i=0}^n \left(a_i . \sum_{k=i}^n b_k \right) = \sum_{k=0}^n \left(b_k . \sum_{i=0}^k a_i \right) \text{ par exemple,}$$

et pire encore pour $\sum_{0 \leq i \leq k \leq n} c_i^k$.

○ **Transformer $\cos(a) + \cos(b)$ en $2.\cos(\dots)$ et autres formules de ce type. Formules réciproques (produits en sommes). Calcul de $\int_a^b \cos(p.\theta). \cos(q.\theta).d\theta$.**

On écrit $\cos(a) = \cos\left(\frac{a+b}{2} + \frac{a-b}{2}\right)$ et $\cos(a) = \cos\left(\frac{a+b}{2} - \frac{a-b}{2}\right)$. On développe, et on trouve quatre formules :

$$\begin{aligned} \cos(a) + \cos(b) &= 2.\cos\left(\frac{a+b}{2}\right).\cos\left(\frac{a-b}{2}\right) \\ \cos(a) - \cos(b) &= -2.\sin\left(\frac{a+b}{2}\right).\sin\left(\frac{a-b}{2}\right) \\ \sin(a) + \sin(b) &= 2.\sin\left(\frac{a+b}{2}\right).\cos\left(\frac{a-b}{2}\right) \\ \sin(a) - \sin(b) &= 2.\cos\left(\frac{a+b}{2}\right).\sin\left(\frac{a-b}{2}\right) \end{aligned}$$

On les vérifie par des cas particuliers tels que $b = 0$ ou $a = b$.

Réciproquement, on a	$\cos(a). \cos(b) = \frac{\cos(a+b) + \cos(a-b)}{2}$
	$\sin(a). \sin(b) = \frac{\cos(a-b) - \cos(a+b)}{2}$
	$\sin(a). \cos(b) = \frac{\sin(a+b) + \sin(a-b)}{2}$

Pour calculer $\int_a^b \cos(p.\theta). \cos(q.\theta).d\theta$, on linéarise la fonction à intégrer en $\frac{\cos((p-q).\theta) + \cos((p+q).\theta)}{2}$,

et on obtient $\left[\frac{\sin((p-q).\theta)}{2.(p-q)} + \frac{\sin((p+q).\theta)}{2.(p+q)} \right]_a^b$ (pour p différent de q évidemment). Si les bornes sont des multiples entiers de π , on trouve 0 et on traduit par l'orthogonalité de deux fonctions pour un produit scalaire naturel.

○ Montrer que le déterminant de deux vecteurs du plan est l'aire algébrique du parallélogramme qu'ils engendrent.

On prend deux vecteurs $\overrightarrow{AB} = a.\overrightarrow{i} + c.\overrightarrow{j}$ et $\overrightarrow{AD} = b.\overrightarrow{i} + d.\overrightarrow{j}$. On construit le parallélogramme (A, B, C, D) . On découpe son aire dans celle d'un grand rectangle de côtés $(a+b)$ et $(c+d)$. On déduit l'aire de deux petits rectangles : $2.b.c$ et de quatre triangles rectangles recomposant deux rectangles : $a.c$ et $b.d$.

Il reste $Aire = (a+b).(c+d) - 2.b.c - a.c - b.d = a.d - b.c = \begin{vmatrix} a & b \\ c & d \end{vmatrix} = \det(\overrightarrow{AB}, \overrightarrow{AD})$.

○ Savoir factoriser $X^2 - 2.X.\cos(\theta) + 1$.

On peut évidemment calculer le discriminant (négatif par théorème de Pythagore), en extraire la racine complexe $2.i.\sin(\theta)$ et obtenir les deux racines.

On peut aussi penser spontanément à $e^{i.\theta}$ et $e^{-i.\theta}$ qu'on vérifie dans l'équation.

On peut plus agréablement penser à ces racines et identifier somme et produit des racines :

$$X^2 - 2.X.\cos(\theta) + 1 = (X - e^{i.\theta}).(X - e^{-i.\theta})$$

Les deux racines sont complexes, conjuguées, de module 1.

On rappelle les cas particuliers $X^2 + 1 = (X - i).(X + i)$ et $X^2 + X + 1 = (X - e^{2.i.\pi/3}).(X - e^{-2.i.\pi/3})$.

On décompose en éléments simples :

$$\frac{2.X - 2.\cos(\theta)}{X^2 - 2.\cos(\theta).X + 1} = \frac{1}{X - e^{i.\theta}} + \frac{1}{X - e^{-i.\theta}} \text{ et } \frac{2.i.\sin(\theta)}{X^2 - 2.\cos(\theta).X + 1} = \frac{1}{X - e^{i.\theta}} - \frac{1}{X - e^{-i.\theta}}.$$

○ Donner la définition de la trace d'une matrice carrée. Démontrer $Tr(A.B) = Tr(B.A)$ quand A et B sont deux matrices de formats compatibles.

Si A est carrée de taille n sur n de terme général a_i^k , sa trace est la somme $\sum_{i \leq n} a_i^i$ (somme des termes diagonaux).

Les matrices non carrées n'ont pas de trace.

On prend A ayant n lignes et q colonnes, et B ayant q lignes et n colonnes (A représente une application linéaire de $\mathbb{R}^q$ dans $\mathbb{R}^n$ et B une application de $\mathbb{R}^n$ dans $\mathbb{R}^q$).

Les deux produits $A.B$ et $B.A$ existent et donnent des matrices carrées ($A.B$ représente un endomorphisme de $\mathbb{R}^n$ et $B.A$ un endomorphisme de $\mathbb{R}^q$).

Le terme général de $A.B$ est $c_i^k = \sum_{j \leq q} a_i^j . b_j^k$ et sa trace est $\sum_{i \leq n} c_i^i = \sum_{i \leq n} \left(\sum_{j \leq q} a_i^j . b_j^i \right)$.

Le terme général de $B.A$ est $\gamma_j^k = \sum_{i \leq n} b_j^i . a_i^k$ et sa trace est $\sum_{j \leq q} \gamma_j^j = \sum_{j \leq q} \left(\sum_{i \leq n} b_j^i . a_i^j \right)$.

Ces deux sommes s'écrivent $\sum_{\substack{i \leq n \\ j \leq q}} a_i^j . b_j^i$ et sont égales.

Le cas particulier $Tr({}^t A.A)$ donne $\sum_{\substack{i \leq n \\ k \leq q}} (a_i^k)^2$.

Avec la relation $Tr(A.B) = Tr(B.A)$ appliquée à $Tr(P^{-1}.M.P) = Tr(P.P^{-1}.M)$, on montre que deux matrices carrées semblables ont la même trace ; la trace est un invariant de similitude.

Si A et B sont carrées de même format, on a aussi $\det(A.B) = \det(B.A)$.

En revanche dans le cas "matrice à n lignes et p colonnes face à matrice à n colonnes et p lignes", on n'a pas en général $\det(A.B) = \det(B.A)$. On montre même aisément en regardant les dimensions des images et noyaux qu'un des deux déterminants $\det(A.B)$ ou $\det(B.A)$ est nul.

○ Montrer que toute application dérivable en un point y est aussi continue. Contre-exemple pour l'absence de réciproque.

Si f est dérivable en a elle y admet un développement limité d'ordre 1 : $f(a+h) = f(a) + h.f'(a) + o(h)_{h \rightarrow 0}$

Quitte à le faire dégénérer en développement d'ordre 0, on obtient $f(a+h) \xrightarrow[h \rightarrow 0]{} f(a)$. On reconnaît la continuité de f en a (sans avoir le moindre besoin de distinguer "à droite/à gauche" comme dans un livre

de Terminale écrit avec les pieds et avec un Inspecteur dans le dos).

L'application $x \mapsto |x|$ est continue en 0 mais n'y est pas dérivable (une demi tangente de chaque côté, mais ce ne sont pas les mêmes).

L'application $x \mapsto \sqrt[3]{x}$ est continue en 0 mais n'y est pas dérivable (une tangente verticale).

On rappelle qu'il existe des applications continues partout, dérivables nulle part.

○ Calculer $1+j+j^2$, exprimer $1/j$ et $\bar{j}$ comme puissances de j . Donner la liste des racines de l'équation $z^6 = 1$ d'inconnue complexe z et les placer sur les cercle trigonométrique.

On rappelle qu'on a défini $j = e^{2.i.\pi/3}$. On a alors $1+j+j^2 = 0$ (factorisation $X^3 - 1 = (X-1).(X^2+X+1)$, somme des racines, série géométrique, simple calcul).

On a aussi $\frac{1}{j} = \bar{j} = j^2$ (complexes de module 1 conjugués, ou relation $j^3 = 1$).

On note que la géométrie dans le triangle équilatéral donne $1+j = e^{i.\pi/3}$.

Comme les racines de $X^6 - 1$ sont les $e^{i.k.\pi/3}$, on peut les écrire de plusieurs façons :

$1 = e^{0.i.\pi/3}$	$e^{i.\pi/3}$	$e^{2.i.\pi/3}$	$-1 = e^{3.i.\pi/3}$	$e^{4.i.\pi/3}$	$e^{5.i.\pi/3}$
1	$1+j = -j^2$	j	-1	j^2	$1+j^2 = -j$

Ce sont les six sommets d'un hexagone régulier inscrit dans le cercle trigonométrique.

○ Exprimer l'équation de la tangente en a au graphe d'une application dérivable.

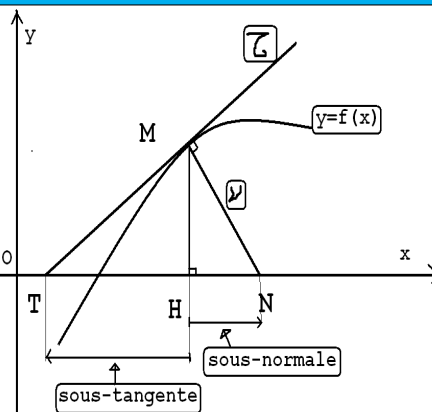
Le coefficient directeur est $f'(a)$, et l'ordonnée à l'origine est... non, ça on s'en fout, c'est vraiment un truc qui ne sert à rien en maths.

L'équation est directement $y = f'(a).(x - a) + f(a)$.

On l'écrit aussi géométriquement $\frac{y - f(a)}{x - a} = f'(a)$.

Les quantités mathématiquement parfois utiles sont alors la sous-tangente et la sous-normale (distance entre le pied H et l'intersection de la tangente avec l'axe Ox). Sans aucun calcul là encore mais par simple lecture de coefficient directeur :

$TH = \frac{f(a)}{f'(a)}$ et $HN = f(a).f'(a)$.



○ Donner plusieurs définitions de “ α est la borne supérieure de la partie A de $\mathbb{R}$ ”. Montrer l'unicité en cas d'existence.

La borne supérieure est le plus petit majorant de l'ensemble (si elle existe).

C'est un majorant	les autres majorants sont plus grands
$\forall a \in A, a \leq \alpha$	$\forall M \in \mathbb{R}, (\forall a \in A, a \leq M) \Rightarrow (\alpha \leq M)$

On prend deux nombres α et α qui se disent borne supérieure :

$\forall a \in A, a \leq \alpha$	$\forall M \in \mathbb{R}, (\forall a \in A, a \leq M) \Rightarrow (\alpha \leq M)$
$\forall a \in A, a \leq \alpha$	$\forall M \in \mathbb{R}, (\forall a \in A, a \leq M) \Rightarrow (\alpha \leq M)$

En croisant $\forall a \in A, a \leq \alpha$ et $\forall M \in \mathbb{R}, (\forall a \in A, a \leq M) \Rightarrow (\alpha \leq M)$ avec α dans le rôle de M , on trouve $\alpha \leq \alpha$.

En prenant l'autre croisement $\forall a \in A, a \leq \alpha$ et $\forall M \in \mathbb{R}, (\forall a \in A, a \leq M) \Rightarrow (\alpha \leq M)$, on trouve $\alpha \leq \alpha$.

Par antisymétrie α et α sont égaux.

On a d'autres quantifications :

C'est un majorant	un nombre plus petit n'est plus un majorant
$\forall a \in A, a \leq \alpha$	$\forall \mu \in \mathbb{R}, (\mu < \alpha) \Rightarrow (\exists a \in A, \mu < a)$
$\forall a \in A, a \leq \alpha$	$\forall \varepsilon > 0, \exists a \in A, \alpha - \varepsilon < a$

Si la partie A admet un plus grand élément (un maximum) alors c'est lui la borne supérieure (atteinte).

Une variante de quantification est aussi

α est un majorant de A	il existe une suite (a_n) d'éléments de A qui converge vers α
---------------------------------	--

Toute partie non vide majorée de $\mathbb{R}$ admet une borne supérieure.

Par convention naturelle, une partie non majorée a pour borne supérieure $+\infty$.

Par convention naturelle, l'ensemble vide a pour borne supérieure $-\infty$.

Toute suite réelle croissante majorée converge vers sa borne supérieure.

○ **Connaître la forme du $n^{ième}$ terme d'une suite u définie par $\forall n, u_{n+2} = a.u_{n+1} + b.u_n$ avec u_0 et u_1 donnés.**

La théorie dit tout de suite que l'espace des solutions est un espace vectoriel de dimension 2 engendré par des suites géométriques (sauf dans le cas de la racine double).

On écrit directement l'équation caractéristique $\lambda^2 = a.\lambda + b$.

On note λ_0 et λ_1 ses deux racines (éventuellement complexes).

Si ces deux valeurs propres sont distinctes, alors les suites sont des combinaisons de $((\lambda_0)^n)$ et de $((\lambda_1)^n)$.

Plus précisément : $\exists(A, B), \forall n, u_n = A.(\lambda_0)^n + B.(\lambda_1)^n$.

Les deux nombres A et B dépendent des conditions initiales et dictent le comportement à l'infini.

Si la suite est réelle mais que les valeurs propres λ_0 et λ_1 sont complexes conjuguées, alors A et B sont aussi complexes conjugués et $A.(\lambda_0)^n + \overline{A}.(\overline{\lambda_0})^n$ est toujours réel.

Si la racine est double, alors l'une des suites géométriques se dédouble et crée un terme polynômial : $\exists(A, B), \forall n, u_n = A.(\lambda_0)^n + B.n.(\lambda_0)^{n-1}$.

○ **Factoriser $a^2 - b^2$ par $a - b$. Factoriser $a^3 - b^3$ par $a - b$. Factoriser $a^n - b^n$ par $a - b$ (dans un anneau commutatif). Factoriser $a^3 + b^3$ et $a^5 + b^5$ par $a + b$.**

On connaît bien sûr $a^2 - b^2 = (a - b).(a + b)$ et aussi $a^3 - b^3 = (a - b).(a^2 + a.b + b^2)$.

La formule générale est celle de la série géométrique : $a^n - b^n = (a - b).(a^{n-1} + a^{n-2}.b + a^{n-3}.b^2 + \dots + b^{n-1})$.

On l'écrit aussi $a^n - b^n = (a - b). \left(\sum_{k=0}^n a^{n-k-1}.b^k \right) = (a - b). \left(\sum_{i+j=n-1} a^i.b^j \right)$.

On retient : que des signes +, et formule homogène.

On factorise en particulier : $X^n - 1 = (X - 1).(X^{n-1} + X^{n-2} + \dots + 1)$ avec le cas particulier $n = 3$.

En remplaçant b par $-b$ si l'exposant est négatif : $a^3 + b^3 = (a + b).(a^2 - a.b + b^2)$ et aussi $a^5 + b^5 = (a + b).(a^4 - a^3.b + a^2.b^2 - a.b^3 + b^4)$.

On retient l'alternance de signes, qui doit "retomber sur ses pattes" par imparité de l'exposant.

○ **Exprimer $\cos(\theta)$, $\sin(\theta)$, $\tan(\theta)$ à l'aide de $\tan(\theta/2)$ (avec domaine de définition). Expression $d\theta = \frac{2.dt}{1+t^2}$. Application aux intégrales.**

Pour que $\tan(\theta/2)$ existe, il faut éviter à θ d'être un multiple impair de π .

Le domaine naturel est $] -\pi, \pi[$.

On a immédiatement (et géométriquement) :

$\cos(\theta) = \frac{1-t^2}{1+t^2}$	$\sin(\theta) = \frac{2.t}{1+t^2}$	$\tan(\theta) = \frac{2.t}{1-t^2}$	$d\theta = \frac{2.dt}{1+t^2}$
--------------------------------------	------------------------------------	------------------------------------	--------------------------------

On peut retrouver rapidement ces formules : $\cos(\theta) = 2.\cos^2(\theta/2) - 1$ avec $\frac{1}{\cos^2(\theta/2)} = 1 + \tan^2(\theta/2)$,

de même $\sin(\theta) = 2.\sin(\theta/2).\cos(\theta/2) = 2.\tan(\theta/2).\cos^2(\theta/2)$.

Pour l'existence de $\tan(\theta)$ on doit éviter aussi les valeurs du type $\frac{\pi}{2} + k.\pi$ pour lesquelles t vaut 1 ou

-1, ce qui annule le dénominateur de $\frac{2.dt}{1-t^2}$.

On peut aussi écrire : $e^{i.\theta} = \frac{e^{i.\theta/2}}{e^{-i.\theta/2}} = \frac{\cos(\theta/2) + i.\sin(\theta/2)}{\cos(\theta/2) - i.\sin(\theta/2)} = \frac{1+i.t}{1-i.t}$. Il ne reste plus qu'à conjuguer et extraire partie réelle et partie imaginaire.

On dérive ensuite $t = \tan(\theta/2)$ par rapport à θ : $\frac{dt}{d\theta} = 1 + \tan^2(\theta/2)$, et on isole bien $dt = \frac{2.dt}{1+t^2}$.

On l'a aussi en écrivant $\theta = 2.Arctan(t)$ et en dérivant.

Ce sont ces formules qu'on utilise pour le changement de variable universel dans les intégrales du type

$\int_a^b R(\cos(\theta).\sin(\theta)).d\theta$. Il ne faut pas oublier alors de remplacer les bornes et l'élément différentiel.

La plus classique est $\int_a^b \frac{d\theta}{\sin(\theta)} = \ln \left(\frac{\tan(/b/2)}{\tan(a/2)} \right)$ quand l'intervalle $[a, b]$ ne contient aucun multiple

de π .

○ Montrez que l'ensemble des $e^{2.i.k.\pi/n}$ pour k de 0 à $n-1$ est un groupe commutatif.

C'est le groupe des solutions de l'équation $x^n = 1$ de racine complexe x . On peut constater que si a et b vérifient $a^n = 1$ et $b^n = 1$ alors on a bien $(a.b)^n = 1$ et $(a^{-1})^n = 1$.

D'autre part, ce groupe est isomorphe au groupe additif $\text{range}(n)$ pour l'addition modulo n .

○ Montrer que le graphe d'une application deux fois dérivable à dérivée seconde positive est au dessus de ses tangentes.

Il s'agit d'applications convexes, mais le programme n'utilise pas le mot.

Si l'application est de classe C^2 , on utilise la formule de Taylor avec reste intégrale :

$f(a+h) = f(a) + h.f'(a) + h^2 \cdot \int_0^1 (1-t).f''(a+th).dt$. Dans l'intégrale, tout est positif, devant l'intégrale aussi ; le terme correctif $+h^2 \cdot \int_0^1 (1-t).f''(a+th).dt$ est positif : $f(a+h) \geq f(a) + h.f'(a)$. On retraduit $f(x) \geq f(a) + f'(a).(x-a)$. C'est bien le graphe et la tangente au graphe en a .

Si l'application n'est que D^2 sans hypothèse de continuité de la dérivée seconde, on peut utiliser la formule de Taylor avec reste de Lagrange : $\exists \theta, f(a+h) = f(a) + h.f'(a) + \frac{h^2}{2}.f''(a+\theta.h)$. Le terme correctif est encore positif.

Mais cette formule de Lagrange n'est pas au programme, il faut donc utiliser une autre méthode.

On définit l'application différence : $f(x) - f(a) - (x-a).f'(a)$ que l'on note φ . On la dérive deux fois, sa dérivée seconde est positive. Par théorème des accroissements finis, la dérivée première est croissante, et ne peut changer qu'une fois de signe. Elle le fait justement en a . φ' est négative avant a et positive après. Par théorème des accroissements finis, φ est décroissante puis croissante. Comme elle est nulle en a , φ est positive partout.

Pour montrer une relation telle que $e^x \geq 1+x$ pour tout x , il suffit de parler de convexité.

○ Donner la définition d'une suite convergente. Montrer que la limite (si elle existe) est unique.

La définition est un classique à maîtriser par cœur :

$\forall \varepsilon > 0, \exists N_\varepsilon \in \mathbb{N}, \forall n \in \mathbb{N}, (n \geq N_\varepsilon \Rightarrow |a_n - \lambda| \leq \varepsilon)$

On croquera aussi $\forall \varepsilon > 0, \exists N_\varepsilon \in \mathbb{N}, \forall n \geq N_\varepsilon, |a_n - \lambda| \leq \varepsilon$.

On pourra aussi terminer par $|a_n - \lambda| < \varepsilon$ quitte à remplacer ε par sa moitié pour que l'inégalité devienne stricte.

On suppose que la suite a admet deux limites λ et μ :

$\forall \varepsilon > 0, \exists N_\varepsilon, \forall n \geq N_\varepsilon, |a_n - \lambda| \leq \varepsilon$ et $\forall \varepsilon > 0, \exists M_\varepsilon, \forall n \geq M_\varepsilon, |a_n - \mu| \leq \varepsilon$

Pour n plus grand que $\text{Max}(N_{\varepsilon/2}, M_{\varepsilon/2})$, on a à la fois $|a_n - \lambda| \leq \varepsilon/2$ et $|a_n - \mu| \leq \varepsilon/2$.

Par inégalité triangulaire : $|\lambda - \mu| \leq |\lambda - a_n| + |a_n - \mu| \leq \varepsilon$.

Le réel positif $|\lambda - \mu|$ est petit que n'importe que ε . La seule solution est qu'il soit nul.

Sinon, prendre le cas particulier $\varepsilon = |\lambda - \mu|$ et constater la contradiction.

○ Calculer $\sum_{k=0}^n k$, $\sum_{k=0}^n k^2$ et $\sum_{k=0}^n k^3$. Connaître une méthode pour calculer de proche en proche les sommes $\sum_{k=0}^n k^p$.

On donne directement

$\sum_{k=0}^n k^0$	$\sum_{k=0}^n k^1$	$\sum_{k=0}^n k^2$	$\sum_{k=0}^n k^3$
$= n+1$	$= \frac{n.(n+1)}{2}$	$= \frac{n.(n+1).(2.n+1)}{6}$	$= \left(\frac{n.(n+1)}{2}\right)^2$

La première est un simple compteur.

La seconde est visuelle.

La troisième se prouve par récurrence.

La dernière s'obtient par télescopage de $\left(\frac{k.(k+1)}{2}\right)^2 - \left(\frac{(k-1).k}{2}\right)^2$ ou par un découpage judicieux.

La somme $\sum_{k=0}^n k^p$ est un polynôme en n de degré $p+1$ à coefficients rationnels, de terme dominant

$\frac{n^{p+1}}{p+1}$ (qu'on obtient par comparaison série intégrale et comportement à l'infini).

En télescopant la formule $\binom{k+1}{p+1} - \binom{k}{p+1} = \binom{k}{p}$, on a la formule de Zhu-Shi-Jie :

$$\sum_{k=0}^n \binom{k}{p} = \binom{n+1}{p+1}.$$

On l'écrit avec la définition des coefficients binomiaux (valable aussi hors triangle) :

$$\sum_{k=0}^n \frac{k.(k-1)\dots(k-p+1)}{p!} = \frac{(k+1).k.(k-1)\dots(k-p+1)}{(p+1)!}.$$

On développe le premier membre, et on isole le terme $\sum_{k=0}^n k^{p+1}$, face à des termes s'exprimant à l'aide

des $\sum_{k=0}^n k^q$ avec q plus petit que p . On a exprimé une somme à l'aide des précédentes.

○ Montrer qu'une somme de réels positifs est nulle si et seulement si chaque réel est nul.

Un seul sens est digne d'intérêt : si la somme est nulle avec des termes tous positifs, alors forcément chaque terme est nul.

On part de $\sum_{k=1}^n a_k = 0$ avec des a_k tous réels positifs (ce peut être classiquement $\sum_{k=1}^n (\alpha_k)^2 = 0$).

On est tenté de montrer par l'absurde que chaque a_k est nul.

Mais il y a plus simple : $0 \leq a_1 \leq a_1 + \sum_{k=2}^n a_k = 0$, il n'y a plus qu'à conclure par antisymétrie.

○ Montrer qu'une suite complexe converge si et seulement si sa partie réelle et sa partie imaginaire convergent.

On suppose que la suite (z_n) (d'écriture cartésienne $z_n = x_n + i.y_n$) converge vers λ (d'écriture cartésienne $\alpha + i.\beta$). On traduit :

$$\forall \varepsilon > 0, \exists N_\varepsilon, \forall n, n \geq N_\varepsilon \Rightarrow |z_n - \lambda| \leq \varepsilon$$

En utilisant les relations classiques $|\Re(z)| \leq |z|$ et $|\Im(z)| \leq |z|$, on obtient

$$\forall \varepsilon > 0, \exists N_\varepsilon, \forall n \geq N_\varepsilon, |x_n - \alpha| \leq \varepsilon \text{ et } \forall \varepsilon > 0, \exists N_\varepsilon, \forall n \geq N_\varepsilon, |y_n - \beta| \leq \varepsilon$$

On reconnaît la convergence de (x_n) vers α et la convergence de (y_n) vers β .

Réciproquement, si on a

$$\forall \varepsilon > 0, \exists R_\varepsilon, \forall n \geq R_\varepsilon, |x_n - \alpha| \leq \varepsilon \text{ et } \forall \varepsilon > 0, \exists I_\varepsilon, \forall n \geq I_\varepsilon, |y_n - \beta| \leq \varepsilon,$$

alors pour ε donné, on a, à partir du rang $\max(R_{\varepsilon/2}, I_{\varepsilon/2})$: $|x_n - \alpha| \leq \varepsilon/2$ et $|y_n - \beta| \leq \varepsilon/2$, puis par inégalité triangulaire : $|x_n + i.y_n - \alpha - i.\beta| \leq \varepsilon$. C'est la quantification de la convergence de $(x_n + i.y_n)$ vers $\alpha + i.\beta$.

Attention, la convergence de la suite (z_n) n'entraîne pas la convergence de son module et de son argument.

Si (z_n) converge vers 0, son module tend vers 0, mais son argument peut faire ce qu'il veut.

D'autre part, l'argument est toujours défini de manière ambiguë : la suite $-1 + (-2)^n.i$ converge vers -1 avec un argument qui hésite entre π et $-\pi$.

○ Calculer géométriquement $\int_a^b (\alpha.t + \beta).dt$ et $\int_0^1 \sqrt{1-x^2}.dx$.

La première intégrale se calcule en recourant si on veut aux primitives. Mais comme le calcul d'intégrale (c'est à dire d'aire) par variation de primitive (opération réciproque de la dérivation) nécessite de gros théorèmes de l'analyse, autant revenir à la définition.

C'est l'aire comprise sous une portion de droite. C'est donc l'aire d'un trapèze droit. Sa base a pour longueur $b - a$ et ses deux hauteurs valent $\alpha.a + \beta$ et $\alpha.b + \beta$. Géométriquement, on a (base) fois (moyenne des hauteurs), ce qui donne $\alpha \cdot \frac{(b-a).(b+a)}{2} + \beta.(b-a)$.

Pour $\int_0^1 \sqrt{1-x^2}.dx$, il faut reconnaître dans $y = \sqrt{1-x^2}$ l'équation d'un demi cercle, de centre (0,0)

et de rayon 1 (pour le cercle : $x^2 + y^2 = 1$). L'aire de la portion de disque vaut alors $\frac{\pi \cdot 1^2}{4}$.

Le résultat peut se généraliser à $\int_0^a \sqrt{1-x^2} \cdot dx$ en découpant en portion de disque et triangle. On trouve plus laborieusement $\frac{a \cdot \sqrt{1-a^2}}{2} - a \cdot \text{Arcsin}(a)$.

On peut aussi intégrer par parties :

1	$\leftrightarrow$	x
$\sqrt{1-x^2}$	$\hookrightarrow$	$-\frac{x}{\sqrt{1-x^2}}$

, et séparer le terme $\int_0^a \frac{x^2}{\sqrt{1-x^2}} \cdot dx$

en $\int_0^a \frac{x^2-1}{\sqrt{1-x^2}} \cdot dx + \int_0^a \frac{1}{\sqrt{1-x^2}} \cdot dx$ qui ait revenir l'intégrale elle-même (mais avec un signe moins qui devient un plus de l'autre côté) et un terme à intégrer en Arcsin . Obtenir $I = \dots + \alpha \cdot I$ en intégrant par parties est classique... en Terminale surtout.

On peut aussi changer de variable : $x = \sin(\theta)$ et l'intégrale devient $\int_0^{\text{Arcsin}(a)} \cos^2(t) \cdot dt$, facile à calculer. Il restera à remplacer $\cos(\text{Arcsin}(a))$ par $\sqrt{1-a^2}$.

On rappelle au passage :

cercle	courbe	le bord	$2 \cdot \pi \cdot R$	et	sphère	surface	le bord	$4 \cdot \pi \cdot R^2$
disque	surface	l'intérieur	$\pi \cdot R^2$		boule	volume	l'intérieur	$4 \cdot \pi \cdot R^3 / 3$

○ Montrer qu'un polynôme de degré n a au plus n racines.

Ce résultat est évident et n'a rien à voir avec le puissant théorème de d'Alembert-Gauss qui lui, garantit l'existence de racines.

On prend un polynôme P de degré n : $P = \sum_{k=0}^n a_k \cdot X^k$ avec a_n non nul.

Pour chaque racine α , P se factorise par $X - \alpha$.

Si P admet plus de n racines α_1 à α_N , alors il se factorise $P(X) = a_n \cdot \prod_{k=1}^N (X - \alpha_k)$ et il y a une contradiction sur les degrés.

On notera que ce résultat n'est valable que dans un anneau commutatif, afin de pouvoir pratiquer la division euclidienne de polynômes. Dans le corps (hors-programme) $(\mathbb{H}, +, \times)$ des quaternions de Hamilton, le polynôme $X^2 + 1$ a quand même beaucoup de racines...

○ Donner la définition de $n!$ pour n entier naturel. Connaître éventuellement l'équivalent de $n!$ quand n tend vers l'infini (formule de Stirling, de démonstration hors programme).

Pour n dans $\mathbb{N}$, on définit de manière récursive : $n! = n \cdot (n-1)!$ avec $0! = 1$.

Ou alors directement : $n! = \prod_{k=1}^n k$ avec la convention qu'un produit vide vaut 1.

La formule $0! = 1$ n'est même pas une convention, elle est logique en termes de dénombrement de permutations d'une liste vide.

On pourrait prendre aussi $n! = \int_0^{+\infty} x^n \cdot e^{-x} \cdot dx = \int_0^1 |\ln(t)|^n \cdot dt$ et étendre la définition à $\mathbb{R}^+$.

On pourra prendre pour convention que les entiers négatifs ont une factorielle infinie.

On montre en introduisant une suite ad-hoc et les intégrales de Wallis : $n! \sim_{n \rightarrow +\infty} \left(\frac{n}{e}\right)^n \cdot \sqrt{2 \cdot n \cdot \pi}$. C'est cette formule qui porte le nom de Stirling.

En revanche, il peut être demandé par une comparaison série intégrale un encadrement de $\ln(n!)$ par $\int_0^n \ln(t) \cdot dt$ et $\int_1^{n+1} \ln(t) \cdot dt$.

○ Montrer que si une suite strictement positive u est telle que $\frac{u_{n+1}}{u_n}$ converge vers un réel α strictement plus petit que 1, alors la suite (u_n) converge vers 0.

Intuitivement, la suite (u_n) est presque géométrique de raison α , elle fait donc comme la suite géométrique de raison plus petite que 1 en valeur absolue : elle converge vers 0.

Comme la suite $\left(\frac{u_{n+1}}{u_n}\right)$ converge vers a strictement plus petit que 1, alors à partir d'un rang $N_{(1-a)/2}$, ses termes sont entre $a - \frac{1-a}{2}$ et $a + \frac{1-a}{2}$. La majoration $\frac{u_{n+1}}{u_n} \leq \frac{1+a}{2} < 1$ dit alors que la suite u est décroissante à partir du rang N . Comme elle est minorée par 0, elle converge.

On note λ sa limite. Si λ est non nulle, alors, par passage à la limite dans $\frac{u_{n+1}}{u_n} \rightarrow a$, on a $\frac{\lambda}{\lambda} = a$, ce qui est contradictoire. Par élimination, la suite converge vers 0.

On peut aussi mettre en boucle la majoration $\frac{u_{n+1}}{u_n} \leq \frac{1+a}{2}$ valable à partir du rang N et obtenir $0 \leq a_n \leq u_N \cdot \left(\frac{1+a}{2}\right)^{n-N}$. On fait tendre n vers l'infini, le membre de droite tend vers 0, par encadrement, le terme a_n tend vers 0 aussi.

Ce résultat porte le nom de comparaison logarithmique, et ceci peut se comprendre aussi en utilisant la suite $\left(\ln\left(\frac{u_{n+1}}{u_n}\right)\right)$. Elle converge vers $\ln(a)$. Par théorème (qui n'est qu'au programme de seconde année), sa moyenne de Césaro converge aussi vers $\ln(a)$. Mais cette moyenne de Césaro télescope en $\frac{\ln(u_n) - \ln(u_0)}{n}$. Par produit, $\ln(u_n)$ est équivalente à $n \cdot \ln(a)$. Comme a est entre 0 et 1, la suite $(\ln(u_n))$ converge vers $-\infty$ et la suite (u_n) converge vers 0. Il faudra traiter à part le cas $a = 0$ dans cette preuve.

C'est avec ce théorème qu'on démontre la plupart des relations de comparaison (croissances comparées) sur les suites réelles infiniment grandes :

$\ln(n)$	n^p	n^q	a^n	b^n	$n!$	n^n
	polynômes	$0 < p < q$	géométriques	$1 < a < b$		

○ Montrer que toute suite convergente est bornée.

On écrit la quantification de la convergence de la suite (a_n) vers α :

$$\forall \varepsilon > 0, \exists N_\varepsilon \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq N_\varepsilon \Rightarrow |a_n - \alpha| \leq \varepsilon$$

et on l'applique au cas particulier $\varepsilon = 1$ (ou toute autre valeur strictement positive). On a donc pour n plus grand que N_1 :

$$|a_n| \leq |a_n - \alpha| + |\alpha| \leq 1 + |\alpha|$$

La suite est bornée à partir du rang N_1 ; avant, il n'y a qu'un nombre fini de termes. Globalement, la suite est bornée (par $\max(1 + |\alpha|, |a_0|, \dots, |a_{N_1-1}|)$).

Ce raisonnement est valable dans un espace métrique, c'est à dire muni d'une distance, et pas seulement dans $\mathbb{R}$.

En revanche, des raisonnements du même type, avec choix d'un ε particulier ne sont valables que dans $\mathbb{R}$ (si une suite réelle converge vers une limite strictement positive λ , alors elle est strictement plus grande que $\lambda/2$ à partir d'un certain rang ; toute suite réelle de limite $+\infty$ est minorée...).

○ Donner la définition de $P(E)$ si E est un ensemble.

$P(E)$ est à son tour un ensemble ("un étage plus haut") dont les éléments sont les parties (ou sous-ensembles) de E .

$$\boxed{A \subset E} \mid \boxed{A \in P(E)} \text{ et en particulier } \boxed{a \in E} \mid \boxed{\{a\} \subset E} \mid \boxed{\{a\} \in P(E)}$$

On a toujours $E \in P(E)$ et $\emptyset \in P(E)$.

Et on a aussi $\emptyset \subset P(E)$ puisque $P(E)$ est un ensemble.

	$a \in E$	$A \subset E$
On évitera la locution ambiguë "est dans" :	a est un élément de E	A est une partie de E
	a appartient à E	A est inclus dans E

Avec des patates (plus proprement "diagrammes de Venn ou de Lewis Carroll), pour A, B et C parties de E :

$A \Delta B \in P(E)$	internes	$A \cap B \in P(E)$
$A \Delta (B \Delta C) = (A \Delta B) \Delta C$	associatives	$A \cap (B \cap C) = (A \cap B) \cap C$
$A \Delta B = B \Delta A$	commutatives	$A \cap B = B \cap A$
$A \Delta \emptyset = A$	neutres	$A \cap E = A$
$A \Delta A = \emptyset$	symétriques pour Δ	
	distributivité de $\cap$ sur Δ	$A \cap (B \Delta C) = (A \cap B) \Delta (A \cap C)$

Les démonstrations se font plus proprement en passant par les indicatrices, qui créent un isomorphisme entre $(P(E), \Delta, \cap)$ est $(\{0, 1\}^n, +_{mod 2}, \times)$.

La relation $A \cap B = \emptyset$ entraîne pas forcément $A = \emptyset$ ou $B = \emptyset$ dès lors que le cardinal de E atteint ou dépasse 2 ; l'anneau n'est pas intègre.

○ Donner le polynôme P de degré inférieur ou égal à 3 vérifiant $P(a) = \alpha$, $P(b) = \beta$ et $P(c) = \gamma$.

On peut évidemment tel un obéissant et gentillet élève de Terminale poser $P(X) = \lambda + \mu.X + \nu.X^2$ et

résoudre un système de trois équations à trois inconnues $\begin{cases} \lambda + \mu.a + \nu.a^2 = \alpha \\ \lambda + \mu.b + \nu.b^2 = \beta \\ \lambda + \mu.c + \nu.c^2 = \gamma \end{cases}$. Son détermi-

nant est $(c-a).(c-b).(b-a)$. Il est non nul, il y a une unique solution. Mais ce serait s'embourber définitivement dans le mauvais chemin : croire qu'un polynôme ne s'écrit que sur la base canonique et que les mathématiques sont des résolutions d'équations.

On crée d'abord trois polynômes parfaitement adaptés chacun à une partie de notre question :

polynôme	$\frac{(X-b).(X-c)}{(a-b).(a-c)}$	$\frac{(X-a).(X-c)}{(b-a).(b-c)}$	$\frac{(X-a).(X-b)}{(c-a).(c-b)}$
valeur en a	1	0	0
valeur en b	0	1	0
valeur en c	0	0	1

Il ne reste plus qu'à les combiner : $\alpha \cdot \frac{(X-b).(X-c)}{(a-b).(a-c)} + \beta \cdot \frac{(X-a).(X-c)}{(b-a).(b-c)} + \gamma \cdot \frac{(X-a).(X-b)}{(c-a).(c-b)}$.

Ce polynôme est de degré inférieur ou égal à 2 (dans quel cas est il de degré 1 ? c'est là un exercice intéressant). Et il vérifie exactement les trois conditions.

Si nécessaire, on montre que c'est le seul. On peut invoquer le fait que le système écrit plus haut admet une unique solution, par inversibilité de la matrice.

On peut aussi prendre deux solutions P et Q ; le polynôme $P - Q$ est alors de degré inférieur ou égal à 2 et est nul en trois points (a , b et c) ; il est nul.

Si on veut tous les polynômes solutions, on prend les

$\alpha \cdot \frac{(X-b).(X-c)}{(a-b).(a-c)} + \beta \cdot \frac{(X-a).(X-c)}{(b-a).(b-c)} + \gamma \cdot \frac{(X-a).(X-b)}{(c-a).(c-b)} + (X-a).(X-b).(X-c).T(X)$ avec T quelconque.

Le phénomène se généralise à de plus grandes listes "abscisse/ordonnée". Ce sont les polynômes interpolateurs de Lagrange.

Dans le cas plus simple où on cherche une fonction affine P vérifiant $P(a) = \alpha$ et $P(b) = \beta$, on peut proposer directement suivant son approche

$\alpha \cdot \frac{X-b}{a-b} + \beta \cdot \frac{X-a}{b-a}$	$\frac{\beta-\alpha}{b-a} \cdot (X-a) + \alpha$	$\frac{\beta.(X-a) + \alpha.(b-X)}{b-a}$
interpolateurs de Lagrange	coefficient directeur	barycentre

○ Montrer : $e^t \geq 1 + t$ pour tout t .

C'est évidemment une inégalité de convexité. On peut se contenter que la dérivée de l'exponentielle est croissante, le graphe est donc au dessus des tangentes. C'est ici juste une comparaison en 0.

On peut aussi étudier l'application différence : $x \mapsto e^x - 1 - x$. Elle se dérive en $x \mapsto e^x - 1$. La dérivée est d'abord négative puis positive. L'application est décroissante puis croissante. Elle admet un minimum en 0 et ce minimum est nul.

On peut aussi écrire $e^t - 1 = \int_0^t e^x dx$ et étudier les deux cas.

Si t est positif, les e^x sous le signe somme sont plus grands que 1, l'intégrale est plus grande que $\int_0^t dx$ de valeur t .

Si t est négatif, on majore e^x par 1, mais on prend garde au sens de l'intervalle d'intégration.

○ Exprimer $Card(A \cup B)$ à l'aide de $Card(A)$, $Card(B)$ et $Card(A \cap B)$.

On a toujours $Card(A \cup B) = Card(A) + Card(B) - Card(A \cap B)$.

On compte les éléments de A , ceux de B et on décompte ceux qu'on a comptés deux fois.

On prouve aussi cette formule en écrivant $1_{A \cup B} = 1_A + 1_B - 1_{A \cap B}$ et en sommant en faisant défiler tous les éléments de E .

On dispose de la généralisation trois ensembles

$$Card(A \cup B \cup C) = \begin{matrix} Card(A) + & -Card(B \cap C) \\ +Card(B) & -Card(C \cap A) & +Card(A \cap B \cap C) \\ +Card(A) & -Card(A \cap CB) \end{matrix} \quad \text{et d'une formule encore}$$

$$\text{plus générale } \sum_k (-1)^k \cdot \sum_{i_1 < i_2 < \dots < i_k} Card(A_{i_1} \cap \dots \cap A_{i_k}).$$

Le cardinal d'un ensemble A sera noté suivant le cadre $Card(A)$, $|A|$ ou $\#A$, et s'appelle même "ordre" quand l'ensemble A est un groupe.

○ Donner la définition de $\binom{n}{k}$. Savoir démontrer la relation de Pascal : $\binom{n}{k} + \binom{n}{k+1} = \binom{n+1}{k+1}$. Prouver que $\binom{n}{k}$ est toujours un entier. Calculer le proche en proche les coefficients le long d'une ligne, et éventuellement le long d'une colonne.

Pour n et k entiers, avec k plus petit que n , on pose $\binom{n}{k} = \frac{n!}{k!(n-k)!}$. Pour n négatif ou k plus grand que n , on pose $\binom{n}{k} = 0$ (coefficients binomiaux aberrants), que l'on prenne la définition par dénombrement, ou en tenant compte du fait que $p!$ est infini pour p négatif.

On retiendra plutôt la forme $\binom{n}{k} = \frac{n(n-1)\dots(n-k+1)}{1.2\dots k}$ dans laquelle il y a k termes au numérateur et autant au dénominateur (ne pas retenir $n-k+1$ mais compter les termes). C'est ainsi qu'on a $\binom{n}{2} = \frac{n(n-1)}{1.2}$ et $\binom{n}{3} = \frac{n(n-1)(n-2)}{1.2.3}$ plutôt que $\frac{n!}{2!(n-2)!}$ et $\frac{n!}{3!(n-3)!}$.

C'est sous cette forme qu'on retrouve que $\binom{n}{k}$ dénombre les sous-ensembles à k éléments issus d'un ensemble à n éléments. Au numérateur, on a n choix pour le premier élément, $n-1$ choix pour le second (différent du premier) jusqu'à $n-(k-1)$ choix pour le k^{ieme} (distinct des $k-1$ déjà tirés). On divise par le nombre de façons d'ordonner la liste des k éléments distincts obtenue.

La relation de Pascal se démontre avec la définition factorielle en mettant en commun les bonnes factorielles :

$$\begin{aligned} \binom{n}{k} + \binom{n}{k+1} &= \frac{n!}{k!(n-k)!} + \frac{n!}{(k+1)!(n-k-1)!} \\ \binom{n}{k} + \binom{n}{k+1} &= \frac{n!}{(k+1)!(n-k)!} \cdot ((k+1) + (n-k)) = \frac{n!(n+1)}{k!((n+1)-(k+1))!} \end{aligned}$$

On peut aussi, dans un ensemble $\{a_0, \dots, a_n\}$ à $n+1$ éléments dénombrer de deux façons les parties à $k+1$ éléments contenant a_0 :

il y a $\binom{n+1}{k+1}$ parties à $k+1$ éléments, mais il faut décompter les parties ne contenant pas a_0 ; ce

sont des parties à $k+1$ éléments parmi n seulement : on aboutit à $\binom{n+1}{k+1} - \binom{n}{k+1}$

mais on peut aussi rendre des parties à k éléments issus de $\{a_1, \dots, a_n\}$ et y ajouter a_0 : il y en a $\binom{n}{k}$

On notera que la relation de Pascal est vraie aussi sur le bord du triangle : $\binom{n}{n} + \binom{n}{n+1} = \binom{n+1}{n+1}$.

Déjà, les coefficients binomiaux sont des rationnels. La propriété que l'on démontre ensuite par récurrence sur n est définie par $H_n : \forall k \in \mathbb{N}, \binom{n}{k} \in \mathbb{N}$.

Elle s'initialise avec $n = 0$: les $\binom{0}{k}$ valent 0 ou 1 et sont entiers.

On suppose que pour un n donné, les $\binom{n}{k}$ sont tous entiers. On se place sur la ligne d'indice $n + 1$ et on écrit pour tout k : $\binom{n+1}{k} = \binom{n}{k} + \binom{n}{k-1}$. En tant que somme d'entiers naturels, c'est toujours un entier naturel.

La même relation $\binom{n+1}{k+1} - \binom{n}{k+1} = \binom{n}{k}$ permet par télescopage

$$\text{d'obtenir } \sum_{n=0}^N \binom{n}{k} = \binom{N+1}{k+1}.$$

Pour calculer les premiers coefficients binomiaux, on utilise la relation

$$\binom{n}{k-1} + \binom{n}{k} = \binom{n+1}{k}.$$

1					
1	1				
1	2	1			
1	3	3	1		
1	4	6	4	1	
1	5	10	10	5	1
1	6	15	20	15	6

Si on doit juste calculer les coefficients d'une ligne (par exemple pour développer $(a+b)^{11}$), on initialise avec $\binom{n}{0} = 1$ et on avance avec la formule $\binom{n}{k+1} = \frac{n-k}{k+1} \cdot \binom{n}{k}$:

multiplication		n		$n-1$		$n-2$		$n-3$	
coefficient	1		n		$\frac{n \cdot (n-1)}{2}$		$\frac{n \cdot (n-1) \cdot (n-2)}{6}$		$\frac{n \cdot (n-1) \cdot (n-2) \cdot (n-3)}{24}$
division		1		2		3		4	

On retrouve que la suite $\left(\binom{n}{k} \right)_{k \in \mathbb{N}}$ est d'abord croissante, puis décroissante, avec un maximum "en milieu de ligne".

En colonne : $\binom{n+1}{k} = \frac{n+1}{n-k+1} \cdot \binom{n}{k}$ et en diagonale $\binom{n+1}{k+1} = \frac{n+1}{k+1} \cdot \binom{n}{k}$:

$\binom{n}{k}$	$\rightarrow \times \frac{n-k}{k+1}$	$\binom{n}{k+1}$
$\downarrow \times \frac{n+1}{n-k+1}$	$\searrow \times \frac{n+1}{k+1}$	
$\binom{n+1}{k}$		$\binom{n+1}{k+1}$

○ Calculer combien il y a d'applications de E dans F quand E et F sont deux ensembles finis. Calculer le nombre d'applications injectives.

On suppose que E est formé de n éléments $e_1, \dots, e_n$ et que F est formé de p éléments $f_1, \dots, f_p$. Pour définir une application, il faut et il suffit de choisir l'image de e_1 (p choix), puis l'image de e_2 (p choix encore) et ainsi de suite. Le produit $p \times p \times \dots \times p$ se simplifie en p^n .

On a donc $\text{Card}(F)^{\text{Card}(E)}$ applications de E dans F (arrivée^{départ}).

On note $\text{Appl}(E, F)$ ou $F(E, F)$ ou F^E l'ensemble des applications de E dans F .

Cette notation se justifie par le dénombrement indiqué plus haut.

Elle se justifie aussi par le fait que F^n (ensemble des n uplets d'éléments $(x_1, \dots, x_n)$ de F) est aussi $F^{\{1, \dots, n\}}$.

La relation $\text{Card}(F^E) = \text{Card}(F)^{\text{Card}(E)}$ est aussi valable pour des ensembles infinis, mais l'interprétation est délicate.

Pour les applications injectives, le choix des images diminue au fur et à mesure :

p choix pour l'image de e_1 , $p-1$ choix pour l'image de e_2 et ainsi de suite, jusqu'à $p \cdot (p-1) \dots (p-n+1)$

formé de n entiers consécutifs.

Cette formule est cohérente pour n plus grand que p : il n'y a plus d'injections, et il y a un 0 dans le produit.

La formule compacte est $\frac{p!}{(p-n)!}$, appelée nombre d'arrangements.

○ **Démontrer la relation $\text{Arcsin}(x) + \text{Arccos}(x) = \frac{\pi}{2}$ pour tout x de $[-1, 1]$. Démontrer la formule $\text{Arctan}(x) + \text{Arctan}\left(\frac{1}{x}\right) = \frac{\pi}{2}$ pour x réel strictement positif.**

On se donne x entre -1 et 1 et on pose $\theta = \text{Arcsin}(x)$. On traduit : $\sin(\theta) = x$ et $-\pi/2 \leq \theta \leq \pi/2$. On a alors : $\cos\left(\frac{\pi}{2} - \theta\right) = \sin(\theta) = x$ et aussi : $0 \leq \frac{\pi}{2} - \theta \leq \pi$. On reconnaît : $\frac{\pi}{2} - \theta = \text{Arccos}(x)$. On a finalement : $\frac{\pi}{2} - \text{Arcsin}(x) = \text{Arccos}(x)$.

On prend x strictement positif et on définit $\alpha = \text{Arctan}(x)$. On traduit : $\tan(\theta) = x$ et $0 < \theta < \pi/2$. On a immédiatement : $\tan\left(\frac{\pi}{2} - \theta\right) = \frac{1}{\tan(\theta)} = \frac{1}{x}$ (échange du sinus et du cosinus) et encore $0 < \frac{\pi}{2} - \theta < \frac{\pi}{2}$. On reconnaît : $\frac{\pi}{2} - \text{Arctan}(x) = \text{Arctan}\left(\frac{1}{x}\right)$.

Attention, pour x négatif, on a $\text{Arctan}(x) + \text{Arctan}\left(\frac{1}{x}\right) = -\frac{\pi}{2}$ par imparité.

On résume en $\text{Arctan}(x) + \text{Arctan}\left(\frac{1}{x}\right) = \text{Sgn}(x) \cdot \frac{\pi}{2}$ ou même $\frac{|x|}{x} \cdot \frac{\pi}{2}$.

○ **Donner le développement limité en 0 de $\frac{1}{1+h}$ (ordre n).**

On peut évidemment l'obtenir en calculant les dérivées successives en 0 et utiliser la formule de Taylor Young. Mais il est plus judicieux de penser à la série géométrique :

$$\frac{1}{1+h} = \frac{1 - (-h)^{n+1}}{1+h} + h^n \cdot \frac{(-1)^{n+1} \cdot h}{1+h} = 1 - h + h^2 \dots + (-h)^n + h^n \cdot \frac{(-1)^{n+1} \cdot h}{1+h}.$$

Le terme $h^n \cdot \frac{(-1)^{n+1} \cdot h}{1+h}$ s'écrit bien $o(h^n)$ quand h tend vers 0.

On identifie le développement limité cherché : $\frac{1}{1+h} = \sum_{k=0}^n (-h)^k \cdot o(h^n)_{h \rightarrow 0}$.

La formule pour $\frac{1}{1-h}$ est encore plus simple.

Si on en a besoin, on a quand même $\text{dsp}\left(t \mapsto \frac{1}{t}\right)^{(n)} = \left(t \mapsto \frac{(-1)^n \cdot n!}{t^{n+1}}\right)$ qu'on obtient par récurrence en écrivant $(-1)^n \cdot n! \cdot t^{-n-1}$ pour dériver facilement.

○ **Prouver qu'une application dérivable croissante sur un intervalle a une dérivée positive.**

C'est le sens naturel qui ne demande aucun effort.

Si l'application est croissante ($y > x \Rightarrow f(y) \geq f(x)$), alors tous les taux d'accroissement $\frac{f(x) - f(a)}{x - a}$ sont positifs. Par passage à la limite (l'hypothèse dit qu'elle existe), la limite $f'(a)$ est positive ou nulle.

L'application peut être strictement croissante et avoir une dérivée nulle en certains points, comme par exemple $x \mapsto x^3$. D'autre part, il existe des applications strictement croissantes n'ayant pas une dérivée positive (il suffit pour cela qu'elles ne soient pas dérivables, ni même d'ailleurs continues).

○ **Donner la dérivée $n^{\text{ième}}$ du produit de deux fonctions n fois dérivables (formule de Leibniz).**

On rappelle déjà la formule pour la dérivée première : $(u.v)' = u'.v + u.v'$ obtenue par exemple par limite des taux d'accroissement : $\frac{u(x).v(x) - u(a).v(a)}{x - a} = u(x) \cdot \frac{v(x) - v(a)}{x - a} + v(a) \cdot \frac{u(x) - u(a)}{x - a}$ ou par produit des deux développements limités d'ordre 1 :

$$u(a+h).v(a+h) = (u(a) + h.u'(a) + o(h)).(v(a) + h.v'(a) + o(h)).$$

On démontre par récurrence sur n la formule qu'on vient d'initialiser $(u.v)^{(n)} = \sum_{k=0}^n \binom{n}{k} \cdot u^{(n-k)} \cdot v^{(k)}$

dans laquelle u et v jouent bien des rôles symétriques.

On suppose pour un n donné que la formule est correcte au rang n , et on redérive car tout est encore au moins une fois dérivable. Par linéarité de la dérivation $(u.v)^{(n+1)} = \sum_{k=0}^n \binom{n}{k} \cdot (u^{(n-k)} \cdot v^{(k)})'$.

Par la formule de dérivation d'un produit déjà démontrée lors de l'initialisation :

$$(u.v)^{(n+1)} = \sum_{k=0}^n \binom{n}{k} \cdot u^{(n-k+1)} \cdot v^{(k)} + \sum_{k=0}^n \binom{n}{k} \cdot u^{(n-k+1)} \cdot v^{(k+1)}$$

On décale les indices dans la seconde :

$$(u.v)^{(n+1)} = \sum_{k=0}^n \binom{n}{k} \cdot u^{(n-k+1)} \cdot v^{(k)} + \sum_{k'=1}^{n+1} \binom{n}{k'-1} \cdot u^{(n-k'+1)} \cdot v^{(k')}$$

On isole les termes extrêmes et on fusionne les parties communes de la somme :

$$(u.v)^{(n+1)} = \binom{n}{0} \cdot u^{(n+1)} \cdot v^{(0)} + \sum_{k=1}^n \left(\binom{n}{k} + \binom{n}{k-1} \right) \cdot u^{(n-k+1)} \cdot v^{(k)} + \binom{n}{n} \cdot u^{(0)} \cdot v^{(n+1)}$$

On rappelle à ce stade de la rédaction que $u^{(0)}$ signifie simplement u . On utilise la relation de Pascal

$$\binom{n}{k} + \binom{n}{k-1} = \binom{n+1}{k} \text{ et on incorpore les termes extrêmes qui vont correspondre aux indices}$$

$k=0$ et $k=n+1$ (les coefficients binomiaux en jeu valent 1, qu'on les écrive $\binom{n}{0}$ ou $\binom{n+1}{0}$).

$$\text{On aboutit bien à } (u.v)^{(n+1)} = \sum_{k=0}^{n+1} \binom{n+1}{k} \cdot u^{(n+1-k)} \cdot v^{(k)}.$$

On pouvait aussi pour l'hérédité partir de $(u.v)' = u' \cdot v + u \cdot v'$ et dériver n fois en appliquant la formule de rang n au couple (u', v) et au couple (u, v') . Le regroupement se fait alors quand même sur le même principe.

On peut aussi se placer en un point et écrire deux formules de Taylor Young (usage licite car u et v sont de classe D^n) : $u(a+h) = \sum_{i=0}^n \frac{u^{(i)}(a)}{i!} \cdot h^i + o(h^n)$ et $v(a+h) = \sum_{j=0}^n \frac{v^{(j)}(a)}{j!} \cdot h^j + o(h^n)$. On les multiplie en ne gardant théoriquement que les termes utiles :

$$u(a+h) \cdot v(a+h) = \sum_{\substack{i \leq n \\ j \leq n}} \frac{u^{(i)}(a) \cdot v^{(j)}(a)}{i! \cdot j!} \cdot h^{i+j} + o(h^n)_{h \rightarrow 0}$$

Il ne reste plus qu'à identifier le terme d'exposant n par un argument d'unicité du développement limité :

$$\frac{(u.v)^{(n)}(a)}{n!} = \sum_{i+j=n} \frac{u^{(i)}(a) \cdot v^{(j)}(a)}{i! \cdot j!}.$$

C'est le quotient $\frac{n!}{i! \cdot j!}$ qui recrée le coefficient binomial attendu.

Cette méthode a un défaut quand même, car elle dit juste que $u \cdot v$ a un développement d'ordre n mais n'assure pas qu'il soit forcément donné par une formule de Taylor avec $u \cdot v$ effectivement n fois dérivable.

En utilisant cette formule, on aura intérêt à faire porter sur la fonction la plus simple l'exposant le plus compliqué. Par exemple, pour dériver n fois $x \mapsto (x^3 + 1) \cdot e^x$, on posera $u = x \mapsto x^3 + 1$ et $v = \exp$ pour

$$\text{utiliser } \sum_k \binom{n}{k} \cdot u^{(k)} \cdot v^{(n-k)}.$$

On rappelle enfin que cette formule n'a de sens qu'à l'étage des fonctions, et que

$$\sum_{k=0}^n \binom{n}{k} \cdot (x^3 + 1)^{(k)} \cdot (e^x)^{(n-k)} \text{ n'a pas de sens, même si il pourra être toléré à un oral de concours.}$$

○ Donner le développement limité du sinus, du cosinus et de l'exponentielle en 0. Obtenir leur développement en un point a .

Sans aucune difficulté (et même par définition) : $e^x = 1 + x + \frac{x^2}{2} + \dots + \frac{x^n}{n!} + o(x^n) = \sum_{k=0}^n \frac{x^k}{k!} + o(x^n)_{x \rightarrow 0}$.

En ne gardant que la partie paire ou impaire :

$$\cosh(x) = 1 + \frac{x^2}{2} + \dots + \frac{x^{2 \cdot n}}{(2 \cdot n)!} + o(x^{2 \cdot n}) = \sum_{k=0}^n \frac{x^{2 \cdot k}}{(2 \cdot k)!} + o(x^{2 \cdot n})$$

$$\text{et } sh(x) = x + \frac{x^3}{6} + \dots + \frac{x^{2.n+1}}{(2.n+1)!} + o(x^{2.n+1}) = \sum_{k=0}^n \frac{x^{2.k+1}}{(2.k+1)!} + o(x^{2.n+1})$$

Si l'on veut s'arrêter à une puissance x^n , on pourra écrire $\sum_{\substack{k \leq n \\ k \text{ pair}}} \frac{x^k}{k!}$ ou aussi $\sum_{p=0}^{[n/2]} \frac{x^{2.p}}{(2.p)!}$.

Un schéma similaire sur les dérivées successives donne alors :

$$\cos(x) = 1 - \frac{x^2}{2} + \dots + (-1)^n \cdot \frac{x^{2.n}}{(2.n)!} + o(x^{2.n}) = \sum_{k=0}^n \frac{(-1)^k \cdot x^{2.k}}{(2.k)!} + o(x^{2.n}) \text{ et}$$

$$\sin(x) = x - \frac{x^3}{6} + \dots + (-1)^n \cdot \frac{x^{2.n+1}}{(2.n+1)!} + o(x^{2.n+1}) = \sum_{p=0}^n \frac{(-1)^p \cdot x^{2.p+1}}{(2.p+1)!} + o(x^{2.n+1})$$

L'alternance de signes peut se comprendre par le fait que les deux sommes de séries sont des fonctions bornées sur $\mathbb{R}$.

Ces développements ne sont utilisables qu'en 0 évidemment.

Ailleurs, on utilise les formules $\cos(a+h) = \cos(a) \cdot \cos(h) - \sin(a) \cdot \sin(h)$ et $\sin(a+h) = \cos(a) \cdot \sin(h) + \sin(a) \cdot \cos(h)$ dans lesquelles $\cos(a)$ et $\sin(a)$ sont de simples réels (*ne comptez pas sur moi, même dans ce document pour parler de "constantes"*).

On obtient par exemple : $\cos(a+h) = \cos(h) - \sin(a) \cdot h - \frac{\cos(a)}{2} \cdot h^2 - \frac{\sin(a)}{6} \cdot h^3 + \frac{\cos(a)}{24} \cdot h^4 + o(h^4)_{h \rightarrow 0}$ dans laquelle on sent aussi un grand usage de la formule de Taylor-Young.

○ Calculer la signature d'une permutation.

Il y a d'une part la définition, et ensuite le calcul.

Si σ est une bijection de $\{1, \dots, n\}$ dans lui-même, sa permutation indique la parité du nombre d'inversions, c'est à dire $Sgn(\sigma) = (-1)^{inv(\sigma)}$ avec $inv(\sigma) = Card\{(i, j) \in \{1, \dots, n\}^2 \mid i < j \text{ et } \sigma(i) > \sigma(j)\}$.

$$\prod_{1 \leq i < j \leq n} (\sigma(j) - \sigma(i))$$

On le calcule par $\frac{\prod_{1 \leq i < j \leq n} (j - i)}{\prod_{1 \leq i < j \leq n} (j - i)}$.

Le calcul effectif se fait en décomposant la permutation

- en produit de transpositions $\tau_{i,j}$: la signature est $(-1)^{\text{nombre de transpositions}}$
- en produit de cycles (*de supports disjoints pour la lisibilité*) : la signature d'un cycle de taille p est $(-1)^{p-1}$, et il ne reste plus qu'à multiplier les signatures.

Avec ces deux définitions, on peut travailler sur les permutations d'une liste non numérique.

Dès que n a dépassé 2, il y a autant de permutations de signature 1 (dites "paires") que de permutations de signature -1 (dites "impaires").

○ Démontrer que la variance d'une variable aléatoire est invariante par translation.

La variance mesure l'écart à la moyenne. Or, la moyenne se translate d'autant que ce qu'on a traduit la variable aléatoire. Il est donc normal qu'elle soit invariante.

On note A la variable aléatoire, de moyenne $E(A)$. On la translate de μ .

La nouvelle variable aléatoire a pour moyenne $E(A + \mu) = E(A) + \mu$,

et a pour variance $E\left(\left((A + \mu) - (E(A) + \mu)\right)^2\right)$ dans laquelle on retrouve $E\left((A - E(A))^2\right)$.

On pourrait aussi utiliser la formule

$$Var(A + \mu) = E((A + \mu)^2) - (E(A + \mu))^2 = E(A^2 + 2.A.\mu + \mu^2) - (E(A) + \mu)^2 = E(A^2) - E(A)^2.$$

○ Énoncer le théorème de Bolzano-Weierstrass sur $\mathbb{R}$. Démontrer sa généralisation à $\mathbb{C}$.

Toute suite réelle bornée admet au moins une sous-suite convergente.

Toute suite réelle bornée admet au moins une valeur d'adhérence.

On prend une suite complexe qu'on écrit $z_n = x_n + i.y_n$ (pour tout n , avec (x_n) et (y_n) réelles). On la suppose bornée par μ (en module). Il s'ensuit que (x_n) et (y_n) sont bornées elles aussi par μ (en valeur absolue).

De la suite réelle bornée (x_n) on peut extraire au moins une sous-suite $(x_{\varphi(n)})$ qui converge (vers un réel α).

La suite $(y_{\varphi(n)})$ est à son tour bornée (*extraite de* (y_n)). On peut en extraire une sous-sous-suite $(y_{\varphi(\psi(n))})$ qui converge (*vers* β).

La suite $(x_{\varphi \circ \psi(n)})$ continue à converger vers α et la suite complexe $(z_{\varphi \circ \psi(n)})$ converge donc par théorème algébrique (*vers* $\alpha + i\beta$). Et elle est bien extraite de la suite (z_n) .

La méthode se généralise à un espace produit du type $\mathbb{R}^N$ avec N fixé, par extraction en cascade.

○ **Savoir intégrer** $x \mapsto \frac{1}{x^2 + s.x + p}$ **quand le dénominateur n'a pas de racine réelle.**
Généraliser à $x \mapsto \frac{a.x + b}{x^2 + s.x + p}$.

L'application $x \mapsto \frac{1}{x^2 + s.x + p}$ est alors définie partout et continue, on peut l'intégrer sur tout segment $[a, b]$.

La méthode consiste à mettre le dénominateur sous forme canonique pour y retrouver une dérivée d'arctangente :

$$\frac{1}{x^2 + s.x + p} = \frac{1}{\left(x + \frac{s}{2}\right)^2 + p - \frac{s^2}{4}} = \frac{4}{4.p - s^2} \cdot \frac{1}{1 + \left(\frac{2.x + s}{\sqrt{4.p - s^2}}\right)^2}$$
 sachant que $4.p - s^2$ est justement

strictement positif.

On intègre en $\frac{2}{\sqrt{4.p - s^2}} \cdot \text{Arctan}\left(\frac{2.x + s}{\sqrt{4.p - s^2}}\right)$.

On ne retient évidemment pas la formule toute prête mais la méthode (factorisation canonique, arctangente).

Pour $\frac{a.x + b}{x^2 + s.x + p}$ on essaye de retrouver une forme en $\frac{u'}{u}$ pour commencer et on ajoute ce qu'il faut : $\frac{a.x + b}{x^2 + s.x + p} = \frac{a}{2} \cdot \frac{2.x + s}{x^2 + s.x + p} + \left(b - \frac{a.s}{2}\right) \cdot \frac{1}{x^2 + 2.s.x + p}$. On aura un terme logarithmique et un terme en arctangente.

○ **Montrer que les (éventuelles) primitives d'une application f sur un intervalle ne diffèrent que d'une constante.**

La définition de " F est une primitive de f " est " F est dérivable de dérivée f " (*définition qui n'a a priori pas de rapport avec la notion d'intégrale, mais juste avec un opérateur réciproque de la dérivation*).

On prend deux primitives de f qu'on note F et ϕ . On a alors $F' = f$ et $\phi' = f$. Par linéarité de la dérivation : $(F - \phi)' = 0$. C'est de cela qu'on déduit que $F - \phi$ est constante sur l'intervalle d'étude.

On prend a et b dans I et on applique le théorème de Rolle :

$\exists c \in I$, $(F - \phi)(b) - (F - \phi)(a) = (b - a) \cdot (F - \phi)'(c)$ (c est entre a et b donc dans I). Par hypothèse de nullité de la dérivée : $(F - \phi)(b) - (F - \phi)(a) = 0$ et donc $(F - \phi)(b) = (F - \phi)(a)$.

Rappelons que la définition de " g est constante" n'est ni $g' = 0$, ni $g(t) = C^{te}$, ni même $g = C^{te}$. C'est à la rigueur $\exists C, \forall t, g(t) = C$. La meilleure quantification est $\forall(a, b), g(a) = g(b)$.

○ **Donner la définition de " f est lipschitzienne de I dans $\mathbb{R}$ (ou $\mathbb{C}$)". Stabilités par addition, composition.**

f est lipschitzienne de rapport K de I (intervalle) dans $\mathbb{R}$ si on a : $\forall(a, b) \in I^2, |f(b) - f(a)| \leq K \cdot |b - a|$.

On pourra écrire $\forall(a, b) \in I^2, a \neq b \Rightarrow \frac{|f(b) - f(a)|}{|b - a|} \leq K$ qui possède une interprétation géométrique en termes de taux d'accroissements.

Le plus petit K vérifiant cette quantification est appelé "rapport de Lipschitz de f ".

La définition de " f est lipschitzienne" est $\exists K \in \mathbb{R}^+, \forall(a, b) \in I^2, |f(b) - f(a)| \leq K \cdot |b - a|$.

On prend f et g lipschitziennes, de rapports respectifs K et H . On a alors pour tout couple (a, b) par inégalité triangulaire :

$$|(f + g)(a) - (f + g)(b)| \leq |f(a) - f(b)| + |g(a) - g(b)| \leq K \cdot |b - a| + H \cdot |b - a| = (H + K) \cdot |b - a|$$

La somme $K + H$ n'est plus forcément le meilleur rapport pour l'application $f + g$.

Si ensuite f et φ sont lipschitziennes de rapports respectifs K et κ sur des domaines compatibles. On a alors pour tout couple (a, b) :

$$|f(\varphi(b)) - f(\varphi(a))| \leq K \cdot |\varphi(b) - \varphi(a)| \leq K \cdot \kappa \cdot |b - a|.$$

Le produit de deux applications lipschitziennes ne l'est plus forcément. Un exercice simple montre que si f et g sont lipschitziennes et bornées, alors leur produit l'est aussi.

La notion d'application lipschitzienne se généralise à des espaces munis d'une distance (espaces normés), tant au départ qu'à l'arrivée.

○ Montrer par récurrence sur n que dans tout espace vectoriel $(E, +, \cdot)$ engendré par n vecteurs $(\vec{e}_1, \dots, \vec{e}_n)$, toute famille $(\vec{a}_0, \dots, \vec{a}_n)$ de $n + 1$ vecteurs (et plus) est liée (théorème fondamental de la dimension finie).

On initialise à n égal à 0, par pure logique.

Le seul espace engendré par une famille vide est réduit au seul vecteur nul : $\{\vec{0}_E\}$. La seule famille de 1 vecteur est $(\vec{0}_E)$; elle est liée (par la relation $1 \times \vec{0}_E = \vec{0}_E$ avec $1 \neq 0$).

On traite aussi le cas $n = 1$. On se place dans $(E, +, \cdot)$ engendré par un vecteur $\vec{e}_1$. On prend une famille de deux vecteurs $(\vec{a}_0, \vec{a}_1)$. On les écrit $\vec{a}_0 = \alpha_0 \cdot \vec{e}_1$ et $\vec{a}_1 = \alpha_1 \cdot \vec{e}_1$. Si α_0 est nul, la famille contient le vecteur nul et est donc liée (par $1 \cdot \vec{a}_0 + 0 \cdot \vec{a}_1 = \vec{0}$ si on y tient). Si α_0 est non nul, alors on a $\alpha_1 \cdot \vec{a}_0 - \alpha_0 \cdot \vec{a}_1 = \vec{0}$ avec au moins un coefficient non nul.

On suppose maintenant que dans tout espace vectoriel engendré par $n - 1$ vecteurs, toute famille de n vecteurs est liée. On met cette hypothèse notée H_{n-1} de côté, et on se place dans un espace vectoriel $(E, +, \cdot)$ engendré par n vecteurs : $E = \text{Vect}(\vec{e}_1, \dots, \vec{e}_n)$, et on y prend une famille de $n + 1$ vecteurs $(\vec{a}_0, \dots, \vec{a}_n)$.

$$\begin{aligned} \vec{a}_0 &= \alpha_0^1 \cdot \vec{e}_1 + \alpha_0^2 \cdot \vec{e}_2 + \dots + \alpha_0^n \cdot \vec{e}_n \\ \vec{a}_1 &= \alpha_1^1 \cdot \vec{e}_1 + \alpha_1^2 \cdot \vec{e}_2 + \dots + \alpha_1^n \cdot \vec{e}_n \\ &\vdots \\ \vec{a}_n &= \alpha_n^1 \cdot \vec{e}_1 + \alpha_n^2 \cdot \vec{e}_2 + \dots + \alpha_n^n \cdot \vec{e}_n \end{aligned}$$

On les décompose suivant la famille génératrice :

On distingue alors deux cas :

• tous les α_k^1 sont nuls. Dans ce cas, les n vecteurs $\vec{a}_k$ (pour k de 0 à $n - 1$ déjà) sont dans $\text{Vect}(\vec{e}_2, \dots, \vec{e}_n)$. Par H_{n-1} ils forment une famille liée. On agrandit : la famille des $\vec{a}_k$ (pour k de 0 à n a fortiori) est liée. On a établi la propriété au rang n , sauf qu'il faut encore étudier l'autre cas. • au moins un des α_k^1 est non nul. Quitte à réordonner la famille (ce qui ne change rien au caractère lié/libre), on suppose donc $\alpha_0^1 \neq 0$. On utilise alors $\vec{a}_0$ pour éliminer les "composantes" suivant $\vec{e}_1$

$$\begin{aligned} \vec{a}_0 &= \alpha_0^1 \cdot \vec{e}_1 + \alpha_0^2 \cdot \vec{e}_2 + \dots + \alpha_0^n \cdot \vec{e}_n \\ \vec{b}_1 &= \vec{a}_1 - \frac{\alpha_1^1}{\alpha_0^1} \cdot \vec{a}_0 = 0 \cdot \vec{e}_1 + \left(\alpha_1^2 - \frac{\alpha_1^1}{\alpha_0^1} \cdot \alpha_0^2 \right) \cdot \vec{e}_2 + \dots + \left(\alpha_1^n - \frac{\alpha_1^1}{\alpha_0^1} \cdot \alpha_0^n \right) \cdot \vec{e}_n \\ &\vdots \\ \vec{b}_n &= \vec{a}_n - \frac{\alpha_n^1}{\alpha_0^1} \cdot \vec{a}_0 = 0 \cdot \vec{e}_1 + \left(\alpha_n^2 - \frac{\alpha_n^1}{\alpha_0^1} \cdot \alpha_0^2 \right) \cdot \vec{e}_2 + \dots + \left(\alpha_n^n - \frac{\alpha_n^1}{\alpha_0^1} \cdot \alpha_0^n \right) \cdot \vec{e}_n \end{aligned}$$

La formule est $\vec{b}_k = \vec{a}_k - \frac{\alpha_k^1}{\alpha_0^1} \cdot \vec{a}_0$. Les n vecteurs $(\vec{b}_1, \dots, \vec{b}_n)$ sont dans $\text{Vect}(\vec{e}_2, \dots, \vec{e}_n)$ (engendré par $n - 1$

vecteurs). Ils forment une famille liée, par hypothèse H_{n-1} par une relation de la forme $\sum_{k=2}^n \beta_k \cdot \vec{b}_k = \vec{0}_E$

avec au moins un des β_k non nul. On remplace : $\sum_{k=2}^n \beta_k \cdot \vec{a}_k - \left(\sum_{k=2}^n \frac{\alpha_k^1}{\alpha_0^1} \cdot \beta_k \right) \cdot \vec{a}_0 = \vec{0}_E$ qu'on écrit même

$\sum_{k=1}^n \beta_k \cdot \vec{a}_k = \vec{0}_E$ avec $\beta_1 = - \sum_{k=2}^n \frac{\alpha_k^1}{\alpha_0^1} \cdot \beta_k$. Que ce dernier coefficient soit nul ou non, l'un au moins des β_k est non nul. La famille des $\vec{a}_k$ est liée.

○ Donner des définitions équivalentes de " H est un sous-groupe de $(G, *)$ ". Connaître les sous-groupes de $(\mathbb{Z}, +)$.

La définition la plus simple est " H est une partie de G qui devient elle-même un groupe pour la loi $*$ ". On peut se passer de l'associativité, déjà acquise dans G et se contenter des autres propriétés, tout en sachant que tout a de H admet déjà un symétrique...dans G .

inclusion	$H \subset G$
stabilité	$\forall (a, b) \in H^2, a * b \in H$
neutre	$e \in H$
symétriques	$\forall b \in H, b^{-1} \in H$

	inclusion	$H \subset G$	
On allège en	stabilité “soustractive”	$\forall(a, b) \in H^2, a * b^{-1} \in H$	et même
	neutre	$e \in H$	

inclusion et non vacuité	$\emptyset \neq H \subset G$
stabilité “soustractive”	$\forall(a, b) \in H^2, a * b^{-1} \in H$

La dernière définition est généralement inutile. Pour montrer que H est non vide, le meilleur éléments à y chercher est toujours le neutre e de G .

Les sous groupes de $(\mathbb{Z}, +)$ sont de la forme $a\mathbb{Z}$ pour a dans $\mathbb{N}$. On dit qu'ils sont monogènes, c'est à dire “engendrés par un élément”.

On commence par montrer que tout ensemble de la forme $a\mathbb{Z}$ est un sous-groupe de $(\mathbb{Z}, +)$.

Le cas particulier $a = 0$ donne le sous-groupe trivial $\{0\}$, et le cas particulier $a = 1$ donne le sous-groupe tout aussi trivial $\mathbb{Z}$.

L'ensemble $a\mathbb{Z}$ est donc $\{a.k \mid k \in \mathbb{Z}\}$, formé des multiples de a . Il est inclus dans $\mathbb{Z}$ par stabilité multiplicative de $(\mathbb{Z}, +, \times)$. On y trouve le neutre en prenant le cas particulier $k = 0$. On prend deux éléments de $a\mathbb{Z}$, qu'on écrit $a.k$ et $a.h$ avec h et h dans $\mathbb{Z}$; leur différence est $a.(k - h)$, et c'est encore un multiple de a .

Il faut montrer à présent que tout sous-groupe H de $(\mathbb{Z}, +)$ est de cette forme.

On prend donc un sous-groupe H (ensemble d'entiers contenant 0 et stable par addition et soustraction). On traite déjà un cas particulier.

- Si H ne contient aucun entier strictement positif, alors par passage à l'opposé il ne contient aucun élément strictement négatif. Il se réduit donc à 0 (qu'il doit contenir). On est en droit d'écrire $H = 0\mathbb{Z}$.
- Si H contient des entiers strictement positifs, on pose $A = H \cap \mathbb{N}^*$. C'est une partie non vide de $\mathbb{N}$; on peut en prendre le plus petit élément, qu'on note a . C'est le plus petit élément strictement positif de H .

On va alors prouver que H est formé des multiples de a .

Comme a est dans H , par stabilité, $a + a$, puis $(a + a) + a$ sont dans H . Par récurrence évidente, tous les $n.a$ avec n dans $\mathbb{N}^*$ sont dans H . Par passage à l'opposé, tous les $k.a$ avec k dans $\mathbb{Z}$ sont dans H (le cas particulier $k = 0$ repose sur “ H est un sous-groupe de $(\mathbb{Z}, +)$, il contient donc le neutre”).

Il reste à montrer l'inclusion réciproque : $H \subset a\mathbb{Z}$. On prend un élément n de H . On écrit la division euclidienne de n par a non nul (y compris pour n négatif) : $n = a.q + r$ avec r entre 0 (inclus) et a (exclu). On écrit alors $r = n - a.q$. Les éléments n et $a.q$ sont dans H (pour n , c'est la question précédente, pour $a.q$ c'est l'inclusion précédente). Leur différence est dans H . Et elle est dans $\mathbb{N} \cap [0, a[$. Si elle est non nulle, on a un élément de A plus petit que a lui même, ce qui est contradictoire. Par élimination, $n - a.q$ est nul. On reporte : $n = a.q$. C'est un multiple de a .

On a bien établi que H est formé des multiples de a , et rien que des multiples de a .

○ Montrer qu'il y a équivalence pour une formule bilinéaire entre “antisymétrique” et “alternée”.

On prend donc une forme bilinéaire f sur un espace vectoriel $(E, +, \cdot)$. Elle prend deux vecteurs et calcule un réel.

On la suppose antisymétrique : $\forall(\vec{a}, \vec{b}) \in E^2, f(\vec{b}, \vec{a}) = -f(\vec{a}, \vec{b})$. Dans le cas particulier $\vec{a} = \vec{b}$, on trouve $f(\vec{a}, \vec{a}) = -f(\vec{a}, \vec{a})$ et donc $f(\vec{a}, \vec{a}) = 0$. La forme est alternée.

On la suppose alternée : $\forall \vec{u} \in E, f(\vec{u}, \vec{u}) = 0$. On se donne alors deux vecteurs $\vec{a}$ et $\vec{b}$. Par caractère alterné, la somme $f(\vec{a} + \vec{b}, \vec{a} + \vec{b}) - f(\vec{a}, \vec{a}) - f(\vec{b}, \vec{b})$ est nulle (trois réels nuls). On développe par bilinéarité et on simplifie : $f(\vec{a}, \vec{b}) + f(\vec{b}, \vec{a}) = 0$. On a bien : $\forall(\vec{a}, \vec{b}) \in E^2, f(\vec{b}, \vec{a}) = -f(\vec{a}, \vec{b})$; la forme est antisymétrique.

Ce résultat se généralise aux formes multilinéaires $(\vec{a}_1, \dots, \vec{a}_n) \mapsto f(\vec{a}_1, \dots, \vec{a}_n)$, avec l'équivalence entre les deux définitions :

- alternée : le réel est nul dès qu'il y a ne serait ce que deux vecteurs égaux
- antisymétrique : le signe change si on intervertit deux vecteurs de la liste et même $f(\vec{a}_{\sigma(1)}, \dots, \vec{a}_{\sigma(n)}) = \text{Sgn}(\sigma).f(\vec{a}_1, \dots, \vec{a}_n)$ pour tout n -uplet de vecteurs $(\vec{a}_1, \dots, \vec{a}_n)$ et toute permutation σ .

○ Savoir dériver $x \mapsto \ln(x + \sqrt{x^2 + 1})$.

Si la question se limite à un calcul, on dérive en $\frac{u'}{u}$. On trouve alors $x \mapsto \frac{1 + \frac{2.x}{2.\sqrt{x^2+1}}}{x + \sqrt{x^2+1}}$ et on simplifie en $x \mapsto \frac{1}{\sqrt{1+x^2}}$.

On peut donc dire qu'une primitive de $x \mapsto \frac{1}{\sqrt{1+x^2}}$ est $x \mapsto \ln(x + \sqrt{1+x^2})$.

C'est la primitive préconisée par le programme, même si on peut observer que c'est la fonction réciproque que $t \mapsto \frac{e^t - e^{-t}}{2}$, et qu'elle s'appelle aussi *Argsh*.

On donne d'ailleurs la liste des primitives usuelles

fonction	$x \mapsto \frac{1}{\sqrt{1+x^2}}$	$x \mapsto \frac{1}{\sqrt{x^2-1}}$	$x \mapsto \frac{1}{1-x^2}$
primitive	$x \mapsto \ln(x + \sqrt{1+x^2})$	$x \mapsto \ln(x + \sqrt{x^2-1})$	$x \mapsto \ln \left \frac{x+1}{x-1} \right $
domaine	$\mathbb{R}$	$[1, +\infty[$	$] -\infty, -1[$ ou $] 1, +\infty[$
autre nom	<i>Argsh</i>	<i>Argch</i>	<i>Argth</i> sur $] -1, 1[$

Il reste quand même à s'assurer que $x \mapsto \ln(x + \sqrt{x^2+1})$ est bien définie sur $\mathbb{R}$ (et dérivable).

Une méthode consiste à distinguer deux cas : pour x positif, $x + \sqrt{x^2+1}$ est positif comme somme de réels positifs. Pour x négatif, $\sqrt{x^2+1}$ est plus grand que $\sqrt{x^2}$, et déjà, $x + \sqrt{x^2}$ est nul.

On peut aussi regarder l'équation $T^2 - 2.X.T - 1$ de racines $x + \sqrt{x^2+1}$ et $x - \sqrt{x^2+1}$; leur produit vaut -1 , c'est que l'une est négative et l'autre positive. La positive est donc $x + \sqrt{x^2+1}$ (la plus grande).

Les primitives d'applications de la forme $x \mapsto \frac{1}{\sqrt{a.x^2 + b.x + c}}$ avec $b^2 - 4.a.c < 0$ se ramènent par changement de variable à la forme précédente.

○ Calculer la somme des premiers termes d'une série géométrique. Calculer la somme des termes consécutifs d'une série géométrique.

La somme $\sum_{k=0}^n a.r^k$ est égale à $a \cdot \frac{1-r^{n+1}}{1-r}$.

On peut le prouver par récurrence sur n . On peut.

Le mieux est d'effectuer un produit "en croix" $(1-r) \cdot \sum_{k=0}^n a.r^k$ et d'y voir une somme télescopique.

La formule générale est $\sum_{k=p}^q r^k = \frac{r^p - r^{q+1}}{1-r}$.

Il faut évidemment mettre de côté le cas $r = 1$ et donc n'utiliser la formule qu'après avoir indiqué "série géométrique de raison différente de 1".

Il n'est pas recommandé de retenir une formule qui compte le nombre de termes, sauf à vraiment savoir sans faille que l'intervalle $[p, q]$ contient $q - p + 1$ termes.

Il vaut mieux retenir $\frac{(\text{premier terme écrit}) - (\text{premier terme à venir})}{1 - \text{raison}}$ qui voit la série comme une suite de termes qui peut se prolonger et non comme un objet figé.

Pour r strictement plus petit que 1 en module, on peut faire tendre le nombre de termes vers l'infini

et obtenir $\sum_{k=0}^{+\infty} r^k = \frac{1}{1-r}$.

○ Démontrer que la somme de deux suites convergentes converge vers la somme de leurs limites. Même question avec produit.

On prend deux suite convergentes (a_n) et (b_n) de limites respectives α et β . On quantifie :

$\forall \varepsilon > 0, \exists A_\varepsilon \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq A_\varepsilon \Rightarrow |a_n - \alpha| \leq \varepsilon$

$\forall \varepsilon > 0, \exists B_\varepsilon \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq B_\varepsilon \Rightarrow |b_n - \beta| \leq \varepsilon$

On prouve déjà $\forall \varepsilon > 0, \exists S_\varepsilon \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq S_\varepsilon \Rightarrow |(a_n + b_n) - (\alpha + \beta)| \leq \varepsilon$.

Il suffit en effet pour ε donné de poser $S_\varepsilon = \max(A_{\varepsilon/2}, B_{\varepsilon/2})$ et de vérifier pour n plus grand que ce maximum :

$|(a_n + b_n) - (\alpha + \beta)| \leq |a_n - \alpha| + |b_n - \beta| \leq 2.\varepsilon/2.$

Pour le produit, la chose est plus délicate.

On commence par montrer que (a_n) est bornée par $|\alpha| + 1$ à partir du rang A_1 .

Ensuite, on écrit $|a_n.b_n - \alpha.\beta| = |a_n.(b_n - \beta) + \beta.(a_n - \alpha)| \leq |a_n|. |b_n - \beta| + |\beta|. |a_n - \alpha| \leq (|\alpha| + 1). |b_n - \beta| + |\beta|. |a_n - \alpha|$.

Si l'on prend alors n plus grand que $\max(A_1, B_{\varepsilon/(2.|\alpha|+2)}, A_{\varepsilon/(2.|\beta|+1)})$, alors le membre de droite de la majoration est lui même plus petit que $|\alpha| + 1|. \frac{\varepsilon}{2.|\alpha|+2} + |\beta|. \frac{\varepsilon}{2.|\beta|+1}$ lui même plus petit que ε .

On a donc bien obtenu $\forall \varepsilon > 0, \exists P_\varepsilon \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq P_\varepsilon \Rightarrow |a_n.b_n - \alpha.\beta| \leq \varepsilon$ en posant $P_\varepsilon = \max(A_1, B_{\varepsilon/(2.|\alpha|+2)}, A_{\varepsilon/(2.|\beta|+1)})$.

On notera qu'on a pris $A_{\varepsilon/(2.|\beta|+1)}$ alors que $A_{\varepsilon/(2.|\beta|)}$ semblait suffire. Mais il faut tenir compte du cas (certes évident) où β est nul.

○ Déterminer deux nombres dont on connaît la somme et le produit.

La question revient à résoudre le système $\begin{cases} x + y = S \\ x \times y = P \end{cases}$ avec S et P donnés, et les inconnues x et y qui jouent des rôles symétriques.

On peut bien évidemment exprimer y en fonction de x et reporter une équation dans l'autre. On peut... On considère le polynôme unitaire du second degré dont les racines sont x et y : $(X - x).(Y - y)$. C'est tout simplement après développement $X^2 - S.X + P$.

La réponse est donc : x et y sont les racines du polynôme $X^2 - S.X + P$.

Si on y tient : $S = \left\{ \left(\frac{S + \sqrt{S^2 - 4.P}}{2}, \frac{S - \sqrt{S^2 - 4.P}}{2} \right), \left(\frac{S - \sqrt{S^2 - 4.P}}{2}, \frac{S + \sqrt{S^2 - 4.P}}{2} \right) \right\}$.

On notera que pour qu'il y ait des solutions réelles, on a forcément $S^2 - 4.P \geq 0$, c'est l'inégalité $\frac{a+b}{2} \geq \sqrt{a.b}$. On insistera sur la formulation "sont les solutions", qui donne $S = \{(a, b), (b, a)\}$ et pas $S = \{a, b\}$.

On généralise : les triplets solutions de $\begin{cases} a + b + c = S \\ b.c + a.c + a.b = D \\ a \times b \times c = P \end{cases}$ sont formés des solutions de $X^3 - S.X^2 + D.X + P$.

○ Donner la définition d'une extraction de $\mathbb{N}$ dans $\mathbb{N}$. Établir la propriété $\varphi(n) \geq n$. Définition de "suite extraite d'une suite (u_n) ". Suite extraite d'une sous-suite. Convergence des suites extraites d'une suite convergente.

Une extraction est une application strictement croissante de $\mathbb{N}$ dans $\mathbb{N}$.

On quantifie : $\forall n \in \mathbb{N}, \varphi(n) \in \mathbb{N}$ et $\varphi(n+1) > \varphi(n)$.

L'extraction la plus simple est Id (ce qui permet de dire qu'une suite est toujours une sous-suite d'elle même).

Toute extraction croît au moins aussi vite que Id et vérifie donc $\forall n \in \mathbb{N}, \varphi(n) \geq n$.

On établit cette propriété par récurrence sur n .

Pour n égal à 0, $\varphi(0)$ est un entier naturel, il est donc bien positif ou nul.

On suppose ensuite pour un n donné : $\varphi(n) \geq n$. Par croissance stricte de φ : $\varphi(n+1) > \varphi(n)$. Par transitivité, $\varphi(n+1) > n$. Or, le plus petit entier strictement plus grand que n est $n+1$. On a donc bien $\varphi(n+1) \geq n+1$.

On dit que v est une suite extraite de u (ou que v est une sous-suite de u) si il existe une extraction φ vérifiant $\forall n, v_n = u_{\varphi(n)}$ ou plus abruptement : $v = u \circ \varphi$.

Cette relation "être extraite de" est réflexive et transitive.

On prend b extraite de a via une extraction α et c extraite de b via une extraction β .

On a alors pour tout n : $c_n = b_{\beta(n)} = a_{\alpha(\beta(n))}$.

On reconnaît que c est extraite de a via l'extraction $\alpha \circ \beta$ (et pas $\beta \circ \alpha$).

Le caractère "croissante", "bornée", "positif" se transmet par extraction.

La convergence aussi. On suppose que a converge vers α et que φ est une extraction.

On quantifie : $\forall \varepsilon > 0, \exists N_\varepsilon \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq N_\varepsilon \Rightarrow |a_n - \alpha| \leq \varepsilon$.

On a alors aussi $\forall \varepsilon > 0, \exists N_\varepsilon \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq N_\varepsilon \Rightarrow |a_{\varphi(n)} - \alpha| \leq \varepsilon$.

En effet : $n \geq N_\varepsilon \Rightarrow \varphi(n) \geq n \geq N_\varepsilon$.

○ Donner la forme générale des différentes matrices “de Gauss” qui permettent de permuter deux lignes, multiplier une ligne par un coefficient non nul, effectuer une combinaison linéaire $L_q \leftarrow L_q + \alpha.L_p$ dans un système de k équations à n inconnues. Exprimer l'inverse de chacune.

Un système linéaire de k équations à n inconnues $\left\{ \begin{array}{cccc|c} a_1^1.x_1 & +a_1^2.x_2 & +\dots & +a_1^n.x_n & =b_1 \\ a_2^1.x_1 & +a_2^2.x_2 & +\dots & +a_2^n.x_n & =b_2 \\ \vdots & \vdots & & \vdots & \vdots \\ a_k^1.x_1 & +a_k^2.x_2 & +\dots & +a_k^n.x_n & =b_k \end{array} \right. \quad \text{s'écrit}$

$A.X = B$ et se résout en écrivant peu à peu $G_1.A.X = G_1.B$ jusqu'à arriver à $T.X = G.B$ avec T “simple”.

Pour effectuer des opérations sur les lignes, on multiplie à gauche par des matrices de taille k sur k :

opération	exemple	notation	inverse
permutation	$\begin{pmatrix} 0 & 0 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}$	$T_{i,j}$	elle même
multiplication	$\begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & \alpha & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}$	$M_{i,\alpha}$	$M_{i,1/\alpha}$
combinaison	$\begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & \alpha \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}$	$C_{p,q,\alpha}$	$C_{p,q,-\alpha}$

Par multiplication par ces mêmes matrices à droite, on effectue des opérations sur les colonnes.

Pour calculer un déterminant, on peut mélanger les deux méthodes. Pour résoudre un système, on ne travaille que sur les lignes, car travailler sur les colonnes revient à changer les inconnues.

○ Donner la définition de $(E, +, \cdot)$ est un $\mathbb{R}$ -espace vectoriel. Démontrer $(\alpha.\vec{u} = \vec{0}) \Rightarrow (\alpha = 0 \text{ ou } \vec{u} = \vec{0})$.

Un espace vectoriel est déjà un groupe additif dont les éléments sont appelés “vecteurs”, muni en plus d'une loi dite “externe”.

loi interne	loi externe
$\forall(\vec{a}, \vec{b}) \in E^2, \vec{a} + \vec{b} \in E$	$\forall \vec{a} \in E, \forall \alpha \in \mathbb{R}, \alpha.\vec{a} \in E$
$\forall(\vec{a}, \vec{b}, \vec{c}) \in E^3, (\vec{a} + \vec{b}) + \vec{c} = \vec{a} + (\vec{b} + \vec{c})$	$\forall(\vec{a}, \vec{b}) \in E^2, \forall \alpha \in \mathbb{R}, \alpha.(\vec{a} + \vec{b}) = \alpha.\vec{a} + \alpha.\vec{b}$
$\forall(\vec{a}, \vec{b}) \in E^2, \vec{a} + \vec{b} = \vec{b} + \vec{a}$	$\forall \vec{a} \in E, \forall(\alpha, \beta) \in \mathbb{R}^2, (\alpha + \beta).\vec{a} = \alpha.\vec{a} + \beta.\vec{a}$
$\forall \vec{a} \in E, \vec{a} + \vec{0} = \vec{a}$	$\forall \vec{a} \in E, \forall(\alpha, \beta) \in \mathbb{R}^2, \alpha.(\beta.\vec{a}) = (\alpha \times \beta).\vec{a}$
$\forall \vec{a} \in E, \exists(-\vec{a}) \in E, \vec{a} + (-\vec{a}) = \vec{0}$	$\forall \vec{a} \in E, 1.\vec{a} = \vec{a}$

Le neutre est noté $\vec{0}$ et est appelé vecteur nul.

Le symétrique d'un vecteur $\vec{a}$ est noté $-\vec{a}$ et coïncide avec $(-1).\vec{a}$ comme on peut le démontrer.

On définit aussi des $\mathbb{C}$ -espaces vectoriels, ou plus généralement des $\mathbb{K}$ -espaces vectoriels, où $\mathbb{K}$ est un corps.

On montre que le vecteur nul et le réel nul sont “absorbants” : $\forall \vec{a}, 0.\vec{a} = \vec{0}$ et $\forall \alpha \in \mathbb{R}, \alpha.\vec{0} = \vec{0}$. En effet, on a pour $\vec{a}$ donné : $\vec{a} = 1.\vec{a} = (1 + 0).\vec{a} = 1.\vec{a} + 0.\vec{a} = \vec{a} + 0.\vec{a}$; il ne reste plus qu'à simplifier dans le groupe additif $(E, +)$ par $\vec{a}$ et il reste $0.\vec{a} = \vec{0}$.

De même, $\alpha.\vec{0} = \alpha.(\vec{0} + \vec{0}) = \alpha.\vec{0} + \alpha.\vec{0}$ et on simplifie encore.

Pour montrer la forme de réciproque $(\alpha.\vec{u} = \vec{0}) \Rightarrow (\alpha = 0 \text{ ou } \vec{u} = \vec{0})$, on commence par faire de la logique :

$(n \Rightarrow r \text{ ou } v) = (\bar{n} \text{ ou } r \text{ ou } v) = (\bar{n} \text{ et } \bar{r} \text{ ou } v) = ((n \text{ et } \bar{r}) \Rightarrow v)$.

On va donc prouver $(\alpha.\vec{u} = \vec{0} \text{ et } \alpha \neq 0) \Rightarrow (\vec{u} = \vec{0})$. On prend donc α non nul. L'inverse numérique α^{-1} existe. On multiplie l'hypothèse $\alpha.\vec{u} = \vec{0}$ par α^{-1} . On obtient $\alpha^{-1}(\alpha.\vec{u}) = \alpha^{-1}.\vec{0}$ puis

$(\alpha^{-1} \times \alpha) \cdot \vec{u} = \vec{0}$ et enfin $\vec{u} = \vec{0}$.

○ Donner la définition du degré d'un polynôme. Terme dominant, coefficient dominant. Degré de $P(X).Q(X)$, de $P(Q(X))$, majoration de $\deg(P(X) + Q(X))$.

Formellement, un polynôme est une suite nulle à partir d'un certain rang $(a_0, a_1, \dots, a_d, 0, \dots)$. On l'écrit plus naturellement $\sum_k a_k.X^k$. Il y a a priori une infinité de termes, mais comme les a_k sont nuls à

partir d'un certain rang d , la somme est finie : $\sum_{k=0}^d a_k.X^k$. Mais sous cette forme, le degré n'est pas forcément d . En effet, il est possible que même a_d soit nul.

La bonne définition est $\text{Max}\{k \in \mathbb{N} \mid a_k \neq 0\}$. Cette définition fait que le polynôme nul a pour degré le maximum de l'ensemble vide, c'est donc $-\infty$.

On peut aussi choisir $\text{Min}\{d \in \mathbb{N} \mid \forall k > d, a_k = 0\}$.

Par définition même, $a_{\deg(P)} \neq 0$ (sauf pour le polynôme nul). C'est lui qu'on qualifie de coefficient dominant. Le terme dominant est alors $a_{\deg(P)}.X^{\deg(P)}$.

degré	coefficient dominant	terme dominant
entier	nombre	monôme

On ne confondra donc pas

On écrit donc à présent $P(X) = \sum_{i=0}^p a_i.X^i$ avec $a_p \neq 0$ pour que P soit vraiment de degré p et

$$Q(X) = \sum_{j=0}^q b_j.X^j \text{ avec } b_q \neq 0.$$

La formule la plus légère est celle du produit : $P(X).Q(X) = \sum_{k,k'} a_k.b_{k'}.X^{k+k'}$ qu'on regroupe

$$P(X).Q(X) = \sum_k c_k.X^k \text{ avec } c_k = \sum_{i+j=k} a_i.b_j. \text{ Ces coefficients sont nuls au delà de } k = p + q. \text{ Pour}$$

$k = p + q$ on a $c_{p+q} = a_p.b_q$ (coefficient dominant du produit rapide de $a_p.X^p + \dots + a_0$ par $b_q.X^q + \dots + b_0$). Par intégrité de la multiplication, ce coefficient est non nul, et $P.Q$ est de degré $p + q$. On a donc $\deg(P \times Q) = \deg(P) + \deg(Q)$.

Pour composer, on imagine qu'on développe $\sum_{i=0}^p a_i \cdot \left(\sum_{j=0}^q b_j.X^j \right)^i$. On obtient des monômes dont le

degré ne peut excéder $p.q$. Il y a même un terme de degré $p.q$: il vient de $a_p \cdot \left(\sum_{j=0}^q b_j.X^j \right)^p$ et il vaut $a_p.(b_q)^p$. Il est non nul. On a donc $\deg(P \circ Q) = \deg(P) \times \deg(Q)$.

On notera que $Q \circ P$ a aussi le degré $\deg(P) \times \deg(Q)$.

Pour sommer, on revient à la forme générale "tous nuls à partir d'un certain rang" : $P(X) + Q(X) = \sum_k (a_k + b_k).X^k$. On est sûr qu'à partir du rang $\text{Max}(p, q)$, les termes $a_k + b_k$ sont nuls. ON a bien un polynôme. Mais son degré n'est pas forcément $\text{Max}(\deg(P), \deg(Q))$. Il est inférieur ou égal à cette valeur (et en général il lui est quand même égal. Quand p est différent de q , le terme $a_{\text{Max}(p,q)} + b_{\text{Max}(p,q)}$ est égal à a_p ou à b_q ; il est non nul. Mais il reste le cas où les deux polynômes sont de même degré. Les termes dominants peuvent se simplifier mutuellement, et le degré peut diminuer.

On résume

polynôme	$P(X) + Q(X)$			$P(X) \times Q(X)$	$P(Q(X))$
degré	$\leq \text{Max}(\deg(P), \deg(Q))$			$\deg(P) + \deg(Q)$	$\deg(P) \times \deg(Q)$
	$\deg(P) < \deg(Q)$	$\deg(P) = \deg(Q)$	$\deg(P) > \deg(Q)$		
	$\deg(Q)$	$\text{Max}(i \leq \deg(P) \mid a_i \neq -b_i)$	$\deg(P)$		

Évidemment, on traite de manière similaire le degré de $P(X) - Q(X)$ et de $\alpha.P(X) + \beta.Q(X)$, puis on trouve la formule pour le degré de $A(X).P(X) - B(X).Q(X)$.

○ Donner la définition de $\text{Vect}(\vec{a}_1, \dots, \vec{a}_p)$ dans un espace vectoriel $(E, +, \cdot)$.

La définition du sous-espace engendré par une partie est "le plus petit sous-espace vectoriel de $(E, +, \cdot)$ contenant cette partie".

Ici, on cherche le plus petit sous-espace contenant tous les $\vec{a_k}$.

Il doit contenir les combinaisons $\sum_{k=1}^p \alpha_k \cdot \vec{a_k}$ par stabilité.

Il suffit ensuite de se convaincre que si l'on s'arrête à l'ensemble de tous les $\sum_{k=1}^p \alpha_k \cdot \vec{a_k}$ on a bien un sous-espace vectoriel de $(E, +, \cdot)$.

On a donc $Vect(\vec{a_1}, \dots, \vec{a_p}) = \left\{ \sum_{k=1}^p \alpha_k \cdot \vec{a_k} \mid (\alpha_1, \dots, \alpha_p) \in \mathbb{R}^p \right\}$

C'est un espace vectoriel dont la dimension ne pourra dépasser p .

Elle sera même égale à p si et seulement si la famille $(\vec{a_1}, \dots, \vec{a_p})$ est libre.

On note que $(\alpha_1, \dots, \alpha_p) \mapsto \sum_{k=1}^p \alpha_k \cdot \vec{a_k}$ construit une application linéaire surjective de $\mathbb{R}^p$ dans $Vect(\vec{a_1}, \dots, \vec{a_p})$.

Si la famille est infinie, on prend l'espace formé de toutes les combinaisons linéaires de toutes les sous-familles finies.

○ Donner la définition d'une suite périodique. Montrer que l'ensemble des suites périodiques est un anneau pour les lois usuelles.

Une suite a est dite périodique de période p si p est un entier naturel non nul et qu'on a $\forall n \in \mathbb{N}, a_{n+p} = a_n$.

L'hypothèse "non nul" est évidemment capitale.

On a aussi $\forall (n, k) \in \mathbb{N}^2, a_{n+k \cdot p} = a_n$.

Une suite est dite périodique si elle admet une période : $\exists p \in \mathbb{N}^*, \forall n, a_{n+p} = a_n$.

On cherchera en général la plus petite période d'une suite périodique :

$Min\{p \in \mathbb{N}^* \mid \forall n \in \mathbb{N}, a_{n+p} = a_n\}$.

On prend deux suites a et b périodiques de périodes respectives p et q . On note P le p.p.c.m. de p et q (plus petit multiple commun ou même produit des deux si on n'a pas envie de chercher si fin). On a alors pour tout n : $a_{n+P} = a_n$ et $b_{n+P} = b_n$. En additionnant les deux relations, on trouve que $a + b$ est encore périodique de période P (pas forcément la plus petite période), puis en multipliant, on reconnaît aussi que $a \cdot b$ est aussi P -périodique.

On montre ensuite que la suite nulle est périodique (période 1 par exemple), et on déduit que les suites périodiques est un sous-anneau de l'anneau des suites.

On montre aussi la stabilité par multiplication par un réel, et on a une algèbre.

Pour les applications, comme il n'y a plus de notion de p.p.c.m. sur les périodes, on n'a plus de telles stabilités, comme l'atteste l'exemple de $x \mapsto \cos(x) + \cos(\sqrt{2} \cdot x)$.

○ Exprimer le terme général d'une suite de la forme $u_{n+1} = a \cdot u_n + b$ avec u_0 donné, en se ramenant à une suite géométrique par translation.

On peut de proche en proche calculer tous les termes de la suite et voir se mettre en place une série géométrique de raison a . Mais le mieux est d'effectuer une translation nettement compréhensible par représentation graphique.

On met de côté le cas trivial $a = 1$ qui fait de la suite u une suite arithmétique, dont le terme général est $u_0 + n \cdot b$.

On cherche déjà la valeur de u_0 qui rend la suite constante (point fixe) en résolvant l'équation $\alpha = a \cdot \alpha + b$.

On crée alors la suite translatée : $v_n = u_n - \frac{b}{1-a}$. On montre alors rapidement : $v_{n+1} = a \cdot v_n$. La suite v est géométrique de raison a . On connaît son terme général : $v_n = a \cdot v_0$ et on revient à u_n :

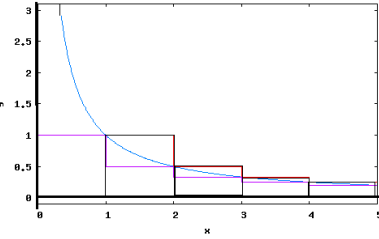
$u_n = a^n \cdot \left(u_0 - \frac{b}{1-a}\right) + \frac{b}{1-a}$ ou même $u_n = a^n \cdot u_0 - b \cdot \frac{1-a^n}{1-a}$

On peut retenir directement la formule générale qui n'est pas si compliquée. Mais il faut aussi connaître la méthode et ne pas se contenter d'une preuve par récurrence.

○ Démontrer que $\sum_{k=1}^n \frac{1}{k}$ est équivalent à $\ln(n)$ quand n tend vers l'infini par une comparaison série intégrale.

On montre : $\ln(n+1) \leq \sum_{k=1}^n \frac{1}{k} \leq \ln(n) + 1$ par ladite comparaison série-intégrale dont la preuve première est visuelle.

Proprement, on se donne un entier naturel k et on étudie $1/t$ quand t décrit $[k, k+1]$: $\frac{1}{k+1} \leq \frac{1}{t} \leq \frac{1}{k}$. On intègre de k à $k+1$: $\frac{1}{k+1} \leq \ln(k+1) - \ln(k) \leq \frac{1}{k}$.



On somme ces relations de 1 à n ou même de 1 à $n-1$. Le terme du milieu télescope. On isole si nécessaire dans la somme le terme $k=0$ pour ne pas faire intervenir $\int_0^1 \frac{dt}{t}$.

On a bien l'encadrement $\ln(n+1) \leq H_n \leq \ln(n) + 1$ en notant H_n la série harmonique.

On divise par $\ln(n)$ strictement positif : $\frac{\ln(n+1)}{\ln(n)} \leq H_n \leq 1 + \frac{1}{\ln(n)}$.

On écrit le terme de gauche $\frac{\ln(n) + \ln\left(1 + \frac{1}{n}\right)}{\ln(n)}$ pour le voir tendre vers 1 quand n tend vers l'infini.

Par encadrement, $\frac{H_n}{\ln(n)}$ converge vers 1 quand n tend vers l'infini. C'est la définition de l'équivalent.

On peut aussi montrer un résultat plus fin : $H_n - \ln(n)$ est non seulement un $o(\ln(n))$ mais converge même vers une limite finie quand n tend vers l'infini.

○ Donner les relations coefficients racines pour un polynôme de degré 2, de degré 3 et de degré 4. Exprimer la somme et le produit des racines d'un polynôme de degré n à l'aide de ses coefficients (polynôme de coefficient dominant pouvant être autre que 1). Retrouver la somme des carrés des racines et la somme de leurs inverses.

En identifiant le polynôme écrit sur la base canonique et le polynôme écrit sous forme factorisée, on obtient les relations coefficients racines

base canonique	forme factorisée	relations
$a.X^2 + b.X + c$	$a.(X - \alpha).(X - \beta)$	$\alpha + \beta = -b/a$ $\alpha \times \beta = c/a$
$a.X^3 + b.X^2 + c.X + d$	$a.(X - \alpha).(X - \beta).(X - \gamma)$	$\alpha + \beta + \gamma = -b/a$ $\beta.\gamma + \alpha.\gamma + \alpha.\beta = c/a$ $\alpha \times \beta \times \gamma = -d/a$
$a.X^4 + b.X^3 + c.X^2 + d.X + e$	$a.(X - \alpha).(X - \beta).(X - \gamma).(X - \delta)$	$\alpha + \beta + \gamma + \delta = -b/a$ $\alpha.\beta + \alpha.\gamma + \alpha.\delta + \beta.\gamma + \beta.\delta + \gamma.\delta = c/a$ $\beta.\gamma.\delta + \alpha.\gamma.\delta + \alpha.\beta.\delta + \alpha.\beta.\gamma = -d/a$ $\alpha \times \beta \times \gamma \times \delta = e/a$

On retiendra pour les polynômes unitaires (coefficient dominant égal à 1) :

$$\boxed{X^2 - S.X + P} \quad \boxed{X^3 - S.X^2 + D.X - P} \quad \boxed{X^4 - S.X^3 + D.X^2 - T.X + P}$$

Pour les polynômes non unitaires, il ne faudra pas oublier la division par le coefficient dominant.

Pour un polynôme de degré n quelconque, on identifie $a_n.X^n + a_{n-1}.X^{n-1} + \dots + a_0$ avec $a_n.(X - \alpha_1) \dots (X - \alpha_n)$ et on a pour le moins $\sum_{k=1}^n \alpha_k = -\frac{a_{n-1}}{a_n}$ et $\prod_{k=1}^n \alpha_k = (-1)^n \frac{a_0}{a_n}$.

On n'oublie pas le $(-1)^n$ dans le produit des racines, issu du calcul en 0 de $(X - \alpha_1) \dots (X - \alpha_n)$.

Les autres relations sont de la forme $\sum_{j_1 < j_2 < \dots < j_k} \alpha_{j_1} \alpha_{j_2} \dots \alpha_{j_k} = (-1)^k \frac{a_{n-k}}{a_n}$ qu'on démontre par récurrence sur le nombre de racines (et donc le degré du polynôme).

En particulier pour les produits à deux termes : $\sum_{i < j} \alpha_i \cdot \alpha_j = \frac{a_{n-2}}{a_n}$ et pour les produits à $n - 1$ termes

$$: \sum_{i=1}^n \prod_{j \neq i} \alpha_j = \frac{a_1}{(-1)^{n-1} \cdot a_n}.$$

En divisant cette dernière formule par le produit des racines (*en supposant qu'aucune n'est nulle*) :

$$\frac{1}{\alpha_1} + \frac{1}{\alpha_2} + \dots + \frac{1}{\alpha_n} = -\frac{a_{n-1}}{a_n} \text{ (qu'on peut voir aussi avec les racines du polynôme "renversé" } a_n + a_{n-1} \cdot X + \dots + a_1 \cdot X^{n-1} + a_0 \cdot X^n \text{).}$$

D'autre part, en élevant au carré $(\alpha_1 + \dots + \alpha_n)^2$ et en soustrayant les "double-produits" :

$$dsp(\alpha_1)^2 + \dots + (\alpha_n)^2 = \left(\frac{a_{n-1}}{a_n}\right)^2 - 2 \cdot \frac{a_{n-2}}{a_n}.$$

○ Inverser une matrice de taille 2 de déterminant non nul. Résoudre un système linéaire de deux équations à deux inconnues.

On suppose que la matrice $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$ a un déterminant non nul : $a \cdot d - b \cdot c \neq 0$.

On vérifie par simple calcul que $\frac{1}{a \cdot d - b \cdot c} \cdot \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}$ est l'inverse (à droite comme à gauche) de $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$.

On retient : échanger les deux coefficients de la diagonale, change le signe des deux autres, et enfin diviser par le déterminant.

On note que cette formule coïncide avec la formule par transposée de la comatrice.

On écrit le système $\begin{cases} a \cdot x + b \cdot y = \alpha \\ c \cdot x + d \cdot y = \beta \end{cases}$ sous la forme $\begin{pmatrix} a & b \\ c & d \end{pmatrix} \cdot \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \alpha \\ \beta \end{pmatrix}$ et on le résout $\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} a & b \\ c & d \end{pmatrix}^{-1} \cdot \begin{pmatrix} \alpha \\ \beta \end{pmatrix}$.

On a donc sous forme compactée : $x = \frac{\begin{vmatrix} \alpha & b \\ \beta & d \end{vmatrix}}{\begin{vmatrix} a & b \\ c & d \end{vmatrix}}$ et $y = \frac{\begin{vmatrix} a & \alpha \\ c & \beta \end{vmatrix}}{\begin{vmatrix} a & b \\ c & d \end{vmatrix}}$ qu'on appelle formules de Cramer.

Il est d'ailleurs aussi facile de résoudre par combinaisons $\begin{cases} a \cdot x + b \cdot y = \alpha \\ c \cdot x + d \cdot y = \beta \\ \delta \cdot x = d \cdot \alpha - b \cdot \beta \end{cases} \begin{matrix} \times d \\ \times (-b) \end{matrix}$

avec $\delta = a \cdot d - b \cdot c$.

La résolution par substitutions est évidemment possible, mais contraire à l'approche naturelle sauf si on a été formé à l'approche très réductrice de l'enseignement des mathématiques dans le secondaire.

○ Donner la définition de $(\vec{u}_1, \dots, \vec{u}_p)$ engendre l'espace vectoriel $(E, +, \cdot)$. Donner la définition de "famille génératrice minimale".

La définition purement rhétorique est simple : "tout vecteur de E est combinaison linéaire des vecteurs de la famille".

On quantifie : $\forall \vec{u} \in E, \exists (\alpha_1, \dots, \alpha_p) \in \mathbb{R}^p, \vec{u} = \sum_{k=1}^p \alpha_k \cdot \vec{u}_k$.

Si la famille génératrice est infinie, la définition est "tout vecteur de $(E, +, \cdot)$ est combinaison linéaire d'une sous-famille finie de la grande famille".

La locution "famille génératrice" n'a pas de sens si on ne précise pas de quoi. Toute famille est en effet génératrice de l'espace vectoriel qu'elle engendre.

On pourra tout aussi bien dire "la famille est génératrice de $(E, +, \cdot)$ " que "la famille engendre $(E, +, \cdot)$ ".

Une famille génératrice minimale est une famille génératrice dont aucune sous-famille stricte n'est génératrice de $(E, +, \cdot)$.

La famille est génératrice, et dès qu'on ôte un vecteur, elle cesse de l'être.

On montre (par l'absurde) qu'il y a équivalence entre "famille génératrice minimale de $(E, +, \cdot)$ " et "famille génératrice de $(E, +, \cdot)$ et libre".

On prend une famille génératrice minimale. Elle est alors génératrice. Supposons la non libre. Alors elle est liée, et l'un des vecteurs (disons $\vec{e}_k$) est combinaison des autres. On peut alors l'enlever de la

famille, et elle reste génératrice (dans les combinaisons de $(\vec{e}_1, \dots, \vec{e}_n)$, on remplace $\vec{e}_k$ par sa combinaison des autres), et ceci contredit le caractère minimal.

On prend une famille libre et génératrice. Elle est alors génératrice. Si elle n'est pas génératrice minimale, on peut enlever l'un d'entre eux (disons $\vec{e}_k$) et la famille réduite reste génératrice. Tout vecteur de $(E, +, \cdot)$ s'écrit alors comme combinaison de la famille réduite. En particulier $\vec{e}_k$ qui est alors combinaison des autres. Il s'ensuit que la famille est liée, ce qui contredit l'hypothèse.

○ Montrer qu'un polynôme $P(X)$ est factorisable par $(X - a)$ si et seulement si a est racine de $P(X)$. Montrer que a est racine double si et seulement si $P(X)$ est factorisable par $(X - a)^d$. Équivalence avec $P(a) = P'(a) = \dots = P^{(d-1)}(a) = 0$.

La définition de " a est racine de $P(X)$ est $P(a) = 0$ ".

Si P se factorise par $(X - a)$, il s'écrit $P(X) = (X - a).Q(X)$ avec Q polynôme (de degré $\deg(P) - 1$). On remplace alors $P(a) = (a - a).Q(a) = 0$.

On suppose maintenant $P(a) = 0$. On commence par poser la division euclidienne de P par $X - a$: quotient Q et reste R : $P(X) = (X - a).Q(X) + R(X)$. Par définition de la division euclidienne, le degré de R est strictement plus petit que celui de $X - a$. C'est donc que $R(X)$ est en fait un polynôme constant (éventuellement nul). On remplace alors : $P(a) = (a - a).Q(a) + R(a)$. Par hypothèse, $P(a)$ est nul, il s'ensuit que $R(a)$ est nul. Comme le polynôme est constant, il est nul. On reporte : $P(X) = (X - a).Q(X)$. Le polynôme se factorise.

On rappelle qu'avec les polynômes, on peut écrire indifféremment P ou $P(X)$, car X est une variable formelle ; au contraire des fonctions où il ne faut pas confondre f et $f(x)$.

On passe au cas de la racine de multiplicité d .

On suppose que P se factorise sous la forme $P(X) = (X - a)^d.Q(X)$. On calcule alors $P(a) = 0$. On peut même dériver k fois avec k strictement plus petit que d . On utilise la formule de Leibniz :

$$P^{(k)}(X) = \sum_{i=0}^k \binom{k}{i}.Q^{(k-i)}(X). \frac{d!. (X - a)^{d-i}}{(d-i)!} \quad (\text{on rappelle } (X^d)^{(i)} = \frac{d!}{(d-i)!}.X^{d-i}).$$

On calcule en a .

Les exposants sont strictement positifs, il reste $P^{(k)}(a) = 0$. Comme ceci est vrai pour k de 0 à $d - 1$, on reconnaît que a est racine de P avec pour le moins multiplicité d .

On suppose maintenant que $P(a)$ jusqu'à $P^{(d-1)}(a)$ sont tous nuls. On écrit une formule de Taylor pour le polynôme P à l'ordre p où p est le degré de P . Comme la formule de Taylor tombe juste pour un polynôme écrit à un ordre suffisant :

$$P(X) = \sum_{i=0}^p \frac{(X - a)^i.P^{(i)}(a)}{i!}$$

Dans cette somme, les premiers termes sont nuls, on peut factoriser :

$$P(X) = \sum_{i=d}^p \frac{(X - a)^i.P^{(i)}(a)}{i!} = (X - a)^d \cdot \sum_{i=d}^p \frac{(X - a)^{i-d}.P^{(i)}(a)}{i!}.$$

Il ne reste plus qu'à dire que $\sum_{i=d}^p \frac{(X - a)^{i-d}.P^{(i)}(a)}{i!}$ est bien un polynôme.

On notera que si l'on veut exprimer que a est racine de P avec multiplicité exactement égale à d , on dispose de plusieurs formulations :

- $P(a) = P'(a) = \dots = P^{(d-1)}(a) = 0 \neq P^{(d)}(a)$
- $\exists Q \in \mathbb{R}[X], P(X) = (X - a)^d.Q(X)$ et $Q(a) \neq 0$
- $\text{p.g.c.d.}(P(X), (X - a)^{d+1}) = (X - a)^d$

○ Calculer le terme général d'une matrice $A.B$, produit d'une matrice A à n lignes et p colonnes et d'une matrice B à p lignes et q colonnes. Démontrer l'associativité du produit matriciel.

Le terme de ligne i et colonne k de la matrice $A.B$ (produit calculable car les formats sont compatibles)

$$\text{est } \sum_{j=1}^p a_i^j.b_j^k.$$

On prend trois matrices dont les formats sont compatibles :

matrice	A	B	C	$A.B$	$B.C$	$(A.B).C$ et $A.(B.C)$
lignes	p	q	r	p	q	p
colonnes	q	r	s	r	s	s
terme général	a_i^j	b_j^k	c_k^l	α_i^k	β_j^l	à calculer

On détermine le terme général de $(A.B).C$ ($i \leq p$ et $l \leq s$), puis de $A.(B.C)$:

$$\sum_{k=1}^r \alpha_i^k . c_k^l = \sum_{k=1}^r \left(\sum_{j=1}^q a_i^j . b_j^k \right) . c_k^l = \sum_{\substack{k \leq r \\ j \leq q}} a_i^j . b_j^k . c_k^l$$

$$\sum_{j=1}^q a_i^j . \beta_j^l = \sum_{k=1}^r a_i^j . \left(\sum_{k=1}^r b_j^k . c_k^l \right) = \sum_{\substack{k \leq r \\ j \leq q}} a_i^j . b_j^k . c_k^l.$$

Il y a égalité des termes généraux (et des formats) ; les deux matrices sont égales.

○ Donner la définition de l'espérance et de la variance d'une variable aléatoire réelle. Montrer que la variance est toujours positive (cas de nullité).

L'espérance est simplement la valeur moyenne : $\sum_i a_i . P(A = a_i)$ où les a_i sont les valeurs prises par la variable aléatoire.

L'opérateur est linéaire, et positif (l'espérance d'une variable aléatoire positive est positive).

La variance est l'espérance de l'écart quadratique à la moyenne (espérance du carré de $A - E(A)$) : $E((A - E(A))^2)$.

Sous cette forme, c'est évidemment un réel positif (somme de carrés pondérés par des nombres positifs).

Les professeurs de mathématiques ont trouvé le moyen de donner un nom à la formule $Var(A) = E(A^2) - E(A)^2$.

On l'a en posant un court instant $E(A) = \mu$ et en développant :

$E((A - \mu)^2) = E(A^2 - 2.\mu.A + \mu^2) = E(A^2) - 2.\mu.E(A) + \mu^2$. Il ne reste qu'à noter $\mu = E(A)$.

Sous cette forme, on a $Var(A) = \sum_i (a_i)^2 . P(A = a_i) - \left(\sum_i a_i . P(A = a_i) \right)^2$. On peut alors prendre

deux vecteurs de $\mathbb{R}^n$ (est égal au nombre d'états de la variable aléatoire réelle) :

$U = (a_1 . \sqrt{P(A = a_1)}, a_2 . \sqrt{P(A = a_2)}, \dots, a_n . \sqrt{P(A = a_n)})$

et $V = (\sqrt{P(A = a_1)}, \sqrt{P(A = a_2)}, \dots, \sqrt{P(A = a_n)})$

Le premier vecteur a pour norme $\sqrt{E(A^2)}$.

Le second vecteur a pour norme 1 (la somme des probabilités vaut 1)

Le produit scalaire des deux vecteurs est $E(A)$.

L'inégalité de Cauchy Schwarz donne $(E(A))^2 \leq E(A^2) . 1$. C'est la positivité de la variance.

La variance est nulle si et seulement si l'espérance de la variable aléatoire positive $(A - E(A))^2$ est nulle. Ceci force la variable à être nulle : $A - E(A) = 0$. La variable aléatoire A est constante.

Avec la démonstration du type Cauchy-Schwarz, la nullité de la variable aléatoire donne la colinéarité des deux vecteurs U et V , qui se traduit encore par "variable aléatoire constante".

○ Démontrer les formules $P(X) = \sum_{k=0}^d \frac{P^{(k)}(a)}{k!} . (X - a)^k$ et $P(X + a) = \sum_{k=0}^d \frac{P^{(k)}(a)}{k!} . X^k$ pour un polynôme de degré inférieur ou égal à d .

Ces deux formules sont la même, quitte à simplement changer de variable.

Si l'on fait de l'analyse, il suffit d'écrire une formule de Taylor avec reste intégrale entre a et $a + X$:

$$P(a + X) = \sum_{k=0}^n \frac{P^{(k)}(a)}{k!} . X^k + \frac{X^{d+1}}{d!} . \int_{t=0}^1 (1-t)^d . P^{(d+1)}(a + t.X) . dt$$

Mais comme on a dérivé le polynôme $d+1$ fois dans le reste intégrale, il ne reste plus rien ; "la formule tombe juste".

Si l'on travaille sur un autre corps sur lequel on ne peut pas faire d'analyse et écrire une formule de Taylor, il suffit de faire de l'algèbre linéaire.

La formule $P(X+a) = \sum_{k=0}^d \frac{P^{(k)}(a)}{k!} \cdot X^k$ est vraie pour $P = 1$, $P = X$, $P = X^2$ et ainsi de suite jusqu'à $P = X^d$.

C'est en effet la formule du binôme de Newton

$$(X+a)^p = \sum_{k=0}^p \frac{p!}{k!(p-k)!} \cdot a^{p-k} \cdot X^k = \sum_{k=0}^p \frac{1}{k!} \cdot ((X^p)^{(k)}(a)) \cdot X^k \text{ et les termes de } p+1 \text{ à } d \text{ sont nuls.}$$

On étend à tous les polynômes de degré inférieur ou égal à d par linéarité.

○ Démontrer le théorème des valeurs intermédiaires pour une application numérique continue sur un intervalle. Démontrer que l'image d'un intervalle par une application numérique continue est un intervalle. Démontrer qu'une application continue sur un segment qui ne s'annule pas est de signe constant.

Dans le théorème, les mots "*continue*", "*intervalle*" et "*numérique*" (*c'est à dire à valeurs dans $\mathbb{R}$*) sont capitaux comme le montrent des contre-exemples classiques ($x \mapsto [x]$, $x \mapsto 1/x$, $x \mapsto e^{i \cdot x}$).

On prend f continue de $[a, b]$ dans $\mathbb{R}$ avec $f(a)$ et $f(b)$ de signes opposés. Quitte à prendre $-f$ à la place de f on va supposer $f(a) < 0$ et $f(b) > 0$. Il faut montrer que f s'annule au moins une fois sur $[a, b]$.

On a une démonstration par "*borne supérieure*".

On pose $A = \{x \in [a, b] \mid f(x) < 0\}$. C'est une partie de $\mathbb{R}$ par définition. Elle est non vide (*elle contient a*) et majorée (par b). Elle admet donc une borne supérieure qu'on va noter c (*dernier instant où f est négative... ou nulle*).

Par définition de la borne supérieure, il existe une suite de points de A (c_n) qui converge vers c . Par continuité de f en tout point, la suite $f(c_n)$ converge vers $f(c)$. Mais chaque $f(c_n)$ est négatif (*appartenance à A*). Par passage à la limite, $f(c)$ est négatif ou nul (*les inégalités strictes sont larges par passage à la limite*).

Comme $f(c)$ est négatif ou nul, c ne peut pas être égal à b . Comme b était un majorant naturel de A , c'est que c est strictement plus petit que b . La suite $\left(\frac{n \cdot c + b}{n+1}\right)$ est donc une suite de réels plus grands que c (et plus petits que b) qui converge vers c . Ces réels sont donc hors de A , leurs images sont positives ou nulles. Par passage à la limite (*continuité de f*) les $f\left(\frac{n \cdot c + b}{n+1}\right)$ convergent vers $f(c)$ qui est maintenant positif ou nul.

Par antisymétrie, $f(c)$ est nul.

On peut aussi faire appel aux suites adjacentes.

On définit $a_0 = a$ et $b_0 = b$ puis pour tout n : $c_n = \frac{a_n + b_n}{2}$.

Si $f(c_n)$ est négatif, on pose $a_{n+1} = c_n$ et $b_{n+1} = b_n$.

Si $f(c_n)$ est positif, on pose $a_{n+1} = a_n$ et $b_{n+1} = c_n$.

On peut alors montrer pour tout n : $f(a_n) \leq 0$ et $f(b_n) \geq 0$ mais aussi $a_0 \leq a_1 \leq \dots \leq a_n \leq a_{n+1} \leq b_{n+1} \leq b_n \leq \dots \leq b_0$.

De plus, $b_n - a_n = \frac{b-a}{2^n}$.

Les deux suites (a_n) et (b_n) (*respectivement croissante majorée et décroissante minorée*) convergent, vers la même limite, qu'on note c . Par passage à la limite dans $f(a_n) \leq 0$, on déduit que $f(c)$ est négatif ou nul. Par passage à la limite dans $f(b_n) \geq 0$, $f(c)$ est positif ou nul. Par antisymétrie, $f(c)$ est nul.

On a donc prouvé : une application numérique continue qui change de signe sur un intervalle s'annule au moins une fois.

Par contraposée, on obtient qu'une application numérique continue sur un segment qui ne s'annule pas reste de signe constant.

On montre aussi qu'une application numérique continue sur un segment $[a, b]$ atteint au moins une fois chaque valeur γ entre $f(a)$ et $f(b)$. Soit en effet γ entre $f(a)$ et $f(b)$; on construit l'application $x \mapsto f(x) - \gamma$. Elle est encore continue sur $[a, b]$, mais elle prend deux signes différents en a et b . Elle s'annule au moins une fois en un point c de $[a, b]$ qui vérifie donc $f(c) = \gamma$.

On prend un intervalle I et une application continue de I dans $\mathbb{R}$. Il faut montrer que $f(I)$ est encore un intervalle. Mais comme il y a plusieurs types d'intervalle (*ouvert, fermé, semi-ouvert, infini d'un côté, de l'autre, des deux...*) il faut prendre la définition la plus générale : tout point entre deux points de I est encore dans I .

On prend deux éléments α et β de I ; on les écrit $\alpha = f(a)$ et $\beta = f(b)$. Comme a et b sont dans l'intervalle I , le segment $[a, b]$ est inclus dans I . Toute valeur entre $f(a)$ et $f(b)$ est donc atteinte au moins une fois sur $[a, b]$ donc sur I .

En combinant théorème des valeurs intermédiaires et compacité, on montre que l'image d'un segment $[a, b]$ par une application continue est encore un segment (*pas forcément $[f(a), f(b)]$ si l'on n'a pas d'hypothèse de monotonie...*).

○ Montrer que dans un espace vectoriel engendré par une famille finie, toutes les bases ont le même cardinal.

On suppose donc $(E, +, \cdot)$ engendré par une famille $(\vec{a}_1, \dots, \vec{a}_n)$ (ce n'est pas forcément une base, cette famille peut être liée).

Par théorème fondamental, les familles libres ne peuvent pas avoir plus de n vecteurs.

Il s'ensuit que les bases (si il y en a) ont un cardinal fini, moindre que n .

Il existe des bases. On part de la famille vide, qui est libre. Tant qu'on peut l'agrandir en famille libre, on agrandit. On est obligé de s'arrêter car le cardinal des familles libres est plafonné. Quand on s'arrête, on a une famille libre maximale. C'est donc une base.

On prend alors deux bases de $(E, +, \cdot)$: $(\vec{e}_1, \dots, \vec{e}_p)$ et $(\vec{\varepsilon}_1, \dots, \vec{\varepsilon}_q)$. Comme le cardinal des familles liées est majoré par le cardinal des familles génératrices, on peut écrire

$(\vec{e}_1, \dots, \vec{e}_p)$ génératrice	$(\vec{\varepsilon}_1, \dots, \vec{\varepsilon}_q)$ libre	$q \leq p$
$(\vec{\varepsilon}_1, \dots, \vec{\varepsilon}_q)$ génératrice	$(\vec{e}_1, \dots, \vec{e}_p)$	$p \leq q$

Les deux bases ont le même cardinal, c'est lui qu'on appelle la dimension de l'espace vectoriel.

Par usage un petit peu plus fin, on peut ensuite montrer

hypothèse de dimension	hypothèse avec bon cardinal	conclusion
$(\vec{e}_1, \dots, \vec{e}_n)$ base	$(\vec{a}_1, \dots, \vec{a}_n)$ libre	$(\vec{a}_1, \dots, \vec{a}_n)$ base
$(\vec{e}_1, \dots, \vec{e}_n)$ base	$(\vec{a}_1, \dots, \vec{a}_n)$ génératrice	$(\vec{a}_1, \dots, \vec{a}_n)$ base

Dans ce tableau, il faut évidemment lire "base" en "base de $(E, +, \cdot)$ " et "base" en "base de $(E, +, \cdot)$ ".

○ Donner la définition de " f est uniformément continue de I dans $\mathbb{R}$ (ou $\mathbb{C}$). Établir les stabilités par addition, multiplication par un réel et composition.

La définition de l'uniforme continuité est plus exigeante que celle de la continuité en tout point :

$$\forall \varepsilon > 0, \exists \eta_\varepsilon > 0, \forall (a, b) \in I^2, |b - a| \leq \eta_\varepsilon \Rightarrow |f(b) - f(a)| \leq \varepsilon.$$

Il découle de cette définition que toute application uniformément continue est continue en tout point. L'uniforme continuité n'a de sens que sur un domaine entier, en général un intervalle, mais il peut être intéressant d'avoir l'uniforme continuité sur une partie dense pour l'étendre au domaine entier.

Enfin, les applications lipschitziennes (rapport K) sont uniformément continue sur leur domaine (en prenant $\eta_\varepsilon = \frac{\varepsilon}{K}$).

On prend f et g uniformément continues de I dans $\mathbb{R}$, et λ dans $\mathbb{R}$:

$$\forall \varepsilon > 0, \exists \eta_\varepsilon > 0, \forall (a, b) \in I^2, |b - a| \leq \eta_\varepsilon \Rightarrow |f(b) - f(a)| \leq \varepsilon$$

$$\forall \varepsilon > 0, \exists \mu_\varepsilon > 0, \forall (a, b) \in I^2, |b - a| \leq \mu_\varepsilon \Rightarrow |g(b) - g(a)| \leq \varepsilon$$

On montre alors :

$$\forall \varepsilon > 0, \exists \sigma_\varepsilon > 0, \forall (a, b) \in I^2, |b - a| \leq \sigma_\varepsilon \Rightarrow |(f(b) + g(b)) - (f(a) + g(a))| \leq \varepsilon$$

$$\text{et } \forall \varepsilon > 0, \exists \pi_\varepsilon > 0, \forall (a, b) \in I^2, |b - a| \leq \pi_\varepsilon \Rightarrow |\lambda \cdot f(b) - \lambda \cdot f(a)| \leq \varepsilon$$

Pour λ nul, il n'y a rien à prouver, sinon on prend $\pi_\varepsilon = \eta_{\varepsilon/|\lambda|}$.

Et d'autre part, on prend $\sigma_\varepsilon = \min(\eta_{\varepsilon/2}, \mu_{\varepsilon/2})$. On vérifie pour a et b dans I :

$$|b - a| \leq \min(\eta_{\varepsilon/2}, \mu_{\varepsilon/2}) \Rightarrow \begin{array}{l} |b - a| \leq \eta_{\varepsilon/2} \Rightarrow |f(b) - f(a)| \leq \varepsilon/2 \\ |b - a| \leq \mu_{\varepsilon/2} \Rightarrow |g(b) - g(a)| \leq \varepsilon/2 \end{array} \Rightarrow |(f(b) + g(b)) - (f(a) + g(a))| \leq \varepsilon$$

Pour la composition, on suppose φ uniformément continue de $f(I)$ dans $\mathbb{R}$:

$$\forall \varepsilon > 0, \exists \tau_\varepsilon > 0, \forall (\alpha, \beta) \in f(I)^2, |\beta - \alpha| \leq \tau_\varepsilon \Rightarrow |\varphi(\beta) - \varphi(\alpha)| \leq \varepsilon$$

$$\text{On montre alors } |b - a| \leq \eta_{\tau_\varepsilon} \Rightarrow |f(b) - f(a)| \leq \tau_\varepsilon \Rightarrow |\varphi(f(b)) - \varphi(f(a))| \leq \varepsilon.$$

Même si il n'est pas sain ni cohérent d'enchaîner les implications, on reconnaît l'uniforme continuité de $\varphi \circ f$.

Le produit de deux applications uniformément continues ne l'est plus forcément comme l'atteste l'exemple $Id \times Id$.

C'est en revanche un exercice classique que de montrer que le produit de deux applications bornées uniformément continues est encore uniformément continu.

○ Montrer que si l'on agrandit une famille libre $(\vec{a}_1, \dots, \vec{a}_p)$ avec un vecteur $\vec{v}$, elle reste libre si et seulement si $\vec{v}$ n'est pas dans l'espace vectoriel engendré par $(\vec{a}_1, \dots, \vec{a}_p)$ (lemme d'agrandissement des familles libres).

Il y a une hypothèse faite une fois pour toutes : $(\vec{a}_1, \dots, \vec{a}_p)$ est libre (si l'on a $\alpha_1 \vec{a}_1 + \dots + \alpha_p \vec{a}_p = \vec{0}$ alors les α_i sont nuls). Il faut alors prouver une équivalence.

On suppose que $\vec{v}$ n'est pas combinaison de $(\vec{a}_1, \dots, \vec{a}_p)$. Il faut montrer que $(\vec{a}_1, \dots, \vec{a}_p, \vec{v})$ est encore libre. On prend donc $p+1$ réels α_1 à α_p et β , et on suppose $\alpha_1 \vec{a}_1 + \dots + \alpha_p \vec{a}_p + \beta \vec{v} = \vec{0}$. On

veut montrer que les α_i et β sont nuls. Si β est non nul, on isole : $\vec{v} = \sum_{i=1}^p \frac{-\alpha_i}{\beta} \vec{a}_i$, ce qui contredit

l'hypothèse " $\vec{v}$ n'est pas combinaison de $(\vec{a}_1, \dots, \vec{a}_p)$ ". β étant nul, on reporte : $\alpha_1 \vec{a}_1 + \dots + \alpha_p \vec{a}_p = \vec{0}$. Par liberté prise en hypothèse initiale les α_i aussi sont nuls.

On suppose que la nouvelle famille $(\vec{a}_1, \dots, \vec{a}_p, \vec{v})$ est libre. Aucun vecteur n'est combinaison des autres. En particulier, $\vec{v}$ n'est pas combinaison des $\vec{a}_i$.

○ Définir la covariance d'un couple de variables aléatoires (A, B) . Montrer qu'elle est nulle si les deux variables sont indépendantes.

On sait que si A est une variable aléatoire prenant des valeurs a_i et B une variable aléatoire prenant des valeurs b_j , alors $A.B$ est une variable aléatoire pouvant prendre pour valeurs les $a_i.b_j$. On définit alors la covariance : $Cov(A, B) = E(A.B) - E(A).E(B)$.

C'est une définition proche de l'algèbre bilinéaire et des produits scalaires. On note que $Cov(A, A)$ est la variance de A .

Cette forme est effectivement symétrique, linéaire par rapport à chacune des variables, et positive, puisque la variance est positive ou nulle.

La formule explicite est $\left(\sum_{\substack{i \in I \\ j \in J}} a_i.b_j.P(A = a_i, B = b_j) \right) - \left(\sum_{i \in I} a_i.P(A = a_i) \right) \cdot \left(\sum_{j \in J} b_j.P(B = b_j) \right)$.

Si les variables sont indépendantes, la première somme vaut $\sum_{\substack{i \in I \\ j \in J}} a_i.b_j.P(A = a_i).P(B = b_j)$ et coïncide

avec le produit de $\sum_{i \in I} a_i.P(A = a_i)$ et $\sum_{j \in J} b_j.P(B = b_j)$. La covariance est nulle.

Il n'y a pas de réciproque. Deux variables non indépendantes peuvent avoir une covariance nulle (on peut construire un contre-exemple).

On note toutefois que si on a deux variables à deux états (du type variable de Bernoulli), la nullité de la covariance entraîne l'indépendance (exercice classique).

○ Montrer que la dimension d'un sous-espace vectoriel est toujours plus petite que la dimension de l'espace vectoriel. Cas d'égalité.

On se place dans un espace vectoriel $(E, +, \cdot)$. On le suppose de dimension finie, sinon la question n'a guère de sens. On le munit donc d'une base $(\vec{e}_1, \dots, \vec{e}_n)$.

On prend un sous-espace vectoriel F . Il ne peut pas être de dimension infinie, sinon il contiendrait des familles libres aussi grandes qu'on veut, donc par exemple de cardinal $n+1$. Or, une famille libre de F reste une famille libre de E , et on a une contradiction alors.

Si l'on prend ensuite une base $(\vec{\varepsilon}_1, \dots, \vec{\varepsilon}_p)$ de $(F, +, \cdot)$. Elle engendre F , mais surtout elle est libre. La liberté ne dépendant pas de l'espace dans lequel on travaille, elle reste libre dans E , et son cardinal ne peut pas dépasser la dimension de $(E, +, \cdot)$. On a donc bien $p \leq n$, c'est à dire $\dim(F) \leq \dim(E)$.

Si l'on suppose en plus $\dim(F) = \dim(E)$, alors la base de F $(\vec{\varepsilon}_1, \dots, \vec{\varepsilon}_n)$ est libre dans $(E, +, \cdot)$. Mais comme son cardinal est la dimension de $(E, +, \cdot)$ c'est une base de $(E, +, \cdot)$. Tout vecteur $\vec{v}$ de E se

décompose sous la forme $\sum_{i \leq n} \alpha_i \cdot \vec{\varepsilon}_i$. Étant combinaison de vecteurs de F , ce vecteur générique de E est

dans f . On a l'inclusion de E dans F .

Avec l'hypothèse "sous-espace vectoriel", on avait déjà F inclus dans $(E, +, \cdot)$.

Un sous-espace vectoriel qui a la même dimension que l'espace vectoriel ambiant est égal à cet espace vectoriel.

○ Montrer que si f est lipschitzienne de I dans I avec un rapport de Lipschitz strictement plus petit que 1 (*application contractante*), alors elle admet un unique point fixe, obtenu par suite récurrente $u_{n+1} = f(u_n)$, vitesse de convergence (*démonstration hors programme en Sup*).

Attention, la formulation de la question est bien existence et unicité.

On note μ le rapport strictement plus petit que 1.

Pour l'unicité, on suppose qu'il y a deux points fixes distincts a et b . On a alors $f(a) = a$ et $f(b) = b$.

On soustrait et passe à la valeur absolue : $|a - b| = |f(a) - f(b)| \leq \mu \cdot |a - b| < |a - b|$. On aboutit à une contradiction (*que l'on n'aurait pas eue avec $a = b$ puisque la dernière inégalité n'aurait pas été valable*).

La suite $u_{n+1} = f(u_n)$ est bien définie puisque f va de I dans I .

On montre par récurrence évidente sur k : $|u_{k+1} - u_k| \leq \mu^k \cdot |u_1 - u_0|$ (l'hérédité utilise le caractère lipschitzien de f).

On définit deux séries : $U_n = \sum_{k=0}^{n-1} ((u_{k+1} - u_k) + |u_{k+1} - u_k|)$ et $V_n = \sum_{k=0}^{n-1} |u_{k+1} - u_k|$.

La série V est croissante car à termes positifs. On la majore :

$$V_n = \sum_{k=0}^{n-1} |u_{k+1} - u_k| \leq \sum_{k=0}^{n-1} \mu^k \cdot |u_1 - u_0| = \frac{1 - \mu^n}{1 - \mu} \cdot |u_1 - u_0| \leq \frac{1}{1 - \mu} \cdot |u_1 - u_0|.$$

Le majorant ne dépend plus de n , la série V croissante majorée converge vers son plus petit majorant.

Le terme général $(u_{k+1} - u_k) + |u_{k+1} - u_k|$ est positif (*de la forme $h + |h|$*) ; la série U est croissante.

On majore cette fois avec les mêmes arguments par $2 \cdot \frac{1}{1 - \mu} \cdot |u_1 - u_0|$. La série U converge.

Par soustraction, $(U_n - V_n)$ converge. Mais c'est alors une somme télescopique : $U_n - V_n =$

$$\sum_{k=0}^{n-1} (u_{k+1} - u_k) = u_n - u_0. \text{ Sa convergence entraîne celle de } u.$$

On note α sa limite.

La suite u vérifie $u_{n+1} = f(u_n)$ et converge. On passe à la limite (*f est continue car lipschitzienne*) : $\alpha = f(\alpha)$.

On vient de détecter un point fixe de f .

Si on veut la vitesse de convergence, on prend deux indices p et q avec p plus petit que q . On écrit

$$|u_q - u_p| = \left| \sum_{k=p}^{q-1} (u_{k+1} - u_k) \right| \leq \sum_{k=p}^{q-1} |u_{k+1} - u_k| \leq |a_1 - a_0| \cdot \sum_{k=p}^{q-1} \mu^k \leq |a_1 - a_0| \cdot \frac{\mu^p - \mu^q}{1 - \mu}.$$

On fait tendre q vers l'infini, maintenant qu'on sait que la suite converge.

Par passage à la limite : $|\alpha - u_p| \leq \mu^p \cdot \frac{|a_1 - a_0|}{1 - \mu}$. On dispose ainsi de la vitesse de convergence.

Si nécessaire, on majore $|a_1 - a_0|$ par la longueur de l'intervalle I si il est borné.

Si l'on travaille sur $\mathbb{C}$ et non sur $\mathbb{R}$ (*ou même sur des espaces plus compliqués*), on ne peut pas passer par des séries croissantes majorées. On reprend toutefois la majoration $|u_q - u_p| \leq |a_1 - a_0| \cdot \frac{\mu^p - \mu^q}{1 - \mu}$ pour déduire que la suite est de Cauchy :

$\forall \varepsilon > 0, \exists K_\varepsilon \in \mathbb{N}, \forall (p, q) \in \mathbb{N}^2, K_\varepsilon \leq p \leq q \Rightarrow |u_q - u_p| \leq \varepsilon$ (*il suffit d'avoir $|a_1 - a_0| \cdot \frac{\mu^p - \mu^q}{1 - \mu} \leq \varepsilon$, ce qui est explicite*).

En tant que suite de Cauchy dans $\mathbb{R}$ ou $\mathbb{C}$ (*ou dans un espace dit "complet"*), elle converge.

○ Donner la définition de (u_n) et (v_n) sont équivalentes en $+\infty$. Passage aux produits, aux quotients. Contre-exemples pour les sommes de suites équivalentes.

Deux suites u et v jamais nulles à partir d'un certain rang si et seulement si le quotient $\frac{u_n}{v_n}$ converge (*d'existence garantie à partir d'un certain rang*) converge vers 1 quand n tend vers l'infini. La vraie défi-

dition qui autorise les suites à s'annuler est "il existe une suite de limite 1 vérifiant $u_n = a_n \cdot v_n$ pour tout n .

On montre que cette relation est à la fois réflexive ($\frac{u_n}{v_n}$ vaut 1 et converge vers 1), symétrique (si $\frac{u_n}{v_n}$ converge vers 1 alors $\frac{v_n}{u_n}$ converge vers 1) et transitive (si $\frac{u_n}{v_n}$ et $\frac{v_n}{w_n}$ convergent vers 1 alors leur produit le fait aussi).

On suppose a_n équivalente à b_n et u_n équivalente à v_n ; alors $\frac{a_n \cdot u_n}{b_n \cdot v_n}$ et $\frac{a_n / u_n}{b_n / v_n}$ convergent vers 1, ce qui prouve que les équivalents passent aux produits et aux quotients. Par récurrence sur l'exposant k , on montre que $((a_n)^k)$ et $((b_n)^k)$ sont aussi équivalentes.

Attention, les équivalents ne passent qu'aux puissances d'exposant fixe. Par exemple : $\left(1 + \frac{1}{n}\right)$ et $\left(1 - \frac{1}{n}\right)$ sont équivalentes en $+\infty$, alors que $\left(\left(1 + \frac{1}{n}\right)^n\right)$ et $\left(\left(1 - \frac{1}{n}\right)^n\right)$ ne le sont pas.

De même, il ne faut pas soustraire des équivalents de même ordre. Par exemple $(n^2 + n)$ est équivalente à (n^2) , de même que (n^2) est équivalente à $(n^2 + 1)$; tandis que $(n^2 + n - n^2)$ n'est pas équivalent à $(n^2 - n^2 - 1)$.

En revanche, quand on a des équivalents d'ordres différents, on peut additionner, mais c'est sans grand intérêt. Si l'on a (a_n) équivalent à (b_n) , puis (u_n) équivalent à (v_n) mais négligeable devant (a_n) (et donc devant (b_n)) alors $(a_n + u_n)$ est équivalent à $(b_n + v_n)$ tout en sachant que l'on ne voit même plus u_n et v_n dans ces formules. De même, on peut additionner des équivalents de même ordre.

○ Donner la définition de "f est un homéomorphisme de I dans J" Condition suffisante sur $\mathbb{R}$ par croissance stricte..

Un homéomorphisme est une application continue bijective dont le réciproque est aussi continue.

L'existence d'homéomorphismes entre deux ensembles est une relation d'équivalence.

On peut reformuler

réflexivité	Id est un homéomorphisme de I dans I
symétrie	la réciproque d'un homéomorphisme de I dans J est un homéomorphisme de J dans I
transitivité	la composée de deux homéomorphismes est encore un homéomorphisme

La définition de "homéomorphisme" n'a de sens que si on précise de quoi dans quoi.

En Terminale, un théorème portant le nom de "théorème de la bijection" montrait que sous des hypothèses de continuité et stricte monotonie de f , une équation $f(x) = a$ (de paramètre a et d'inconnue x) admettait une unique solution x . Sa version dans l'enseignement supérieur porte le nom de théorème de l'homéomorphisme et donne le résultat plus puissant : la solution x dépend continuellement du paramètre a .

Pour les applications numériques (de $\mathbb{R}$ dans $\mathbb{R}$), on dispose de critères par stricte monotonie.

On a en effet une implication : toute application strictement croissante est injective.

On a l'implication réciproque avec une hypothèse en plus : toute application continue injective est forcément monotone (théorème classique faisant l'objet d'un autre point du bilan de compétence).

Les théorèmes de base de l'analyse permettent de dire que toute application numérique continue strictement monotone sur un intervalle I réalise un homéomorphisme de l'intervalle de départ I vers l'intervalle d'arrivée $f(I)$.

Soit I un intervalle de $\mathbb{R}$ (ouvert, fermé, semi ouvert, semi fermé, c'est sans importance). Par stricte croissance (en cas de stricte décroissance, le raisonnement est le même), l'application est injective. Par théorème des valeurs intermédiaires, l'ensemble image est un intervalle et on a une application réciproque (définie sans ambiguïté par injectivité). L'application réciproque f^{-1} est alors aussi monotone, de même sens de variation que f .

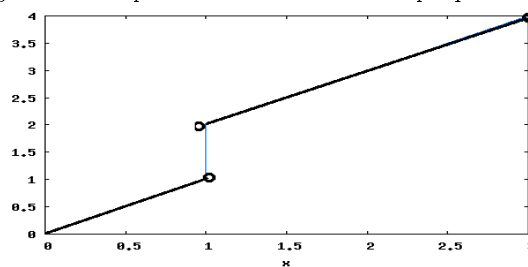
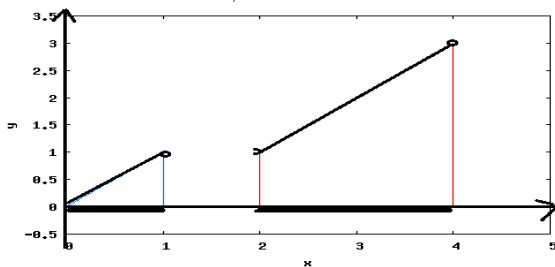
Il reste à montrer que l'application réciproque est continue aussi. On revient à la définition par continuité à droite et à gauche.

$$\forall \alpha \in f(I), \forall \varepsilon > 0, \exists \delta_\varepsilon > 0, \forall y \in f(I), \alpha \leq y \leq \alpha + \delta_\varepsilon \Rightarrow |f(\alpha) - f^{-1}(y)| \leq \varepsilon$$

$$\forall \alpha \in f(I), \forall \varepsilon > 0, \exists \gamma_\varepsilon > 0, \forall y \in f(I), \alpha - \gamma_\varepsilon \leq y \leq \alpha \Rightarrow |f(\alpha) - f^{-1}(y)| \leq \varepsilon$$

Il suffit, pour ε donné, de prendre $\delta_\varepsilon = f(f^{-1}(\alpha) + \varepsilon) - \alpha$ et $\gamma_\varepsilon = \alpha - f(f^{-1}(\alpha) - \varepsilon)$. Il ne reste plus qu'à passer par exemple de $\alpha \leq y \leq f(f^{-1}(\alpha) + \varepsilon)$ à $f^{-1}(\alpha) \leq f^{-1}(y) \leq f^{-1}(\alpha) + \varepsilon$ par stricte monotonie de f^{-1} .

Sans le mot intervalle, la continuité et la croissance ne garantissent pas la continuité de la réciproque.



○ Donner des définitions équivalentes de “ p est un projecteur de $(E, +, \cdot)$ ”. Exprimer sa matrice en dimension finie sur une base adaptée, calculer sa trace.

Un projecteur p de $(E, +, \cdot)$ est un endomorphisme (application linéaire de E dans lui même) vérifiant $p \circ p = p$.

Id et $\vec{u} \mapsto \vec{0}$ sont les deux projecteurs les plus simples et Id est le seul projecteur bijectif.

En développant $(Id - p) \circ (Id - p) = Id - 2p + p \circ p$, on montre que p est un projecteur si et seulement si $Id - p$ est aussi un projecteur. On note au passage : $Id - (Id - p) = p$, ce qui permet de dire que p et $Id - p$ sont associés.

Un cas particulier du théorème dit “des noyaux” montre que p est un projecteur si et seulement si on a $E = Ker(Id - p) \oplus Ker(p)$.

Supposons que p soit un projecteur ($p \circ p = p$ plus la linéarité).

Tout vecteur $\vec{u}$ se décompose sous la forme $\vec{u} = (p(\vec{u})) + (\vec{u} - p(\vec{u}))$.

On vérifie alors : $(Id - p)(p(\vec{u})) = \vec{0}$ et $p(\vec{u} - p(\vec{u})) = \vec{0}$, ce qui permet d’avoir d’ores et déjà $E = Ker(Id - p) + Ker(p)$.

Pour montrer que la somme est directe, on montre que l’intersection des deux noyaux est réduit à $\vec{0}$: $(Id - p)(\vec{v}) = \vec{0}$ et $p(\vec{v}) = \vec{0}$ implique $\vec{v} = \vec{0}$. On peut aussi montrer que si $\vec{0}$ se décompose sous la forme $\vec{a} + \vec{b}$ avec $\vec{a}$ dans $Ker(Id - p)$ et $\vec{b}$ dans $Ker(p)$, il n’y a que ces deux vecteurs possibles.

Réciproquement, on suppose $E = Ker(Id - p) \oplus Ker(p)$. Pour tout vecteur $\vec{a}$ de $Ker(Id - p)$, on a $(p - p^2)(\vec{a}) = \vec{0}$. Pour tout vecteur $\vec{b}$ de $Ker(p)$, on a $(p - p^2)(\vec{b}) = \vec{0}$. Par linéarité, le résultat s’étend : pour tout vecteur $\vec{u}$ de $Ker(Id - p) + Ker(p)$, on a $(p - p^2)(\vec{u}) = \vec{0}$. On reconnaît $p - p^2 = 0$.

On montre alors même : $Im(p) = Ker(Id - p)$ (et en permutant les rôles : $Ker(p) = Im(Id - p)$).

Tout vecteur $\vec{u}$ de $Im(p)$ s’écrit $p(\vec{a})$ pour au moins un vecteur $\vec{a}$ et vérifie $(Id - p)(\vec{u}) = (Id - p) \circ p(\vec{a}) = \vec{0}$.

Tout vecteur $\vec{u}$ de $Ker(Id - p)$ vérifie $\vec{u} - p(\vec{u}) = \vec{0}$ et s’écrit donc $\vec{u} = p(\vec{a})$ pour $\vec{a} = \vec{u}$ déjà.

On fusionne :

$Ker(p)$	$= Im(Id - p)$	$\vec{u} - p(\vec{u})$
$Im(p)$	$= Ker(Id - p)$	$p(\vec{u})$

Attention, le critère $E = Im(p) \oplus Ker(p)$ n’est pas une condition suffisante pour que p soit un projecteur. Par exemple les automorphismes vérifient ce critère sans pour autant être des projecteurs.

Réciproquement, si E s’écrit comme somme de deux supplémentaires A et B ($E = A \oplus B$, tout vecteur $\vec{u}$ s’écrit d’une façon unique $\vec{a} + \vec{b}$ avec $\vec{a}$ dans A et $\vec{b}$ dans B), alors on peut définir deux applications de E dans E : $p_A = \vec{u} \mapsto \vec{a}$ et $p_B = \vec{u} \mapsto \vec{b}$. Elles sont alors linéaires, et sont même des projecteurs.

projecteur	image	noyau
p_A	A	B
p_B	B	A

La connaissance du seul noyau ou de la seule image ne détermine pas totalement le projecteur. Il faut les deux. Il existe plusieurs projecteurs qui projettent sur A . La locution précise “parallèlement à B ” qui se traduit par “de noyau B ”.

On travaille en dimension finie. On se donne un projecteur p d’image A de noyau B .

Comme A et B sont supplémentaires, en concaténant une base $(\vec{e}_1, \dots, \vec{e}_r)$ de A et une base $(\vec{e}_{r+1}, \dots, \vec{e}_n)$

de B , on a une base de $(E, +, \cdot)$. Le projecteur est alors $\sum_{k=1}^n a_k \cdot \vec{e}_k \mapsto \sum_{k=1}^r a_k \cdot \vec{e}_k$.

Sa matrice sur cette base est de la forme $\begin{pmatrix} 1 & 0 & 0 & \dots & 0 & 0 \\ 0 & 1 & 0 & \dots & 0 & 0 \\ 0 & 0 & 1 & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 \end{pmatrix}$. Le nombre de 1 sur la diagonale

correspond à la dimension r de A (qui n'est autre que $Im(p)$).

La trace de cette matrice est égale au nombre de 1.

Comme toutes les matrices représentant p ont la même trace, on déduit que $Tr(p)$ est un entier naturel, égal à sa dimension : $Tr(p) = rg(p)$.

Attention, ce n'est que sur une base adaptée que la matrice est de cette forme, sinon, elles s'écrivent $P \cdot J_r \cdot P^{-1}$.

○ Montrer que l'intersection de deux sous-groupes d'un groupe $(G, *)$ est encore un sous-groupe de $(G, *)$. Définition du p.g.c.d. et du p.p.c.m. de deux entiers naturels non nuls a et b par $a\mathbb{Z} \cup b\mathbb{Z}$ et $a\mathbb{Z} \cap b\mathbb{Z}$.

Soient E et F deux sous-groupes de $(G, *)$ (parties de G contenant le neutre, stables par la loi $*$ et par passage au symétrique).

$E \cap F$ est alors encore une partie de G , contenant le neutre.

On prend a et b dans $E \cap F$. Ils sont donc dans E , leur "produit" $a * b$ est encore dans E . Ils sont tous deux dans F , $a * b$ est dans F . Étant dans E et dans F , $a * b$ est dans $E \cap F$.

De plus, comme a est dans E , son symétrique a^{-1} est dans E . Comme il est dans F , son symétrique est aussi dans F . On trouve bien que a^{-1} est dans $E \cap F$.

On a un résultat plus général : "si chaque E_i est un sous-groupe de $(G, *)$ pour tout i de I , alors $\bigcap_{i \in I} E_i$ est encore un sous-groupe de $(G, *)$ ".

Par définition, cette intersection généralisée $\{x \in G \mid \forall i \in I, x \in E_i\}$ est une partie de G contenant encore le neutre e .

On prend a et b dans cette intersection généralisée.

Par définition de cette appartenance : $\forall i \in I, a \in E_i, b \in E_i$. Par définition de "chaque E_i est un sous-groupe de $(G, *)$ ", on a $\forall i \in I, a * b^{-1} \in E_i$. On reconnaît que $a * b^{-1}$ est dans $\bigcap_{i \in I} E_i$.

*Un exercice classique montre que si A et B sont deux sous-groupes de $(G, *)$, alors $A \cup B$ n'est plus un sous-groupe de $(G, *)$ (sauf si A est inclus dans B ou B dans A). Il suffit de prendre un a dans A mais pas dans B et un b dans B mais pas dans A , puis de se demander si $a * b$ peut être dans $A \cup B$.*

On prend deux entiers naturels (non nuls pour ne pas compliquer) a et b . Les deux ensembles $a\mathbb{Z}$ et $b\mathbb{Z}$ sont des sous-groupes de $(\mathbb{Z}, +)$ (sens facile du théorème sur les sous-groupes de $(\mathbb{Z}, +)$).

Leur intersection $a\mathbb{Z} \cap b\mathbb{Z}$ est à son tour un sous-groupe de $(\mathbb{Z}, +)$, formé par définition des multiples communs de a et b . Ce sous-groupe est de la forme $m\mathbb{Z}$ (sens moins facile du théorème sur les sous-groupes de $(\mathbb{Z}, +)$), où m est le plus petit élément de $(a\mathbb{Z} \cap b\mathbb{Z}) \cap \mathbb{N}^*$. m est donc le plus petit multiple commun de a et b , qu'on note $p.p.c.m.(a, b)$ ou même $a \vee b$. On a donc $a\mathbb{Z} \cap b\mathbb{Z} = (a \vee b)\mathbb{Z}$.

Cette définition permet de montrer que l'opérateur p.p.c.m. est commutatif, associatif ; on trouve aussi le neutre 1 et l'élément absorbant 0.

On considère ensuite $a\mathbb{Z} \cup b\mathbb{Z}$ qui n'est en général pas un sous-groupe de $(\mathbb{Z}, +)$. On cherche plus petit sous-groupe de $(\mathbb{Z}, +)$ contenant $a\mathbb{Z} \cup b\mathbb{Z}$. Il est de la forme $d\mathbb{Z}$ pour un entier naturel d .

Comme on a $a \in a\mathbb{Z} \subset (a\mathbb{Z} \cup b\mathbb{Z}) \subset d\mathbb{Z}$, d est un diviseur de a .

De même, d est un diviseur de b .

Comme on prend le plus petit sous-groupe, d est le plus grand diviseur commun de a et b .

On a donc la formule $Eng(a\mathbb{Z} \cup b\mathbb{Z}) = (a \wedge b)\mathbb{Z}$ (la notation $Eng(H)$ n'étant pas dans le programme pour désigner le sous-groupe engendré par une partie).

C'est avec ces deux définitions qu'on montre que cette fois le p.g.c.d. est commutatif, associatif et même distributif sur le p.p.c.m. (et vice versa, ce qui fait de $(\mathbb{Z}, \vee, \wedge)$ une structure assez particulière).

Mieux encore, on montre que si une partie de $\mathbb{Z}$ stable par addition et soustraction contient les multiples de a et les multiples de b (c'est à dire contient $a.\mathbb{Z} \cup b.\mathbb{Z}$), alors elle contient tous les entiers de la forme $\alpha.a + \beta.b$ avec a et b dans $\mathbb{Z}$. On montre ensuite que $\{\alpha.a + \beta.b \mid (\alpha, \beta) \in \mathbb{Z}^2\}$ est un sous-groupe de $(\mathbb{Z}, +)$.

On déduit : $\{\alpha.a + \beta.b \mid (\alpha, \beta) \in \mathbb{Z}^2\} = (a \wedge b).\mathbb{Z}$.

Le p.g.c.d. de a et b s'écrit donc sous la forme $\alpha.a + \beta.b$ et c'est même le plus petit entier naturel s'écrivant sous cette forme. C'est l'identité de Bézout.

○ Démontrer que toute suite réelle croissante majorée converge vers son plus petit majorant.

On prend une suite réelle a , croissante, majorée par M .

Le théorème admis en Terminale disait alors "elle converge".

Le théorème de Sup (basé sur le principe de la borne supérieure) dit alors "elle converge vers son plus petit majorant".

C'est ce qui permet d'écrire alors des choses comme $\forall k \in \mathbb{N}, a_k \leq \lim_{n \rightarrow +\infty} a_n \leq M$.

On pose $A = \{a_n \mid n \in \mathbb{N}\}$ (ensemble des valeurs prises par la suite, projection sur Oy).

C'est une partie de $\mathbb{R}$, non vide.

Elle est majorée (par M entre autres, qui est un majorant).

Cette partie admet alors un plus petit majorant, qu'on note α (théorème ou axiome suivant qu'on a construit ou non $\mathbb{R}$).

Par définition de "majorant", on a $a_n \leq \alpha$ pour tout n . On a même $a_n \leq \alpha \leq \alpha + \varepsilon$ pour tout n et tout ε strictement positif.

Prenons d'ailleurs un ε strictement positif. Par définition de "plus petit majorant", $\alpha - \varepsilon$ n'est plus un majorant de A . Ceci se traduit par l'existence d'au moins un élément a_N de A vérifiant $\alpha - \varepsilon \leq a_N$.

En mettant bout à bout ces inégalités, incluant la croissance de la suite a , on a :

$\forall \varepsilon > 0, \exists N \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq N \Rightarrow \alpha - \varepsilon \leq a_n \leq \alpha \leq \alpha + \varepsilon$.

En effaçant ce qu'il faut, on reconnaît la définition de $a_n \rightarrow_{n \rightarrow +\infty} \alpha$.

Une suite réelle croissante n'a que deux possibilités :

- converger vers son plus petit majorant
- diverger vers $+\infty$.

○ Donner la définition de "morphisme d'espaces vectoriels", "endomorphisme", "isomorphisme", "automorphisme". Stabilités diverses. Groupe linéaire.

Un morphisme f entre deux espaces vectoriels $(E, +, \cdot)$ et $(F, +, \cdot)$ (sur le même corps $(\mathbb{K}, +, \times)$) est une application de E dans F "transparente" pour les deux lois d'espace vectoriel :

$\forall (\vec{a}, \vec{b}) \in E^2, f(\vec{a} + \vec{b}) = f(\vec{a}) + f(\vec{b})$ et $\forall \vec{a} \in E, \forall \alpha \in \mathbb{K}, f(\alpha.\vec{a}) = \alpha.f(\vec{a})$
ou même en une fois $\forall (\vec{a}, \vec{b}) \in E^2, \forall (\alpha, \beta) \in \mathbb{K}^2, f(\alpha.\vec{a} + \beta.\vec{b}) = \alpha.f(\vec{a}) + \beta.f(\vec{b})$

Plutôt que "morphisme d'espaces vectoriels", on dira en général "application linéaire".

Pour toute application linéaire f de $(E, +, \cdot)$ dans $(F, +, \cdot)$ on a $f(\vec{0}_E) = \vec{0}_F$.

La seule application linéaire constante est $\vec{a} \mapsto \vec{0}_F$ (qui est application linéaire neutre pour l'addition).

Si f et g sont des applications linéaires de E dans F , alors $f + g$ est aussi une application linéaire.

De même, si μ est un coefficient dans $\mathbb{K}$ alors $\mu.f$ est encore une application linéaire.

On note $L(E, F)$ (et même plus proprement $L_{\mathbb{K}}(E, F)$) l'espace des applications linéaires de $(E, +, \cdot)$ dans $(F, +, \cdot)$. Il s'ensuit que $(L(E, F), +, \cdot)$ est à son tour un espace vectoriel.

Si $(E, +, \cdot)$ est de dimension finie n (muni d'une base $(\vec{e}_1, \dots, \vec{e}_n)$) et $(F, +, \cdot)$ de dimension p (muni d'une base $(\vec{a}_1, \dots, \vec{a}_p)$), alors une application linéaire f de E dans F est totalement déterminée par la donnée des p composantes sur la base $(\vec{a}_1, \dots, \vec{a}_p)$ de chacune des n images $f(\vec{e}_k)$.

C'est ce qui permet d'identifier alors $L(E, F)$ à $M_{n,p}(\mathbb{K})$ (ensemble des matrices à n colonnes et p lignes).

Il s'ensuit : $\dim(L(E, F)) = \dim(E).\dim(F)$.

	nom	définition
On peut préfixer le mot “morphisme”	endomorphisme	d’un espace dans lui même
	isomorphisme	bijectif
	automorphisme	“endo” et “iso”

L’image d’une famille libre par une application linéaire injective est libre.

L’image d’une famille génératrice de $(E, +, \cdot)$ par une application linéaire surjective de $(E, +, \cdot)$ dans $(F, +, \cdot)$ est une famille génératrice de $(F, +, \cdot)$.

Un isomorphisme de $(E, +, \cdot)$ dans $(F, +, \cdot)$ transforme une base de $(E, +, \cdot)$ en base de $(F, +, \cdot)$.

On déduit (*en dimension finie*) qu’il existe un isomorphisme entre E et F si et seulement si $(E, +, \cdot)$ et $(F, +, \cdot)$ ont la même dimension.

Enfin, la composée de deux morphismes est encore un morphisme.

Ceci permet de dire que les endomorphismes de $(E, +, \cdot)$ forment une algèbre (il faut vérifier $(f + g) \circ \phi = f \circ \phi + g \circ \phi$ et $f \circ (\phi + \psi) = f \circ \phi + f \circ \psi$, que l’on vérifie vecteur par vecteur).

Dès qu’on ajoute un critère de bijectivité, on perd l’application nulle, et on perd aussi toute stabilité par addition. Les isomorphismes sont bien plus intéressants à composer et inverser.

Le groupe linéaire est donc un groupe pour la loi de composition et surtout pas pour l’addition.

On peut vérifier que les automorphismes de $(E, +, \cdot)$ forment un groupe pour la loi de composition, non commutatif dès que la dimension atteint et dépasse 2.

applications	ensemble	lois		structure	
morphismes	$L(E, F)$	$+, \cdot$		espace vectoriel	matrices rectangulaires
endomorphismes	$L(E)$	$+, \cdot$	$\circ$	algèbre	matrices carrées
automorphismes	$GL(E)$		$\circ$	groupe	matrices carrées de déterminant non nul

○ Montrer que si f est intégrable et bornée sur $[a, b]$, alors $x \mapsto \int_a^x f(t).dt$ (notée F ou même F_a) est lipschitzienne (donc continues). Passage de f positive à F croissante. Passer de f continue à F dérivable.

L’intégrabilité de f sur $[a, b]$ (par exemple issue de “en escalier” ou “continue” par morceaux), quelle passe par les sommes de Riemann, Darboux ou toute autre définition, se transmet à tout intervalle $[x, y]$ inclus dans $[a, b]$ et en particulier à tout intervalle $[a, x]$ (c’est la première partie de la démonstration de Chasles).

Chaque intégrale $\int_a^x f(t).dt$ existe bien.

Cette application F s’appelle intégrale fonction de la borne supérieure (la borne supérieure c’est x , en tant que bornes de l’intervalle d’intégration).

Quand on calcule $F(y) - F(x)$, la relation de Chasles (encore elle) donne : $F(y) - F(x) = \int_x^y f(t).dt$.

On a supposé f bornée en valeur absolue par M sur $[a, b]$ (donc sur $[x, y]$). Alors, géométriquement ou même par définition des sommes de Riemann et Darboux, on a $\left| \int_x^y f(t).dt \right| \leq M \cdot |y - x|$ (hauteur maximale fois largeur de l’intervalle).

On reconnaît que F est lipschitzienne de rapport M .

Ceci entraîne effectivement que F continue, même si f ne l’était pas.

Si la fonction sous le signe somme est positive, l’intégrale fonction de la borne supérieure est croissante. Ce serait un gaspillage énorme de le prouver en dérivant F en f et de passer de la positivité de F' à la croissance de F . En effet, une telle preuve fait appel à deux gros théorèmes d’analyse, inutilement.

Il suffit, pour f positive, de prendre x et y dans $[a, b]$, avec x plus petit que y , et d’écrire $F(y) - F(x) = \int_x^y f(t).dt$. La fonction est positive, l’intervalle est dans le bon sens ; la différence est positive, du signe de $y - x$.

On passe ensuite au théorème fondamental : si f est positive, alors F est dérivable de dérivée f .

On se donne x et h avec x et $x+h$ dans $[a, b]$ et on calcule la différence $F(x+h) - F(x) = \int_x^{x+h} f(t).dt$.

L’application f est continue sur $[x, x+h]$. Par continuité, elle est bornée et atteint ses bornes sur ce

segment en μ_h et ν_h . On a donc $\forall t \in [x, x+h], f(\mu_h) \leq f(t) \leq f(\nu_h)$. On intègre de x à $x+h$, on encadre donc $\int_x^{x+h} f(t).dt$ par $h.f(\mu_h)$ et $h.f(\nu_h)$. On divise par h : $f(\mu_h) \leq \frac{F(x+h) - F(x)}{h} \leq f(\nu_h)$

(si h est négatif, le premier encadrement est $h.f(\mu_h) \geq \int_x^{x+h} f(t).dt \geq h.f(\nu_h)$ à cause du sens de l'intervalle, mais la division par h remet l'inégalité dans le sens cité).

Le taux d'accroissement $\frac{F(x+h) - F(x)}{h}$ est compris entre deux valeurs atteintes par f sur l'intervalle.

Par théorème des valeurs intermédiaires appliqué à f (continue), cette valeur est réalisée en un point au moins de $[\mu_h, \nu_h]$ donc en un point de $[x, x+h]$, qu'on pourra noter τ_h .

On vient de montrer que la valeur moyenne $\frac{1}{h} \int_x^{x+h} f(t).dt$ est réalisée en au moins un point de l'intervalle pour une application continue.

On fait tendre h vers 0 ; par encadrement τ_h tend vers x et par continuité de f , $f(\tau_h)$ converge vers $f(x)$.

On vient de prouver en utilisant plusieurs théorèmes de continuité que $\frac{F(x+h) - F(x)}{h}$ converge vers $f(x)$ quand h tend vers 0.

C'est la définition même de la dérivabilité de F , de dérivée f .

○ Déterminer le *p.g.c.d.* de deux entiers naturels par algorithme d'Euclide. Remonter cet algorithme pour obtenir une identité de Bézout.

On va commencer par des lemmes.

Si a et b sont deux entiers naturels, alors on a $p.g.c.d.(a, b) = p.g.c.d.(a - b, b)$.

En effet, si un entier divise a et b , alors il divise $a - b$ et b .

De même, si un entier divise $a - b$ et b , par somme, il divise a et b .

Les couples (a, b) et $(a - b, b)$ ont les mêmes diviseurs communs.

Ils ont donc le même plus grand diviseur commun.

En mettant en boucle le lemme précédent, on a $p.g.c.d.(a, b) = p.g.c.d.(a - q.b, b)$ pour tout naturel q .

En effectuant la division euclidienne de a par b ($a = b.q + r$), on a $p.g.c.d.(a, b) = p.g.c.d.(b, r)$.

On peut donc étape par étape remplacer le couple (a, b) (avec rôles dividende, diviseur) par le couple (b, r) (diviseur, reste), jusqu'à aboutir à un couple où le reste est nul (la suite des restes est strictement décroissante, elle finit par arriver sur 0). Pour ce dernier couple $(r_k, 0)$, le *p.g.c.d.* vaut r .

C'est donc que le *p.g.c.d.* est "le dernier reste non nul de l'algorithme d'Euclide".

dividende	diviseur	reste	division		dividende	diviseur	reste	division
$a = a_0$	$b = a_1$	a_2	$a_0 = a_1.q_0 + a_2$					
a_1	a_2	a_3	$a_1 = a_2.q_1 + a_3$		2016	34	10	$2016 = 34.59 + 10$
a_2	a_3	a_4	$a_1 = a_2.q_2 + a_3$		34	10	4	$34 = 10.3 + 4$
$\vdots$	$\vdots$	$\vdots$	$\vdots$		10	4	2	$10 = 4.2 + 2$
a_{k-2}	a_{k-1}	a_k	$a_{k-2} = a_{k-1}.q_{k-2} + a_k$		4	2	0	$4 = 2.2 + 0$
a_{k-1}	a_k	0	$a_{k-1} = a_k.q_{k-1}$					
	<i>p.g.c.d.</i>							

On peut écrire l'algorithme sous forme matricielle $\begin{pmatrix} a \\ b \end{pmatrix} = \begin{pmatrix} q_0 & 1 \\ 1 & 0 \end{pmatrix} \cdot \begin{pmatrix} b \\ r \end{pmatrix}$ et plus généralement

$\begin{pmatrix} a_k \\ b_{k+1} \end{pmatrix} = \begin{pmatrix} q_k & 1 \\ 1 & 0 \end{pmatrix} \cdot \begin{pmatrix} b_k \\ r_k \end{pmatrix}$ avec des matrices dont on vérifie qu'elles sont inversibles.

On peut mettre en boucle : $\begin{pmatrix} a \\ b \end{pmatrix} = \begin{pmatrix} q_0 & 1 \\ 1 & 0 \end{pmatrix} \cdot \begin{pmatrix} q_1 & 1 \\ 1 & 0 \end{pmatrix} \cdots \begin{pmatrix} q_k & 1 \\ 1 & 0 \end{pmatrix} \cdot \begin{pmatrix} r_k \\ 0 \end{pmatrix}$ où r_k est le dernier reste non nul, c'est à dire le *p.g.c.d.*.

On inverse pas à pas : $\begin{pmatrix} r_k \\ 0 \end{pmatrix} = \begin{pmatrix} q_k & 1 \\ 1 & 0 \end{pmatrix}^{-1} \cdots \begin{pmatrix} q_0 & 1 \\ 1 & 0 \end{pmatrix}^{-1} \cdot \begin{pmatrix} a \\ b \end{pmatrix}$

et on explicite : $\begin{pmatrix} p.g.c.d. \\ 0 \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ 1 & -q_k \end{pmatrix} \cdots \begin{pmatrix} 0 & 1 \\ 1 & -q_0 \end{pmatrix} \cdot \begin{pmatrix} a \\ b \end{pmatrix}$

de la forme $\begin{pmatrix} p.g.c.d. \\ 0 \end{pmatrix} = \begin{pmatrix} \alpha_k & \beta_k \\ \gamma_k & \delta_k \end{pmatrix} \cdot \begin{pmatrix} a \\ b \end{pmatrix}$. La première ligne du produit matriciel donne une identité de Bézout calculant le $p.g.c.d.$ de a et b comme combinaison de a et b .

Même sans utiliser le produit matriciel, on peut remonter ligne à ligne

dividende	diviseur	reste	division	renversement	Bézout
2016	34	10	$2016 = 34.59 + 10$	$10 = 2016.(1) - 34.(59)$	$2 = 34.(-2) + [2016.(1) - 34.(59)].(7)$
34	10	4	$34 = 10.3 + 4$	$4 = 34.(1) - 10.(3)$	$2 = 10.(1) - [34 - 10.3].(2) = 34.(-2) + 10.(7)$
10	4	2	$10 = 4.2 + 2$	$2 = 10.(1) - 4.(2)$	
4	2	0	$4 = 2.2 + 0$		

La case en haut à gauche du tableau donne l'identité de Bézout $2 = 34.(-415) + 2016.(7)$.

○ Écrire la relation de Cayley-Hamilton pour une matrice carrée de taille 2 sur 2.

On vérifie aisément pour $M = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$, en posant $Tr(M) = (a + d)$ et $\det(M) = a.d - b.c$:

$$M^2 - Tr(M).M + \det(M).I_2 = 0_2.$$

Ceci permet d'inverser la matrice dès que son déterminant est non nul :

$$\det(M).I_2 = M.(Tr(M).I_2 - M) \text{ puis } I_2 = \frac{Tr(M).I_2 - M}{\det(M)}.M = M.\frac{Tr(M).I_2 - M}{\det(M)}.$$

Ceci permet d'identifier : $M^{-1} = \frac{Tr(M).I_2 - M}{\det(M)}.$

Il est quand même plus facile de se souvenir : $\begin{pmatrix} a & b \\ c & d \end{pmatrix}^{-1} = \frac{1}{a.d - b.c} \cdot \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}$ directement.

Il existe une formule aussi en taille 3 : $M = \begin{pmatrix} a & a' & a'' \\ b & b' & b'' \\ c & c' & c'' \end{pmatrix}$

$$M^3 - Tr(M).M^2 + Min_2(M).M - \det(M).I_3 = 0_3$$

avec $Tr(M) = a + b' + c''$,

$$\det(M) = a.b'.c'' + b.c'.a'' + a'.b''.c - c.b'.a'' - b.a'.c'' - c'.b''.a \text{ et}$$

$$Min_2(M) = \begin{vmatrix} b' & b'' \\ c' & c'' \end{vmatrix} + \begin{vmatrix} a & a'' \\ c & c'' \end{vmatrix} + \begin{vmatrix} a & a' \\ b & b' \end{vmatrix}.$$

La trace et le déterminant sont des termes classiques du programme.

Le terme $Min_2(M)$ est d'un usage plus limité. On peut aussi l'écrire $Tr(Com(M))$ ou $\frac{Tr(M)^2 - Tr(M^2)}{2}$.

C'est aussi par cette formule qu'on peut inverser M en isolant

$$I_3 = \frac{1}{\det(M)}(M^2 - Tr(M).M + Min_2(M).I_3).M.$$

La formule se généralise en grande dimension :

si on note P_M le polynôme caractéristique de la matrice M (d'écriture formelle $\det(M - X.I_n)$), alors on a $P_M(M) = 0_n$.

La démonstration est hors-programme en Sup.

○ Démontrer, en admettant le théorème de Bolzano-Weierstrass, que toute application numérique continue sur un segment est bornée et atteint ses bornes.

On admet donc que de toute suite réelle bornée on peut extraire au moins une sous-suite convergente.

On prend alors une application f de $[a, b]$ dans $\mathbb{R}$ continue en tout point.

On montre par l'absurde qu'elle est majorée.

On quantifie donc " f n'est pas majorée" et on veut arriver à une contradiction :

$$\forall M \in \mathbb{R}, \exists t \in [a, b], f(t) > M \text{ (négation de } \exists M, \forall t \in [a, b], f(t) \leq M)$$

En particulier pour chaque entier naturel n , il existe au moins un réel t de $[a, b]$ qui vérifie $f(t) > n$.

On note t_n un tel réel (car il faut en changer quand on change de n).

On a alors une suite (t_n) de réels de $[a, b]$. C'est donc une suite réelle bornée. On peut en extraire au moins une sous-suite $(t_{\varphi(n)})$ qui converge vers un réel α .

Sachant "pour tout n , $a \leq t_{\varphi(n)} \leq b$ ", on a, par passage à la limite (large) : $a \leq \alpha \leq b$.

Comme α est un point de $[a, b]$, f est continue en α .

Par critère séquentiel, comme la suite $(t_{\varphi(n)})$ converge vers α , la suite $(f(t_{\varphi(n)}))$ converge vers $f(\alpha)$.

Mais dans le même temps, on a $(f(t_{\varphi(n)})) \geq \varphi(n) \geq n$ (définition des t_k et lemme d'extraction). Par minoration, la suite $(f(t_{\varphi(n)}))$ diverge vers $+\infty$.

Cette suite étant à la fois convergente et divergente, on a notre contradiction.
 f est donc majorée.

En appliquant ce même raisonnement à $-f$, f est minorée. Il s'ensuit que f est bornée.
On peut aussi dire que $|f|$ est continue sur $[a, b]$ donc majorée, ce qui suffit pour borner f .

L'ensemble V des valeurs prises par f est donc une partie de $\mathbb{R}$ majorée.

Cet ensemble V contient au moins $f(a)$. C'est donc une partie de $\mathbb{R}$ non vide majorée.

On pose alors $S = \text{Sup}\{f(t) \mid t \in [a, b]\}$. Il faut à présent trouver un élément β de $[a, b]$ vérifiant $f(\beta) = S$.

Par définition de la borne supérieure, pour tout n , $S - 2^{-n}$ n'est plus un majorant de V . Il existe donc au moins un élément de V (noté $f(u_n)$) vérifiant $S - 2^{-n} \leq f(u_n)$.

Les abscisses u_n forment une suite de points de $[a, b]$. On peut en extraire au moins une sous-suite $(u_{\psi(n)})$ qui converge. On note β sa limite. Par passage à la limite sur $\forall n, a \leq u_{\psi(n)} \leq b$, le réel β est dans $[a, b]$. Par continuité de f sur $[a, b]$, f est continue en β et on a donc $f(u_{\psi(n)}) \rightarrow_{n \rightarrow +\infty} f(\beta)$.

Mais on a dans le même temps pour tout n : $S - 2^{-\psi(n)} \leq f(u_{\psi(n)}) \leq S$. Par encadrement : $f(u_{\psi(n)}) \rightarrow_{n \rightarrow +\infty} S$.

Par unicité de la limite en cas de convergence : $f(\beta) = S$. La borne supérieure est un maximum, atteint au moins en β .

On fait de même avec le minimum.

Ce résultat permet de montrer que $f \mapsto \text{Sup}\{|f(t)| \mid t \in [a, b]\}$ est une norme sur l'espace des applications continues de $[a, b]$ dans $\mathbb{R}$. Le résultat précédent donne l'existence (et la positivité).

La démonstration de l'inégalité triangulaire est à maîtriser.

On se donne f et g continues de $[a, b]$ dans $\mathbb{R}$. L'application $f + g$ est à son tour continue et bornée.

D'ailleurs, pour tout x de $[a, b]$, on a par inégalité triangulaire et définition d'un majorant :

$$|f(x) + g(x)| \leq |f(x)| + |g(x)| \leq \text{Sup}\{|f(t)| \mid t \in [a, b]\} + \text{Sup}\{|g(t)| \mid t \in [a, b]\}$$

Le réel $\text{Sup}\{|f(t)| \mid t \in [a, b]\} + \text{Sup}\{|g(t)| \mid t \in [a, b]\}$ est un majorant de $|f + g|$ sur $[a, b]$.

Comme le réel $\text{Sup}\{|f(x)| \mid x \in [a, b]\}$ est par définition "le plus petit majorant", on a bien

$$\text{Sup}\{|f(x) + g(x)| \mid x \in [a, b]\} \leq \text{Sup}\{|f(t)| \mid t \in [a, b]\} + \text{Sup}\{|g(t)| \mid t \in [a, b]\}$$

qu'on écrit aussi $\|f + g\| \leq \|f\| + \|g\|$ ou aussi $N_\infty(f + g) \leq N_\infty(f) + N_\infty(g)$.

Pour l'homogénéité, on se donne f et λ non nul (pour λ nul, c'est sans intérêt).

Pour tout x de $[a, b]$, on a $|(\lambda.f)(x)| = |\lambda| \cdot |f(x)| \leq |\lambda| \cdot \text{Sup}\{|f(t)| \mid t \in [a, b]\} = |\lambda| \cdot N_\infty(f)$.

Le réel positif $|\lambda| \cdot N_\infty(f)$ est un majorant de $|\lambda.f|$ sur $[a, b]$.

Toujours par définition du plus petit majorant, on a $N_\infty(\lambda.f) \leq |\lambda| \cdot N_\infty(f)$.

En appliquant ce même résultat au couple $(\lambda.f, \frac{1}{\lambda})$, on a $N_\infty(\frac{1}{\lambda} \cdot \lambda.f) \leq |\frac{1}{\lambda}| \cdot N_\infty(\lambda.f)$, et on trouve l'autre moitié de l'inégalité.

○ Exprimer les sommes de Riemann droite, gauche et milieu pour une équisubdivision en n parties pour une application f sur un segment $[a, b]$.

On découpe le segment $[a, b]$ en n segments, par $n + 1$ points $a + k \cdot \frac{b-a}{n}$ avec k de 0 à n .

Ces points peuvent aussi être écrits $\frac{(n-k) \cdot a + k \cdot b}{n}$ si on veut les voir comme barycentres plutôt que comme points qui avancent pas à pas.

Pour la somme des aires de rectangles, la largeur est à chaque fois $\frac{b-a}{n}$, et la hauteur se mesure par f .

On a trois sommes de Riemann :

gauche	milieu	droite
$\frac{b-a}{n} \cdot \sum_{k=0}^{n-1} f\left(a + k \cdot \frac{b-a}{n}\right)$	$\frac{b-a}{n} \cdot \sum_{k=0}^{n-1} f\left(a + \frac{(2k+1) \cdot (b-a)}{2n}\right)$	$\frac{b-a}{n} \cdot \sum_{k=1}^n f\left(a + k \cdot \frac{b-a}{n}\right)$
for k in range(n)		$\frac{b-a}{n} \cdot \sum_{k=0}^{n-1} f\left(a + (k+1) \cdot \frac{b-a}{n}\right)$

Géométriquement, en rappelant que l'aire d'un trapèze se calcule par base $\times$ (moyenne des hauteurs), la

moyenne $\frac{R_g(f) + R_d(f)}{2}$ correspond à la méthode des trapèzes.

La formule des trois niveaux dit aussi que la méthode des “arcs de parabole” (dite de Simpson) fait appel à la moyenne $\frac{R_g(f) + 4.R_n(f) + R_d(f)}{6}$.

○ **Cofacteur dans une matrice carrée. Formule du développement par rapport à une ligne ou une colonne. Formule $A \cdot {}^t Com(A) = {}^t Com(A) \cdot A = \det(A) \cdot I_n$. Calcule théorique de l'inverse d'une matrice carrée de déterminant non nul. Formules de Cramer pour la résolution d'un système de n équations à n inconnues.**

On se donne une matrice A carrée de taille n , de terme général a_i^k (ligne i colonne k).

On se donne deux indices p et q , on définit le cofacteur α_p^q , déterminant de la matrice de taille $n - 1$ en effaçant la ligne p et la colonne q .

Le cofacteur pondéré est $(-1)^{p+q} \cdot \alpha_p^q$.

La comatrice est la matrice des cofacteurs pondérés.

Par exemple, en dimensions 2 et 3 :

$A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$	$Com(A) = \begin{pmatrix} d & -c \\ -b & a \end{pmatrix}$
$A = \begin{pmatrix} a & a' & a'' \\ b & b' & b'' \\ c & c' & c'' \end{pmatrix}$	$Com(A) = \begin{pmatrix} \begin{vmatrix} b' & b'' \\ c' & c'' \end{vmatrix} & -\begin{vmatrix} b & b'' \\ c & c'' \end{vmatrix} & \begin{vmatrix} b & b' \\ c & c' \end{vmatrix} \\ -\begin{vmatrix} a' & a'' \\ c' & c'' \end{vmatrix} & \begin{vmatrix} a & a'' \\ c & c'' \end{vmatrix} & -\begin{vmatrix} a & a' \\ c & c' \end{vmatrix} \\ \begin{vmatrix} a' & a'' \\ b' & b'' \end{vmatrix} & -\begin{vmatrix} a & a'' \\ b & b'' \end{vmatrix} & \begin{vmatrix} a & a' \\ b & b' \end{vmatrix} \end{pmatrix}$

On démontre de deux façons différentes la formule de développement par rapport à la première ligne :

$$\det(A) = \sum_{k=1}^n (-1)^k \cdot a_1^k \cdot \alpha_1^k$$

On fait appel à la définition “brutale” du déterminant : $\det(A) = \sum_{\sigma \in S_n} Sgn(\sigma) \cdot a_1^{\sigma(1)} \cdot a_2^{\sigma(2)} \dots a_n^{\sigma(n)}$.

On sépare en fonction de la valeur de $\sigma(1)$: $\det(A) = \sum_{k=1}^n \left(\sum_{\substack{\sigma \in S_n \\ \sigma(1)=k}} Sgn(\sigma) \cdot a_1^{\sigma(1)} \cdot a_2^{\sigma(2)} \dots a_n^{\sigma(n)} \right)$.

On factorise : $\det(A) = \sum_{k=1}^n a_1^k \cdot \left(\sum_{\substack{\sigma \in S_n \\ \sigma(1)=k}} Sgn(\sigma) \cdot a_2^{\sigma(2)} \dots a_n^{\sigma(n)} \right)$.

Chaque somme $\sum_{\substack{\sigma \in S_n \\ \sigma(1)=k}} Sgn(\sigma) \cdot a_2^{\sigma(2)} \dots a_n^{\sigma(n)}$ peut être vue comme le déterminant d'une matrice de taille

$n - 1$ issue de la matrice A . Simplement, il n'y a plus aucun a_1^k ; la ligne 1 de la matrice A a disparu. De même, aucun indice de colonne ne vaut k puisque c'est a_1^k qui a pris l'indice k ; la colonne k de la matrice A a disparu. Il reste à se convaincre que la signature de σ et la signature de la “permutation” de $\{2, 3, \dots, n\}$ sur $1, \dots, k - 1, k + 1, \dots, n\}$ diffèrent d'un facteur $(-1)^k$.

On peut aussi comparer les deux formes n -linéaires $(\vec{u}_1, \dots, \vec{u}_n) \mapsto \det(\vec{u}_1, \dots, \vec{u}_n)$ et $(\vec{u}_1, \dots, \vec{u}_n) \mapsto$

$$\sum_{k=1}^n (-1)^k \cdot \mu(\vec{u}_k) \cdot \det(p(\vec{u}_1), \dots, p(\vec{u}_{k-1}), p(\vec{u}_{k+1}), \dots, p(\vec{u}_n)) \text{ avec } \mu = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \mapsto x_1 \text{ (forme linéaire sur } \mathbb{R}^n \text{)}$$

$$\mathbb{R}^n \text{ et } p = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \mapsto \begin{pmatrix} x_2 \\ \vdots \\ x_n \end{pmatrix} \text{ (de } \mathbb{R}^n \text{ dans } \mathbb{R}^{n-1}).$$

On montre que ces deux applications sont des formes n -linéaires antisymétriques. Elles sont donc proportionnelles. Leur valeur sur la base canonique montre qu'elles sont égales.

La seconde forme est la formule de développement par rapport à la première ligne.

Comme une matrice et sa transposée ont le même déterminant, la formule de développement par rap-

port à la première ligne devient la formule de développement par rapport à la première colonne :

$$(A) = \sum_{i=1}^n (-1)^i . a_i^1 . \alpha_i^1.$$

En échangeant les colonnes par antisymétrie, on démontre aussi la formule de développement par rap-

port à la colonne d'indice q : $(A) = \sum_{i=1}^n (-1)^{i+q} . a_i^q . \alpha_i^q$; et en transposant : $(A) = \sum_{k=1}^n (-1)^{k+p} . a_p^k . \alpha_p^k$.

On considère alors le produit de la matrice A et de la matrice ${}^tCom(A)$. Son terme général est $\sum_{j=1}^n a_i^j . \beta_j^k$

où β_j^k est le terme général de ${}^tCom(A)$. C'est donc $(-1)^{j+k} . \alpha_k^j$.

Le terme de ligne i et colonne k de $A . {}^tCom(A)$ est donc $\sum_{j=1}^n (-1)^{j+k} . a_i^j . \alpha_k^j$.

Il faut alors séparer deux cas :

- $i = k$: on reconnaît alors exactement le développement de $\det(A)$ par rapport à la i^{eme} ligne ;
- $i \neq k$: on reconnaît alors le développement du déterminant d'une matrice A' par rapport à sa k^{ieme} ligne, mais d'une matrice dont la k^{ieme} ligne est la ligne d'indice i de la matrice A ; ce déterminant est donc nul.

On a donc les coefficients de $A . {}^tCom(A)$ qui sont tous nuls, sauf ceux de la diagonale qui valent $\det(A)$.

On reconnaît $A . {}^tCom(A) = \det(A) . I_n$.

Pour saisir en dimension 3 :

$$\begin{pmatrix} a & a' & a'' \\ b & b' & b'' \\ c & c' & c'' \end{pmatrix} \cdot \begin{pmatrix} \begin{vmatrix} b' & b'' \\ c' & c'' \end{vmatrix} & - \begin{vmatrix} a' & a'' \\ c' & c'' \end{vmatrix} & \begin{vmatrix} a' & a'' \\ b' & b'' \end{vmatrix} \\ - \begin{vmatrix} b & b'' \\ c & c'' \end{vmatrix} & \begin{vmatrix} a & a'' \\ c & c'' \end{vmatrix} & - \begin{vmatrix} a & a'' \\ b & b'' \end{vmatrix} \\ \begin{vmatrix} b & b' \\ c & c' \end{vmatrix} & - \begin{vmatrix} a & a' \\ c & c' \end{vmatrix} & \begin{vmatrix} a & a' \\ b & b' \end{vmatrix} \end{pmatrix} = \begin{pmatrix} \begin{vmatrix} a & a' & a'' \\ b & b' & b'' \\ c & c' & c'' \end{vmatrix} & \begin{vmatrix} a & a' & a'' \\ c & c' & c'' \end{vmatrix} & \begin{vmatrix} a & a' & a'' \\ b & b' & b'' \end{vmatrix} \\ \begin{vmatrix} b & b' & b'' \\ c & c' & c'' \end{vmatrix} & \begin{vmatrix} a & a' & a'' \\ b & b' & b'' \end{vmatrix} & \begin{vmatrix} a & a' & a'' \\ c & c' & c'' \end{vmatrix} \\ \begin{vmatrix} c & c' & c'' \\ a & a' & a'' \\ b & b' & b'' \end{vmatrix} & \begin{vmatrix} c & c' & c'' \\ a & a' & a'' \\ c & c' & c'' \end{vmatrix} & \begin{vmatrix} c & c' & c'' \\ b & b' & b'' \\ c & c' & c'' \end{vmatrix} \end{pmatrix} = \begin{pmatrix} \det(A) & 0 & 0 \\ 0 & \det(A) & 0 \\ 0 & 0 & \det(A) \end{pmatrix}$$

Par des arguments similaires, on trouve ${}^tCom(A) . A = \det(A) . I_n$.

Si $\det(A)$ est non nul, on divise par $\det(A)$ et on reconnaît dans $A . \frac{{}^tCom(A)}{\det(A)} = \frac{{}^tCom(A)}{\det(A)} . A = I_n$ le

fait que $\frac{{}^tCom(A)}{\det(A)}$ est bien l'inverse de A .

On rappelle que pour $\det(A)$ nul, un raisonnement par l'absurde prouve que mA n'est pas inversible.

On pourra d'ailleurs retenir la formule "une matrice est inversible si et seulement si elle a un déterminant, et il est non nul".

Il devient alors facile d'inverser le système $A.X = Y$ d'inconnue X et de paramètre Y , avec A carrée inversible.

L'unique solution est alors $\frac{{}^tCom(A)}{\det(A)} . Y$.

Si on calcule le terme de ligne i de ce produit matriciel, on trouve, avec les notations déjà prises :

$$x_i = \sum_{j=1}^k (-1)^{i+j} . \beta_i^j . y_j = \sum_{j=1}^k (-1)^{i+j} . \alpha_j^i . y_j.$$

On reconnaît cette fois encore le développement par rapport à la i^{eme} colonne du déterminant d'une matrice qui ressemble à A (*présence des cofacteurs*) mais dont la i^{eme} colonne est simplement faite des coefficients y_j . C'est donc le déterminant d'une matrice où l'on a remplacé la colonne i (*coefficients de x_i dans le système*) par la colonne des paramètres.

En dimension 3 :

$$\begin{array}{|c|c|c|}
\hline
\left\{ \begin{array}{l} a_1^1.x_1 + a_1^2.x_2 + a_1^3.x_3 = y_1 \\ a_2^1.x_1 + a_2^2.x_2 + a_2^3.x_3 = y_2 \\ a_3^1.x_1 + a_3^2.x_2 + a_3^3.x_3 = y_3 \end{array} \right. & x_1 = & \begin{array}{|c|c|c|} \hline y_1 & a_1^2 & a_1^3 \\ y_2 & a_2^2 & a_2^3 \\ y_3 & a_3^2 & a_3^3 \\ \hline a_1^1 & a_2^1 & a_3^1 \\ a_2^1 & a_2^2 & a_2^3 \\ a_3^1 & a_3^2 & a_3^3 \\ \hline \end{array} \\
\hline
x_2 = & \begin{array}{|c|c|c|} \hline a_1^1 & y_1 & a_1^3 \\ a_2^1 & y_2 & a_2^3 \\ a_3^1 & y_3 & a_3^3 \\ \hline a_1^1 & a_2^1 & a_3^1 \\ a_2^1 & a_2^2 & a_2^3 \\ a_3^1 & a_3^2 & a_3^3 \\ \hline \end{array} & x_3 = & \begin{array}{|c|c|c|} \hline a_1^1 & a_2^1 & y_1 \\ a_2^1 & a_2^2 & y_2 \\ a_3^1 & a_3^2 & y_3 \\ \hline a_1^1 & a_2^1 & a_3^1 \\ a_2^1 & a_2^2 & a_2^3 \\ a_3^1 & a_3^2 & a_3^3 \\ \hline \end{array} \\
\hline
\end{array}$$

Ce sont ces formules qu'on appelle formules de Cramer.

Avec des notations différentielles, on note qu'on a $a_i^k = \frac{\partial y_i}{\partial x_k}$ (dérivation par rapport à x_k , avec les autres x_p "constants"). On note qu'on a aussi $\frac{(-1)^{i+k} \cdot \alpha_k^i}{\det(A)} = \frac{\partial x_k}{\partial y_i}$ (dérivation de x_k par rapport à y_i avec les autres y_q "constants").

C'est ce qui explique qu'on n'ait jamais $\frac{\partial a}{\partial x} = \frac{1}{\frac{\partial x}{\partial a}}$, ni $\frac{\partial x}{\partial \rho} \cdot \frac{\partial \rho}{\partial x} = 1$.

○ Donner la définition du rang d'une famille de vecteurs, d'une matrice, d'une application linéaire. Justifier les majorations élémentaires.

Le rang d'une famille de vecteurs est la dimension de l'espace vectoriel qu'ils engendrent.

Si cette famille est libre, son rang est donc égal à son cardinal.

Sinon, il ne peut que diminuer par rapport à celui ci.

De plus, si les vecteurs sont dans un espace vectoriel de dimension finie n , le sous-espace qu'ils engendrent ne peut pas être de dimension plus grande que n .

Une famille de rang nul n'est formée que de vecteurs nuls.

On peut montrer que le rang de la famille est son cardinal, diminué du nombre de relations de dépendance linéaire linéairement indépendantes entre elles qu'on peut écrire entre ces vecteurs. C'est une des formulations (alambiquée) du théorème du rang.

Le rang d'une application linéaire est la dimension de l'espace vectoriel image.

Cette dimension ne peut pas dépasser celle de l'espace d'arrivée (*l'image est un sous-espace vectoriel de l'espace d'arrivée*). Si il y a égalité, c'est que l'application est surjective.

On sait aussi que l'image d'une base de l'espace de départ est une famille génératrice de l'ensemble image. La dimension de l'espace image ne peut donc pas excéder la dimension de l'espace de départ.

On résume : $rg(f) \leq \min(dim(E), dim(F))$.

majoration	$rg(f) \leq dim(F)$	$rg(f) \leq dim(E)$
cas d'égalité	surjective	injective

Le rang d'une matrice est égal au nombre de colonnes linéairement indépendantes.

On peut montrer que c'est aussi le nombre de lignes linéairement indépendantes.

majoration	$rg(M) \leq NbCols(M)$	$rg(M) \leq NbLig(M)$
cas d'égalité	colonnes indépendantes	lignes indépendantes

C'est donc aussi la dimension de l'espace vectoriel engendré par les colonnes.

C'est aussi la dimension de l'ensemble image de l'application linéaire de dont M est la matrice (une fois qu'on a choisi des bases).

On se donne une matrice A à n colonnes et p lignes. On l'interprète comme la matrice d'une application linéaire de $(\mathbb{R}^n, +, \cdot)$ dans $(\mathbb{R}^p, +, \cdot)$, qu'on note f . On se donne alors une base du noyau, qu'on complète en base de $(\mathbb{R}^n, +, \cdot)$: $(\vec{e}_1, \dots, \vec{e}_r, \vec{e}_{r+1}, \dots, \vec{e}_n)$ (on considère que le noyau a pour base $(\vec{e}_{r+1}, \dots, \vec{e}_n)$). On sait alors que $(f(\vec{e}_1), \dots, f(\vec{e}_r))$ est une base de l'ensemble image. On la complète en base de $(\mathbb{R}^p, +, \cdot)$ (en notant qu'on retrouve $r \leq p$).

La matrice de f de la base $(\vec{e}_1, \dots, \vec{e}_r, \vec{e}_{r+1}, \dots, \vec{e}_n)$ vers la base $(f(\vec{e}_1), \dots, f(\vec{e}_r), \vec{e}_{r+1}, \dots, \vec{e}_p)$ est alors

de la forme $\begin{pmatrix} 1 & \dots & 0 & 0 & \dots & 0 \\ \vdots & \ddots & \vdots & \vdots & & \vdots \\ 0 & \dots & 1 & 0 & \dots & 0 \\ \vdots & & \vdots & & & \vdots \\ 0 & & 0 & 0 & & 0 \end{pmatrix}$. On note $J_{n,p,r}$ une telle matrice dont tous les termes sont

nuls, sauf les r premiers termes de la diagonale qui valent 1, ce qu'on peut écrire $1_{i=k \leq r}$.

En tenant compte des changements de base, A s'écrit $P.J_{n,p,r}.Q$ avec P et Q inversibles, de formats p sur p et n sur n (P change de base à l'arrivée et Q change de base au départ).

Réciproquement, comme la composition par des matrices inversibles ou des applications bijectives ne modifie pas le rang, on a l'équivalence entre A est de rang r et A s'écrit $P.J_{n,p,r}.Q$ avec P et Q inversibles.

On déduit alors pour A de rang r de la forme ci dessus : ${}^tA = {}^t(P.J_{n,p,r}.Q) = {}^tQ.J_{p,n,r}.{}^tP$. On reconnaît que tA est donc aussi de rang r .

C'est l'une des façons de prouver $rg(A) = rg({}^tA)$ et de déduire "le nombre de colonnes indépendantes est égal au nombre de lignes indépendantes".

On pouvait aussi l'obtenir avec "le rang est le format du plus grand déterminant non nul qu'on peut extraire de la matrice A ".

○ **Montrer que si E et F sont de même cardinal fini n , alors toute application injective de E dans F est automatiquement bijective. Même question avec "surjective implique bijective".**

○ Donner la forme "factorisée" des matrices de rang 1.

$\begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{pmatrix} \cdot (c_1 \ c_2 \ \dots \ c_p)$ avec $\begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{pmatrix}$ qui engendre l'image (de dimension 1) et

$c_1.x_1 + c_2.x_2 + \dots + c_p.x_p = 0$ équation du noyau, de dimension $\dim(E) - 1$.

○ .

○ .

○ .

○ .