

Chapitre 4 : Ensembles

PTSI B Lycée Eiffel

16 octobre 2020

La vie n'est bonne qu'à étudier et à enseigner les mathématiques.

Blaise PASCAL.

Ne tenez pour certain que ce qui est démontré.

Isaac NEWTON.

Ce chapitre un peu fourre-tout est prétexte à regrouper trois parties de cours en fait complètement indépendantes (et d'importance variable) : une partie sur le principe de récurrence et ses variantes qui contiendra également en guise d'application des calculs de sommes et de produits (partie de cours à inclure dans les fondamentaux dont la maîtrise technique est obligatoire) ; une petite partie consacrée aux rudiments de l'arithmétique (dont l'objectif est surtout de préparer un chapitre ultérieur qui sera consacré aux polynômes) ; et une partie théorique sur les notions d'applications injective et surjective, permettant de vraiment comprendre la notion essentielle en mathématiques de bijection (là encore des connaissances indispensables, même si elles ne feront pas forcément l'objet d'interrogations spécifiques).

Objectifs du chapitre :

- comprendre réellement le principe de récurrence, et savoir l'appliquer dans toutes les situations où on peut en avoir besoin (mais aussi ne PAS l'appliquer dans les situations auxquelles il ne se prête pas).
- maîtriser les calculs de sommes faisant intervenir les sommes classiques et autres techniques du type décalage d'indices ou sommes télescopiques.
- maîtriser les notions d'application injective et surjective, y compris pour des applications autres que celles de $\mathbb{R}$ dans $\mathbb{R}$.

1 Démonstration par récurrence et applications.

1.1 Principe de récurrence et variations.

Théorème 1. Tout sous-ensemble non vide de $\mathbb{N}$ admet un plus petit élément.

Remarque 1. Ce résultat fondamental pour la structure de l'ensemble $\mathbb{N}$ est plus un axiome qu'un réel théorème.

Proposition 1. Principe de récurrence.

Soit $(P_n)_{n \in \mathbb{N}}$ une suite d'énoncés mathématiques. Si les deux hypothèses suivantes sont vérifiées :

- P_0 est vrai
- pour tout entier naturel n , $P_n \Rightarrow P_{n+1}$

alors toutes les propriétés P_n sont vraies.

Démonstration. C'est en fait équivalent au théorème énoncé précédemment. Notons $A = \{n \in \mathbb{N} \mid P_n \text{ est fausse}\}$. On procède par l'absurde, supposons donc que les propriétés ne sont pas toutes vraies, ce qui revient à dire que l'ensemble A n'est pas vide. D'après le théorème précédent, il y a donc un entier n_0 qui est le plus petit élément de l'ensemble A . Cet entier ne peut pas être nul puisque P_0 est supposée vraie, on en déduit que $n_0 - 1 \in \mathbb{N}$. La propriété P_{n_0-1} est vraie puisque n_0 est le plus petit élément de A , mais P_{n_0} est fausse. C'est impossible à cause de l'hypothèse $P_n \Rightarrow P_{n+1}$, l'ensemble A est donc vide, et les propriétés sont toutes vraies. $\square$

Cette propriété sert également de pense-bête pour bien structurer la rédaction d'une démonstration par récurrence. On procède théoriquement en quatre étapes :

- **Énoncé** clair et précis des propriétés P_n et du fait qu'on va réaliser une récurrence.
- **Initialisation** : on vérifie que P_0 est vraie (habituellement un calcul très simple).
- **Hérédité** : on suppose P_n vraie pour un entier n quelconque (c'est l'hypothèse de récurrence) et on prouve P_{n+1} à l'aide de cette hypothèse (si on n'utilise pas l'hypothèse de récurrence, c'est qu'on n'avait pas besoin de faire une récurrence!).
- **Conclusion** : En invoquant le principe de récurrence, on peut affirmer avoir démontré P_n pour tout entier n .

Exemple : On considère la suite numérique définie de la façon suivante : $u_0 = 4$ et $\forall n \in \mathbb{N}$, $u_{n+1} = \frac{1}{u_n - 2} + 2$. On souhaite prouver que cette suite est minorée par 2, c'est-à-dire que $\forall n \in \mathbb{N}$, $u_n > 2$. Nous allons pour cela, bien évidemment, procéder par récurrence :

- **Énoncé** : Nous allons prouver par récurrence la propriété $P_n : u_n > 2$.
- **Initialisation** : $u_0 = 4 > 2$, donc la propriété P_0 est vérifiée.
- **Hérédité** : Supposons désormais P_n vraie, c'est-à-dire que $u_n > 2$, et essayons de prouver que $u_{n+1} > 2$. C'est en fait assez simple en partant de l'hypothèse de récurrence : $u_n > 2 \Rightarrow u_n - 2 > 0 \Rightarrow \frac{1}{u_n - 2} > 0 \Rightarrow \frac{1}{u_n - 2} + 2 > 2 \Rightarrow u_{n+1} > 2$.
- **Conclusion** : D'après le principe de récurrence, la propriété P_n est vraie pour tout entier n .

En pratique, il n'est pas nécessaire de nommer explicitement les différentes étapes de la rédaction de la récurrence, tant qu'elles sont toutes bel et bien présentes.

Remarque 2. Variations du principe de récurrence :

Le monde mathématique n'étant pas parfait, une récurrence classique n'est hélas pas toujours suffisante pour montrer certaines propriétés. Il faut donc être capable de modifier légèrement la structure dans certains cas :

- si on ne cherche à montrer P_n que lorsque $n \geq n_0$ (n_0 étant un entier fixe dépendant du contexte), on peut toujours procéder par récurrence, mais en initialisant à n_0 (bien entendu, l'hypothèse de récurrence portera alors nécessairement sur un entier $n \geq n_0$).

- il est parfois nécessaire que l'hypothèse de récurrence porte non pas sur une valeur de n , mais sur deux valeurs consécutives. On peut alors effectuer une récurrence double : on vérifie P_0 et P_1 lors de l'étape d'initialisation, et on prouve P_{n+2} à l'aide de P_n et P_{n+1} lors de l'hérédité (on peut de même effectuer des récurrences triples, quadruples, etc. en faisant une initialisation triple ou plus, et en prenant une hypothèse de récurrence triple ou plus ; dans tous les cas on ne démontre qu'une seule propriété lors de l'hérédité).
- on peut même avoir besoin pour prouver l'hérédité que la propriété soit vérifiée pour **tous** les entiers inférieurs. Dans ce cas, on parle de récurrence forte : le plus simple est de modifier la définition de la propriété P_n pour lui donner un énoncé commençant par $\forall k \leq n$. Ainsi, lorsqu'on suppose P_n vérifiée, on a une relation vraie pour toutes les valeurs de k inférieures ou égales à n (les plus malins d'entre vous noteront d'ailleurs qu'on peut toujours rédiger une récurrence sous forme de récurrence forte, ça ne demande pas plus de travail et ça ne peut pas être moins efficace ; c'est toutefois un peu plus lourd et déconseillé sauf nécessité).

Exemple : On considère la suite définie par $u_0 = 0$, $u_1 = 1$ et $\forall n \in \mathbb{N}$, $u_{n+2} = 5u_{n+1} - 6u_n$, et on veut déterminer une expression du terme général de la suite (u_n) . Pour cela (ce n'est pas forcément la meilleure méthode, mais la plus simple pour nous pour l'instant), on calcule les termes suivants de la suite : $u_2 = 5$, $u_3 = 19$, $u_4 = 65$. Une inspiration soudaine nous fait conjecturer que $u_n = 3^n - 2^n$ (si on ne devine pas la formule, on ne pourra jamais faire de récurrence), ce qu'on va prouver par récurrence double. La formule est vraie pour $u_0 : 3^0 - 2^0 = 1 - 1 = 0$ et pour $u_1 : 3^1 - 2^1 = 1$. Supposons-là vérifiée pour u_n et u_{n+1} , alors $u_{n+2} = 5u_{n+1} - 6u_n = 5(3^{n+1} - 2^{n+1}) - 6(3^n - 2^n) = 15 \times 3^n - 10 \times 2^n - 6 \times 3^n + 6 \times 2^n = 9 \times 3^n - 4 \times 2^n = 3^{n+2} - 2^{n+2}$. La formule est donc vérifiée au rang $n + 2$, le principe de récurrence double permet de conclure.

1.2 Sommes classiques.

Définition 1. Le symbole $\sum$ sert en mathématiques à désigner une somme de termes « semblables »

pouvant être décrits à l'aide d'un indice entier i . Plus précisément, la notation $\sum_{i=1}^{i=n} a_i$ se lit par

exemple « somme pour i variant de 1 à n de a_i » et peut se détailler de la façon suivante : $\sum_{i=1}^{i=n} a_i = a_1 + a_2 + \dots + a_n$.

Remarque 3. La notation est en fait exactement similaire à celle qu'on utilise pour les intégrales (ce qui est d'ailleurs tout à fait normal puisqu'une intégrale n'est rien d'autre qu'une « somme continue », d'ailleurs le symbole $\int$ est bel et bien à l'origine un 'S' comme somme). En particulier :

- La lettre i est une variable muette, autrement dit on peut la changer par n'importe quelle autre lettre sans changer la valeur de la somme. On choisit traditionnellement les lettres i , j , k , etc. pour les indices de sommes.
- Dans une somme, la variable muette prend toujours **toutes** les valeurs entières comprises entre la valeur initiale et la valeur finale.
- On ne rappelle pas systématiquement l'indice en haut de la somme puisqu'il s'agit nécessairement de la même variable qu'en-dessous de cette même somme. Ainsi, on écrit volontiers

$\sum_{i=1}^n i^2$ plutôt que $\sum_{i=1}^{i=n} i^2$. Alternativement, on peut aussi décrire cette même somme sous la forme $\sum_{1 \leq i \leq n} i^2$, même si cette dernière notation est plus classique pour les sommes doubles que nous évoquerons un peu plus loin.

Exemples : $\sum_{i=0}^n 3 = 3(n+1)$ (attention à compter correctement le nombre de termes d'une telle

somme). Plus généralement, $\sum_{i=p}^n a = a(n-p+1)$ si a est une constante.

Proposition 2. Règles de calcul sur les sommes. On a le droit d'effectuer les opérations suivantes :

- factoriser par une constante : $\sum_{i=0}^n k \times a_i = k \sum_{i=0}^n a_i$.
- séparer ou regrouper des sommes de mêmes indices : $\sum_{i=0}^n a_i + b_i = \sum_{i=0}^n a_i + \sum_{i=0}^n b_i$.
- appliquer la relation de Chasles : $\sum_{i=1}^n a_i = \sum_{i=1}^p a_i + \sum_{i=p+1}^n a_i$.
- faire un changement d'indice, par exemple : $\sum_{i=1}^n a_i = \sum_{j=0}^{n-1} a_{j+1}$ (on a ici posé $j = i-1$, de façon générale les seules choses qu'on a le droit de faire sont des décalages d'un nombre constant de valeurs, jamais une modification du genre $j = 2i$ qui transgresserait le principe « les indices prennent toujours toutes les valeurs entières possibles »).

Remarque 4. Encore une fois, ces propriétés sont les mêmes que pour les intégrales. Les deux premières peuvent être regroupées sous un même nom, on parle de linéarité du calcul de sommes (un terme que nous aurons tout le loisir d'expliquer plus en détail beaucoup plus tard cette année). La relation de Chasles est la même que pour les intégrales, mais il faut bien faire attention à ne pas compter deux fois le terme correspondant à l'indice frontière entre les deux sommes. Il existe un équivalent du changement d'indice sur les intégrales (qu'on appelle le changement de variable), mais c'est nettement plus complexe (nous l'étudierons dans le prochain chapitre).

Par contre, tenter de simplifier d'une façon ou d'une autre une somme de la forme $\sum_{i=0}^{i=n} a_i b_i$ est une très bonne manière de s'attacher la rancœur tenace de votre professeur. De façon générale, les sommes et produits ne font pas bon ménage.

Proposition 3. Sommes classiques.

- $\forall n \in \mathbb{N}, \sum_{i=1}^n i = \frac{n(n+1)}{2}$
- $\forall n \in \mathbb{N}, \sum_{i=1}^n i^2 = \frac{n(n+1)(2n+1)}{6}$
- $\forall n \in \mathbb{N}, \sum_{i=1}^n i^3 = \frac{n^2(n+1)^2}{4} = \left(\sum_{i=1}^n i \right)^2$
- $\forall q \neq 1, \forall n \in \mathbb{N}, \sum_{i=0}^n q^i = \frac{1-q^{n+1}}{1-q}$

Démonstration.

- Nous allons démontrer par récurrence que la propriété $P_n : \sum_{i=1}^n i = \frac{n(n+1)}{2}$ est vraie pour tout entier n . Pour $n = 0$, nous avons à gauche une somme vide qui vaut donc 0, ce qui est cohérent avec la formule de droite (si vraiment ça ne vous plaît pas d'avoir une somme vide, initialisez à $n = 1$). Supposons P_n vraie pour un entier n quelconque, c'est-à-dire que $\sum_{i=1}^n i = \frac{n(n+1)}{2}$. On peut alors effectuer le calcul suivant : $\sum_{i=1}^{n+1} i = \left(\sum_{i=1}^n i \right) + n + 1 = \frac{n(n+1)}{2} + n + 1 = \frac{n(n+1) + 2(n+1)}{2} = \frac{(n+1)(n+2)}{2}$, ce qui prouve P_{n+1} . On conclut en invoquant le principe de récurrence.
- Nous allons ensuite prouver de même par récurrence la propriété $P_n : \sum_{i=1}^n i^2 = \frac{n(n+1)(2n+1)}{6}$. Pour $n = 0$, c'est le même principe que ci-dessus. Supposons désormais P_n vraie pour un entier n quelconque, on peut alors écrire $\sum_{i=1}^{n+1} i^2 = \sum_{i=1}^n i^2 + (n+1)^2 = \frac{n(n+1)(2n+1)}{6} + (n+1)^2 = \frac{n(n+1)(2n+1) + 6(n+1)^2}{6} = \frac{(n+1)(n(2n+1) + 6n + 6)}{6} = \frac{(n+1)(2n^2 + 7n + 6)}{6} = \frac{(n+1)(n+2)(2n+3)}{6} = \frac{(n+1)((n+1)+1)(2(n+1)+1)}{6}$, donc P_{n+1} est vérifiée. D'après le principe de récurrence, on peut conclure que P_n est vraie pour tout entier naturel n .
- Toujours sur le même modèle, prouvons par récurrence la propriété $P_n : \sum_{i=1}^n i^3 = \frac{n^2(n+1)^2}{4}$. Je n'insiste pas plus lourdement sur l'initialisation, c'est toujours pareil. Supposons désormais P_n vraie pour un entier n quelconque, on peut alors écrire $\sum_{i=1}^{n+1} i^3 = \sum_{i=1}^n i^3 + (n+1)^3 = \frac{n^2(n+1)^2}{4} + (n+1)^3 = \frac{n^2(n+1)^2 + 4(n+1)^3}{4} = \frac{(n+1)^2(n^2 + 4n + 4)}{4} = \frac{(n+1)^2(n+2)^2}{4}$, donc P_{n+1} est vérifiée. D'après le principe de récurrence, on peut conclure que P_n est vraie pour tout entier naturel n .
- Nous allons enfin prouver par récurrence la dernière propriété $P_n : \sum_{i=0}^n q^i = \frac{1-q^{n+1}}{1-q}$. Pour $n = 0$, nous avons $\sum_{i=0}^0 q^i = q^0 = 1$, et $\frac{1-q^1}{1-q} = 1$, donc P_0 est vérifiée. Supposons désormais P_n vraie pour un entier n quelconque, on peut alors écrire $\sum_{i=0}^{n+1} q^i = \sum_{i=0}^n q^i + q^{n+1} = \frac{1-q^{n+1}}{1-q} + q^{n+1} = \frac{1-q^{n+1} + q^{n+1} - q^{n+2}}{1-q} = \frac{1-q^{n+2}}{1-q}$, donc P_{n+1} est vérifiée. D'après le principe de récurrence, on peut conclure que P_n est vraie pour tout entier naturel n . □

Remarque 5. Il existe d'autres façons de démontrer ces formules. Heureusement d'ailleurs, puisque la démonstration par récurrence nécessite de déjà connaître la formule avant de pouvoir la démontrer. Parmi les méthodes célèbres pour la première somme, celle-ci est souvent citée comme ayant été découverte par Gauss à l'âge de deux ans et demi (à peu près, parfois un peu plus selon les sources). On écrit la somme dans un sens : $S_n = 1 + 2 + \dots + n$, puis on écrit la même somme dans l'autre sens : $S_n = n + (n-1) + \dots + 2 + 1$, et on additionne les deux sommes terme à terme pour obtenir $2S_n = (n+1) + (n+1) + \dots + (n+1) = n(n+1)$ dont on déduit immédiatement le résultat.

1.3 Compléments techniques concernant les calculs de sommes.

Définition 2. Une **somme télescopique** est une somme de la forme $\sum_{i=0}^n a_{i+1} - a_i$, que l'on peut simplifier sous la forme $a_{n+1} - a_0$.

En effet, presque tous les termes de la somme « se simplifient ». On peut l'écrire explicitement :

$$\sum_{i=1}^n a_{i+1} - a_i = (a_1 - a_0) + (a_2 - a_1) + (a_3 - a_2) + \cdots + (a_n - a_{n-1}) + (a_{n+1} - a_n) = a_{n+1} - a_0. \text{ On}$$

peut bien sûr rédiger les choses plus rigoureusement en faisant un décalage d'indices : $\sum_{i=0}^n a_{i+1} - a_i =$

$$\sum_{i=0}^n a_{i+1} - \sum_{i=0}^n a_i = \sum_{i=1}^{n+1} a_i - \sum_{i=0}^n a_i = \left(\sum_{i=1}^n a_i \right) + a_{n+1} - a_0 - \left(\sum_{i=1}^n a_i \right) = a_{n+1} - a_0 \text{ (on a « sorti » des deux sommes les termes qui ne sont pas communs aux deux sommes pour ensuite les simplifier, ce qui est la technique habituelle pour gérer ces sommes télescopiques).}$$

Exemple : considérons la somme $S_n = \sum_{i=1}^n \frac{1}{i(i+1)}$. A priori pas évident à calculer, du moins tant

qu'on a pas constaté que $\frac{1}{i} - \frac{1}{i+1} = \frac{i+1-i}{i(i+1)} = \frac{1}{i(i+1)}$. On peut alors faire le calcul suivant :

$$S_n = \sum_{i=1}^n \frac{1}{i(i+1)} = \sum_{i=1}^n \frac{1}{i} - \sum_{i=1}^n \frac{1}{i+1} = \sum_{i=1}^n \frac{1}{i} - \sum_{j=2}^{n+1} \frac{1}{j} = 1 + \sum_{i=2}^n \frac{1}{i} - \sum_{j=2}^n \frac{1}{j} - \frac{1}{n+1} = 1 - \frac{1}{n+1}.$$

Si la fin du calcul ne vous semble pas claire, on peut aussi voir les choses ainsi :

$$\sum_{i=1}^n \frac{1}{i} - \frac{1}{i+1} = 1 - \frac{1}{2} + \frac{1}{2} - \frac{1}{3} + \cdots + \frac{1}{n} - \frac{1}{n+1} = 1 - \frac{1}{n+1}.$$

Exemple : un exemple un peu plus compliqué est le calcul de $S_n = \sum_{i=1}^n \frac{1}{i^3 + 3i^2 + 2i}$. Le seul espoir

pour calculer ce genre de somme est de faire apparaître un télescopage (on ne connaît aucune formule permettant de calculer des sommes faisant intervenir des expressions polynômiales au dénominateur d'une fraction), ce qui sera en fait systématiquement le cas à condition d'arriver à « séparer » la fraction en plusieurs termes pouvant se télescoper. Il existe pour cela un théorème bien utile que nous recroiserons dans le prochain chapitre, celui de la **décomposition en éléments simples**. Comme ce théorème n'est pas au programme dans sa version complète, je vous en donne pour l'instant une version très simplifiée qui suffira largement pour les quelques applications que nous en ferons :

Théorème : si un polynôme P peut se factoriser sous la forme $P(x) = (x - \alpha_1)(x - \alpha_2) \cdots (x - \alpha_n)$, où $\alpha_1, \alpha_2, \dots, \alpha_n$ sont des réels deux à deux distincts, alors il existe n réels $(a_i)_{i \in \{1, 2, \dots, n\}}$ tels que

$$\frac{1}{P(x)} = \frac{a_1}{x - \alpha_1} + \frac{a_2}{x - \alpha_2} + \cdots + \frac{a_n}{x - \alpha_n}.$$

La première étape de notre calcul consiste donc à factoriser notre dénominateur pour faire apparaître les racines α_i du théorème. Ici, c'est facile : $i^3 + 3i^2 + 2i = i(i^2 + 3i + 2) = i(i+1)(i+2)$. Le théorème de décomposition en éléments simples assure alors qu'on peut trouver trois constantes

a, b et c telles que $\frac{1}{i^3 + 3i^2 + 2i} = \frac{a}{i} + \frac{b}{i+1} + \frac{c}{i+2}$. Pour déterminer ces constantes, il existe des méthodes astucieuses, mais nous nous contenterons pour l'instant d'une bête identification en met-

tant les termes du membre de droite de l'égalité au même dénominateur : $\frac{a}{i} + \frac{b}{i+1} + \frac{c}{i+2} = \frac{a(i+1)(i+2) + bi(i+2) + ci(i+1)}{i(i+1)(i+2)} = \frac{(a+b+c)i^2 + (3a+2b+c)i + 2a}{i^3 + 3i^2 + 2i}$. Le numérateur de cette

fraction devant être constant égal à 1, nous n'avons pas d'autre choix que d'imposer les conditions $a + b + c = 3a + 2b + c = 0$ et $2a = 1$, donc $a = \frac{1}{2}$. On en déduit que $b + c = -\frac{1}{2}$ et $2b + c = -\frac{3}{2}$, et une soustraction de ces deux équations donne $b = -1$, puis $c = \frac{1}{2}$. Autrement dit, $\frac{1}{i^3 + 3i^2 + 2i} = \frac{1}{2i} - \frac{1}{i+1} + \frac{1}{2(i+2)}$.

La dernière étape consiste à écrire le télescopage, qui va ici être nécessairement plus compliqué que dans l'exemple précédent puisqu'il va faire intervenir trois termes et non deux (on peut écrire une partie de la somme « avec des pointillés » pour bien visualiser ce qui se passe les premières fois). Les techniques sont toutefois toujours les mêmes, on effectue des décalages d'indices pour avoir la même expression dans chaque somme, puis on sort des sommes les termes qui n'apparaissent pas dans les trois sommes pour pouvoir les simplifier : $S_n = \frac{1}{2} \sum_{i=1}^n \frac{1}{i} - \sum_{i=1}^n \frac{1}{i+1} + \frac{1}{2} \sum_{i=1}^n \frac{1}{i+2} = \frac{1}{2} \sum_{i=1}^n \frac{1}{i} - \sum_{i=2}^{n+1} \frac{1}{i} + \frac{1}{2} \sum_{i=3}^{n+2} \frac{1}{i} = \frac{1}{2} \sum_{i=1}^n \frac{1}{i} - \frac{1}{2} + \frac{1}{2} \left(\sum_{i=3}^n \frac{1}{i} \right) - \frac{1}{2} - \left(\sum_{i=3}^n \frac{1}{i} \right) - \frac{1}{n+1} + \frac{1}{2} \left(\sum_{i=3}^n \frac{1}{i} \right) + \frac{1}{2(n+1)} + \frac{1}{2(n+2)} = \frac{1}{4} - \frac{1}{2n+2} + \frac{1}{2n+4}$.

Définition 3. Une **somme double** est une somme dépendant de deux indices, et habituellement notée sous la forme de deux sommes imbriquées, par exemple $\sum_{i=1}^n \sum_{j=1}^i \frac{i}{j}$, ou sous la forme plus simple d'une seule somme à deux indices, comme $\sum_{1 \leq i, j \leq n} \frac{i}{j}$ (qui désigne exactement la même chose que la notation précédente). La somme double prise en exemple est une somme de n^2 termes qu'on peut représenter dans le tableau à double entrée suivant :

$i \backslash j$	1	2	3	...	$n-1$	n
1	1	$\frac{1}{2}$	$\frac{1}{3}$	...	$\frac{1}{n-1}$	$\frac{1}{n}$
2	2	1	$\frac{2}{3}$	...	$\frac{2}{n-1}$	$\frac{2}{n}$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\vdots$
n	n	$\frac{n}{2}$	$\frac{n}{3}$	...	$\frac{n}{n-1}$	1

Remarque 6. Si on décide d'écrire la somme précédente en utilisant deux symboles somme, on peut au choix l'écrire sous la forme $\sum_{i=1}^n i \times \sum_{j=1}^i \frac{1}{j}$, ou sous la forme $\sum_{j=1}^n \frac{1}{j} \times \sum_{i=j}^n i$. On calculera en pratique d'abord la somme intérieure, ce qui influencer le choix de l'ordre des deux sommes. Dans certains cas, il sera même impossible de calculer la somme double en l'écrivant dans un sens, alors que le calcul est tout à fait faisable dans l'autre (ici, ça ne change rien, la somme dépendant de l'indice j ne se calcule pas, qu'on la mette à l'intérieur ou à l'extérieur).

Exemple : On cherche à calculer la somme double $S_n = \sum_{1 \leq j \leq i \leq n} 3j$. La condition $j \leq i$ qui est imposée sur les deux indices de la somme revient à dire que, si on représente les termes de cette somme dans un tableau à double entrée comme ci-dessus, on ne prendra en compte que les termes situés sous la diagonale du tableau (diagonale incluse). On peut toujours écrire la somme double à l'aide de deux symboles somme en choisissant l'ordre des indices, mais attention à les écrire correctement, la condition $i \leq j$ apparaîtra toujours dans la somme intérieure et jamais dans la somme extérieure.

Autrement dit, on peut écrire au choix $S_n = \sum_{i=1}^n \sum_{j=1}^i 3j$, ou $S_n = \sum_{j=1}^n \sum_{i=j}^n 3j$. Comme il est nettement

plus pratique d'avoir des sommes démarrant à un indice fixe égal à 1 pour appliquer aisément les formules du cours, on préférera en pratique la première écriture.

$$\begin{aligned} \text{On calcule donc } S_n &= 3 \sum_{i=1}^n \sum_{j=1}^i = 3 \sum_{i=1}^n \frac{i(i+1)}{2} = \frac{3}{2} \sum_{i=1}^n i^2 + i = \frac{3}{2} \left(\frac{n(n+1)(2n+1)}{6} + \frac{n(n+1)}{2} \right) = \\ &= \frac{n(n+1)(2n+1) + 3n(n+1)}{4} = \frac{(n+1)(2n^2 + n + 3n)}{4} = \frac{n(n+1)(n+2)}{2}. \end{aligned}$$

Exemple : Si on inverse la condition sur les deux indices dans le calcul de somme double précédent, le résultat n'est plus du tout le même. En effet, $\sum_{1 \leq i \leq j \leq n} 3j = \sum_{j=1}^n \sum_{i=1}^j 3j = \sum_{j=1}^n 3j^2 = \frac{n(n+1)(2n+1)}{2}$.

1.4 Produits.

Définition 4. Suivant le même principe que pour les sommes, le symbole $\prod$ signifie « produit ».

Par exemple, $\prod_{i=1}^n a_i = a_1 \times a_2 \times \cdots \times a_n$.

Proposition 4. Les règles de calcul suivantes peuvent être utiles quand on manipule des produits :

- séparation ou regroupement de produits : $\prod_{i=1}^n a_i \times \prod_{i=1}^n b_i = \prod_{i=1}^n a_i b_i$.
- relation de Chasles : $\prod_{i=1}^n a_i = \prod_{i=1}^p a_i \times \prod_{i=p+1}^n a_i$.
- changement d'indice : $\prod_{i=2}^{n+1} a_i = \prod_{j=1}^n a_{j+1}$

Remarque 7. Bien entendu, tenter de simplifier $\prod_{i=1}^n (a_i + b_i)$ serait une grave erreur que, j'en suis certain, vous ne commettrez pas deux fois (ni même une seule, si possible).

Il n'existe pas d'équivalent pour les produits de la factorisation par une constante, si on a une constante a en facteur d'un produit contenant n termes, c'est un facteur a^n qu'on pourra sortir du produit en séparant simplement le produit en deux produits.

Définition 5. On appelle **factorielle** de l'entier naturel n , et on note $n!$, le nombre $n! = \prod_{i=1}^n i$.

Exemple : Un exemple de calcul de produit peu évident pour terminer ce paragraphe. On veut écrire sous une forme simple le produit des n premiers entiers impairs $P_n = 1 \times 3 \times 5 \times \cdots \times (2n-1)$. La meilleure méthode consiste à « boucher les trous » en multipliant (et en divisant pour compenser) par les entiers pairs, car un produit d'entiers pairs se gère beaucoup plus facilement (on peut mettre des 2 en facteur sur chaque terme !) :

$$\begin{aligned} P_n &= 3 \times 5 \times 7 \times \cdots \times (2n-1) = \frac{2 \times 3 \times 4 \times \cdots \times (2n-1) \times (2n)}{2 \times 4 \times \cdots \times 2n} \\ &= \frac{(2n)!}{(2 \times 1) \times (2 \times 2) \times (2 \times 3) \times \cdots \times (2 \times (n-1)) \times (2 \times n)} = \frac{(2n)!}{2^n \times n!}. \end{aligned}$$

2 Arithmétique.

Définition 6. Soient n et p deux entiers relatifs, n est **divisible par** p (ou p divise n) s'il existe un troisième entier k tel que $n = kp$.

Remarque 8. La relation de divisibilité est une relation d'ordre sur $\mathbb{N}^*$ mais pas sur $\mathbb{Z}^*$, où elle n'est pas antisymétrique. En effet, deux entiers relatifs qui se divisent l'un l'autre sont soit égaux soit opposés.

Théorème 2. Division euclidienne.

Soit $n \in \mathbb{Z}$ et $p \in \mathbb{N}^*$, alors il existe un unique couple $(q, r) \in \mathbb{Z}^2$ tel que $n = pq + r$, et $0 \leq r < p$. L'entier q est appelé **quotient** de la division euclidienne de n par p , et l'entier r **reste** de cette même division.

Démonstration. Commençons par prouver l'existence du couple (q, r) , en supposant $n > 0$ (sinon, ce n'est pas beaucoup plus compliqué). Il existe certainement un entier a à partir duquel $ap > n$ (cette propriété présentée ici comme une évidence s'exprime de la façon suivante en langage mathématique pédant : « $\mathbb{R}$ est archimédien »). Notons alors $q = \max\{a \in \mathbb{N} \mid ap \leq n\}$, et $r = n - pq$. Par définition, $pq \leq n$, donc $r \geq 0$. De plus, par maximalité de q , on doit avoir $(q+1)p > n$, soit $pq + p - n > 0$, ou encore $p > n - pq = r$. Enfin, par définition de r , $n = pq + r$, l'existence du couple est donc prouvée.

Démontrons désormais l'unicité par l'absurde (c'est très classique pour démontrer des résultats d'unicité) en supposant qu'il y a deux couples convenables (q, r) et (q', r') . On a alors $pq + r = pq' + r' = n$, donc $p(q - q') = r' - r$. En particulier $r' - r$ divise p , alors que $-p < r' - r < p$. Ce n'est possible que si $r' - r = 0$, soit $r' = r$, ce qui implique $p(q' - q) = 0$, donc $q = q'$. Les deux couples sont alors identiques. $\square$

Définition 7. Un entier naturel n est **premier** s'il n'est divisible que par 1 et par lui-même.

Remarque 9. Par convention, le nombre 1 n'est pas considéré comme un nombre premier.

Remarque 10. Il n'existe pas de méthode extrêmement simple pour savoir si un entier donné est premier ou non. Pour faire la liste des nombres premiers inférieurs à un entier donné, le plus simple est encore d'utiliser le **crible d'Eratosthène** : on place dans un tableau tous les entiers inférieurs à n (à partir de 2), et à chaque étape on entoure le plus petit entier disponible (qui sera nécessairement premier), et on raye tous ses multiples.

	(2)	(3)	4	(5)	6	(7)	8	9	10
(11)	12	(13)	14	15	16	(17)	18	(19)	20
21	22	(23)	24	25	26	27	28	(29)	30
(31)	32	33	34	35	36	(37)	38	39	40
(41)	42	(43)	44	45	46	(47)	48	49	50
51	52	(53)	54	55	56	57	58	(59)	60
(61)	62	63	64	65	66	(67)	68	69	70
(71)	72	(73)	74	75	76	77	78	(79)	80
81	82	(83)	84	85	86	87	88	(89)	90
91	92	93	94	95	96	(97)	98	99	100

Les nombres restants : 2, 3, 5, 7, 11, 13, 17, 19, 23, 29, 31, 37, 41, 43, 47, 53, 59, 61, 67, 71, 73, 79, 83, 89 et 97 sont tous premiers. Il y a donc 25 nombres premiers inférieurs ou égaux à 100. On peut démontrer plus généralement (mais c'est très compliqué) que le nombre de nombres premiers inférieurs à n devient proche de $\frac{n}{\ln(n)}$ (c'est-à-dire que le quotient des deux tend vers 1, résultat connu sous le nom de théorème des nombres premiers), lorsque n tend vers $+\infty$.

Définition 8. Soit n et p deux entiers naturels non nuls. Le **plus grand commun diviseur** (ou pgcd) de n et p est, comme son nom l'indique, le plus grand entier divisant simultanément n et p . On le note parfois $n \wedge p$. Le **plus petit commun multiple** (ou ppcm) est le plus petit entier naturel que divisent n et p .

Proposition 5. Quels que soient les nombres entiers n et p , le produit du pgcd et du ppcm de n et de p est égal à np .

Définition 9. Deux entiers n et p sont **premiers entre eux** si leur pgcd est égal à 1.

Théorème 3. Théorème de Bezout.

Deux entiers n et p sont premiers entre eux si et seulement s'il existe un couple d'entiers $(a, b) \in \mathbb{Z}^2$ tels que $an + bp = 1$. Plus généralement, il existe toujours un couple d'entiers relatifs tels que $an + bp = n \wedge p$.

Démonstration. La preuve dépasse largement ce qu'on est capables de faire à notre niveau (pour information, le calcul explicite des entiers a et b peut s'effectuer à l'aide d'une variante de l'algorithme d'Euclide, qui est lui-même un algorithme classique de calcul du pgcd de deux entiers). D'ailleurs, le théorème lui-même est vraiment à la limite du programme. $\square$

Théorème 4. Il existe une infinité de nombres premiers.

Démonstration. On va un tout petit peu anticiper sur le dernier résultat du paragraphe. Faisons un raisonnement par l'absurde : il existerait donc une liste finie $p_1, p_2, \dots, p_k$ de nombres premiers.

Notons alors $p = \prod_{i=1}^k p_i + 1$. Cet entier n'est sûrement pas divisible par p_1 (puisque l'entier qui le précède l'est et que $p_1 \geq 2$), ni par aucun des p_i . Soit il est lui-même premier (mais distinct des autres, ce qui est absurde), soit il est divisible par un entier premier (c'est là qu'on utilise la décomposition), qui n'est lui-même aucun des p_i . Dans tous les cas, on aboutit à une contradiction. $\square$

Théorème 5. Décomposition en facteurs premiers.

Tout nombre entier $n \geq 2$ peut se décomposer de façon unique sous la forme $n = \prod_{i=1}^{i=k} p_i^{\alpha_i} = p_1^{\alpha_1} \times \dots \times p_k^{\alpha_k}$, où $p_1 < p_2 < \dots < p_k$ sont des nombres premiers distincts, et $(\alpha_1, \dots, \alpha_k) \in (\mathbb{N}^*)^k$.

Démonstration. La démonstration de ce théorème n'est pas au programme, on va s'en dispenser. On peut la faire par récurrence, mais c'est un peu technique. $\square$

Exemple : La décomposition du nombre 384 est $2^7 \times 3$ (il suffit de diviser par 2 jusqu'à ce que ce ne soit plus possible, de recommencer avec 3, etc), et celle de 660 est $2^2 \times 3 \times 5 \times 11$. Il est assez facile ensuite de calculer le pgcd et le ppcm. Pour le pgcd, il suffit d'élever chaque facteur premier apparaissant à la fois dans la décomposition de n et de p à la puissance égale au minimum des puissances des deux décompositions. Ici, $384 \wedge 660 = 2^2 \times 3 = 12$. Ensuite, le ppcm est égal au produit des deux entiers divisé par le pgcd, ou au produit des facteurs premiers apparaissant dans au moins une des deux décompositions, élevés à la puissance égale au maximum des deux puissances. Ici, $\text{ppcm}(384, 660) = 2^7 \times 3 \times 5 \times 11 = \frac{384 \times 660}{12} = 21\,120$.

3 Applications.

3.1 Définitions.

Une application est une sorte de généralisation de ce que vous avez l'habitude d'appeler une fonction. La différence principale est qu'une application doit être définie sur tout son ensemble de départ (en gros, l'ensemble de définition fait partie de la définition de l'objet mathématique), alors qu'on parle par exemple de fonction de $\mathbb{R}$ dans $\mathbb{R}$ pour la fonction inverse (mais on peut très bien parler de l'application inverse de $\mathbb{R}^*$ dans $\mathbb{R}$). Par ailleurs, les applications peuvent être définies sur des ensembles nettement plus variés et exotiques que $\mathbb{R}$.

Définition 10. Une **application** f est la donnée d'un ensemble E , appelé ensemble de départ de l'application, d'un ensemble F appelé ensemble d'arrivée, et pour chaque élément x de E , d'un unique élément de F noté $f(x)$. On appelle $f(x)$ l'**image** de l'élément x par f , et si $y \in F$, les éléments x de E vérifiant $f(x) = y$ sont appelés **antécédents** de y par f (un élément y peut très bien ne pas avoir d'antécédent, ou au contraire en avoir plusieurs). Une application est en général notée $\begin{cases} E & \rightarrow & F \\ x & \mapsto & f(x) \end{cases}$.

Exemples : Toutes les fonctions usuelles sont des applications, à condition de bien préciser leur ensemble de définition. Ainsi, $\begin{cases} \mathbb{R}^+ & \rightarrow & \mathbb{R} \\ x & \mapsto & \sqrt{x} \end{cases}$. Il existe aussi énormément d'applications définies sur toutes sortes d'ensembles mathématiques, ou même d'ensembles qui n'ont pas grand intérêt mathématique, comme $\begin{cases} E & \rightarrow & \llbracket 1, 15 \rrbracket \\ x & \mapsto & \text{groupe de colles de } x \end{cases}$ est une application définie sur l'ensemble E des élèves de la classe (la notation utilisée pour l'ensemble d'arrivée désigne les entiers compris entre 1 et 15, on pourrait mettre à la place un ensemble d'arrivée égal à $\mathbb{N}$, ça ne poserait pas de problème, l'ensemble d'arrivée n'a aucune raison d'être restreint à l'ensemble des valeurs effectivement prises par l'application).

Remarque 11. Deux applications sont identiques si elles ont même ensemble de départ, même ensemble d'arrivée et envoient un même élément sur une même image. Par exemple, les fonctions d'une variable réelle $f : x \mapsto x - 4$ et $g : x \mapsto \frac{x^2 - 5x + 4}{x - 1}$ sont différentes, même si elles coïncident sur $\mathbb{R} \setminus \{1\}$: elles n'ont pas le même ensemble de définition.

Définition 11. Si E est un ensemble, l'**application identité** sur l'ensemble E est l'application $\begin{cases} E & \rightarrow & E \\ x & \mapsto & x \end{cases}$. On la note id_E , ou simplement id s'il n'y a pas d'ambiguïté possible sur l'ensemble E .

Définition 12. Soit $f : E \rightarrow F$ une application et E' un sous-ensemble de E . L'application $g : \begin{cases} E' & \rightarrow & F \\ x & \mapsto & f(x) \end{cases}$ est appelée **restriction** de f au sous-ensemble E' et notée $f|_{E'}$.

Réciproquement, si $E \subset G$, tout application $h : \begin{cases} H & \rightarrow F \\ x & \mapsto h(x) \end{cases}$ telle que $h(x) = f(x)$ pour tout élément $x \in E$ est un **prolongement** de f à E (il existe énormément de prolongements possibles, mais la restriction est par construction unique).

Exemple : La fonction $x \rightarrow x \ln x$, définie sur $\mathbb{R}^{+*}$, peut se prolonger en une **unique** fonction $\tilde{f}$ définie et continue sur $\mathbb{R}^+$ en posant $\tilde{f}(0) = 0$ (et bien sûr $\tilde{f}(x) = x \ln(x)$ si $x > 0$). Ce prolongement est appelé **prolongement par continuité** de f à $\mathbb{R}^+$ (c'est le fait d'imposer la conservation de la propriété de continuité qui rend ce prolongement unique et donc remarquable). En pratique, on utilise souvent la même notation pour désigner le prolongement que pour la fonction d'origine, même si c'est un abus de notation.

Définition 13. Soit $A \subset E$, la **fonction caractéristique** (ou **fonction indicatrice**) du sous-ensemble A est l'application $1_A : E \rightarrow \{0, 1\}$ définie par $1_A(x) = 1$ si $x \in A$ et $1_A(x) = 0$ si $x \notin A$.

3.2 Injections, surjections, bijections.

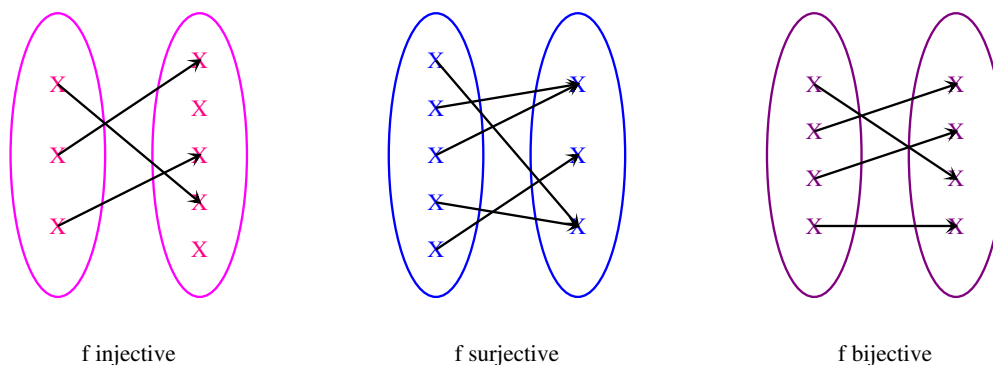
Définition 14. Soit $f : E \rightarrow F$ une application.

- f est **injective** si $\forall (x, x') \in E^2, f(x) = f(x') \Rightarrow x = x'$ (ou, de façon équivalente, si $\forall (x, x') \in E^2, f(x) \neq f(x') \Rightarrow x \neq x'$).
- f est **surjective** si $\forall y \in F, \exists x \in E, f(x) = y$.
- f est **bijjective** si elle est à la fois injective et surjective.

Remarque 12. Autrement dit, f est injective si tout élément de F a **au plus** un antécédent par f , surjective si tout élément de F a **au moins** un antécédent par f , et bijective si tout élément de F a **exactement** un antécédent par f (ce qui correspond bien à la définition de la bijectivité que nous avons déjà vue dans le chapitre consacré aux fonctions usuelles).

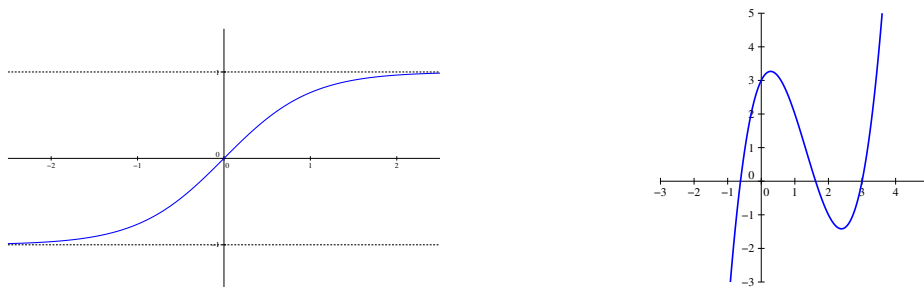
Exemples : L'application $x \mapsto x^2$, qui va de $\mathbb{R}$ dans $\mathbb{R}^+$, est surjective (tout réel positif admet une racine carrée) mais pas injective car par exemple 2 et -2 ont la même image par f . L'application racine carrée est par contre bijective de $\mathbb{R}^+$ dans lui-même.

L'application qui à tout élève de la classe associe son numéro de groupe de colles n'est pas injective (deux élèves différents appartenant au même trinôme auront la même image), elle est surjective si on limite l'ensemble d'arrivée à $\llbracket 1, 15 \rrbracket$, mais pas si l'ensemble d'arrivée est $\mathbb{N}$ tout entier. De façon générale, on peut toujours rendre une application surjective en modifiant son ensemble d'arrivée. Pour la rendre injective, c'est beaucoup plus délicat car c'est dans ce cas l'ensemble de départ qu'il faut restreindre. Une illustration des notions d'injectivité et de surjectivité dans le cas d'ensembles finis :



On notera que, si E et F sont deux ensembles finis, et f une application de E vers F , f ne peut être injective que si F contient au moins autant d'éléments que E (mais ce n'est bien sûr pas une condition suffisante!), et ne peut être surjective que si E contient au moins autant d'éléments que F . En particulier, s'il existe une bijection entre deux ensembles finis, alors ces ensembles ont nécessairement le même nombre d'éléments.

Et pour ceux qui préfèrent avec des fonctions, un exemple de fonction de $\mathbb{R}$ dans $\mathbb{R}$ injective mais pas surjective à gauche (on voit que les valeurs supérieures à 1 par exemple n'ont pas d'antécédent), et un de fonction surjective mais pas injective à droite (par exemple 0 a trois antécédents par cette fonction) :



Proposition 6. Soient $f : E \rightarrow F$ et $g : F \rightarrow G$ deux applications. Si f et g sont injectives alors $g \circ f$ est injective. Si f et g sont surjectives, alors $g \circ f$ est surjective.

Démonstration. Supposons g et f injectives, et soient $(x, x') \in E^2$ tels que $g(f(x)) = g(f(x'))$. Par injectivité de g , on a alors nécessairement $f(x) = f(x')$, puis par injectivité de f , $x = x'$, ce qui prouve l'injectivité de $g \circ f$. Supposons désormais g et f surjectives et soit $z \in G$. Par surjectivité de g , $\exists y \in F, z = g(y)$, puis par surjectivité de f , $\exists x \in E, y = f(x)$. Mais alors $z = g \circ f(x)$, donc z a un antécédent par $g \circ f$, ce qui prouve la surjectivité de cette application. $\square$

Remarque 13. La réciproque de ces propriétés est totalement fausse, comme le prouve l'exemple suivant :

- on définit une première application $f : \begin{cases} \mathbb{N} & \rightarrow & \mathbb{N} \\ n & \mapsto & n^2 \end{cases}$.
- et une deuxième application $g : \begin{cases} \mathbb{N} & \rightarrow & \mathbb{N} \\ n & \mapsto & \text{Ent}(\sqrt{n}) \end{cases}$.

Le but de l'exercice est d'étudier l'injectivité et la surjectivité des quatre applications $f, g, f \circ g$ et $g \circ f$.

- l'application f est certainement injective (deux entiers naturels distincts ne peuvent pas avoir le même carré) mais pas surjective (par exemple 2 n'a pas d'antécédent entier par cette application).
- au contraire, l'application g n'est pas injective (on a par exemple $g(3) = g(2) = 1$), mais elle est bien surjective : quel que soit l'entier naturel p , l'entier p^2 est un antécédent de p par g de façon évidente.
- il est assez facile de comprendre que $f \circ g$ ne peut être ni injective (2 ne peut pas avoir d'antécédent entier par $f \circ g$ puisque $g(n)$ est toujours un entier, dont le carré ne peut donc pas être égal à 2), ni surjective (comme pour g , on a $f \circ g(2) = f \circ g(3)$).
- au contraire, $g \circ f$ est bijective, puisqu'on a tout bêtement $g \circ f = \text{id}_{\mathbb{N}}$.

Proposition 7. Une application $f : E \rightarrow F$ est bijective si et seulement si il existe $g : F \rightarrow E$ telle que $g \circ f = id_E$ et $f \circ g = id_F$. L'application g est alors appelée **bijection réciproque** de f (ou réciproque tout court) et notée f^{-1} .

Remarque 14. Rappelons que cette réciproque, bien que notée f^{-1} , n'a rien à voir avec la fonction inverse de f , que pour cette raison nous noterons toujours $\frac{1}{f}$. Notons au passage que f^{-1} est effectivement bijective, de réciproque f (c'est évident une fois le théorème démontré).

Démonstration. Supposons f bijective et soit $y \in F$. Il existe un unique antécédent x de y par f , on pose simplement $g(y) = x$. On a alors par construction $f \circ g(x) = x$, donc $f \circ g = id_F$. De plus, si $x \in E$, $g(f(x))$ est un antécédent de $f(x)$, mais comme il n'y en qu'un ça ne peut être que x , donc on a aussi $g \circ f = id_E$.

Réciproquement, si $g \circ f = id_E$ et $f \circ g = id_F$, considérons x et x' tels que $f(x) = f(x')$, on a alors $g \circ f(x) = g \circ f(x')$, donc $x = x'$, ce qui prouve l'injectivité de f . Soit maintenant $y \in F$, alors $g(y)$ est un antécédent de y par f puisque $f \circ g(y) = y$, donc f est surjective. L'application f est donc bijective. □

Proposition 8. Soient $f : E \rightarrow F$ et $g : F \rightarrow G$ deux applications bijectives, alors $g \circ f : E \rightarrow G$ est une application bijective et $(g \circ f)^{-1} = f^{-1} \circ g^{-1}$.

Démonstration. f et g étant à la fois injectives et surjectives, $g \circ f$ est à la fois injective et surjective (cf plus haut) donc bijective. De plus, $\forall x \in E$, $f^{-1} \circ g^{-1} \circ g \circ f(x) = f^{-1}((g^{-1} \circ g)(f(x))) = f^{-1}(f(x)) = x$ et de même $\forall x \in G$, $g \circ f \circ f^{-1} \circ g^{-1}(x) = x$. □

Définition 15. Soit $f : E \rightarrow F$ une application et $A \subset E$. On appelle **image** (directe) de A l'ensemble des images des éléments de A : $f(A) = \{y \in F \mid \exists x \in A, f(x) = y\}$. Soit maintenant $B \subset F$, on appelle **image réciproque** de B par F l'ensemble des antécédents d'éléments de B : $f^{-1}(B) = \{x \in E \mid f(x) \in B\}$.

Remarque 15. La deuxième notation n'a pas été choisie de façon contradictoire avec la définition d'application réciproque (encore heureux). Si f est bijective, l'image réciproque d'une partie B de F est confondue avec son image directe par f^{-1} .

Exemple : Considérons l'application $f : x \mapsto x^2$ de $\mathbb{R}$ dans $\mathbb{R}$, alors $f([2, 5]) = [4, 25]$; $f([-1, 3]) = [0, 9]$; $f^{-1}([4, 9]) = [-3, -2] \cup [2, 3]$.

Définition 16. L'ensemble des parties d'un ensemble E est l'ensemble dont les éléments sont tous les sous-ensembles de E (y compris les sous-ensembles « triviaux » que sont le sous-ensemble vide et l'ensemble E tout entier). On le note $\mathcal{P}(E)$.

Exemple : si E est un ensemble à trois éléments $E = \{a, b, c\}$, son ensemble des parties contiendra huit éléments : $\mathcal{P}(E) = \{\emptyset, \{a\}, \{b\}, \{c\}, \{a, b\}, \{a, c\}, \{b, c\}, \{a, b, c\}\}$ (attention aux notations, les éléments de ce nouvel ensemble sont eux-même des ensembles). On peut se convaincre assez facilement que, si E est un ensemble fini à n éléments, $\mathcal{P}(E)$ contiendra 2^n éléments (nous reviendrons sur ce résultat dans le chapitre que nous consacrerons au dénombrement). Ça marche même lorsque l'ensemble E ne contient aucun élément : $\mathcal{P}(\emptyset) = \{\emptyset\}$ est bien un ensemble à $2^0 = 1$ élément, dont l'unique élément est l'ensemble vide ! Et, bien sûr, $\mathcal{P}(\mathcal{P}(\emptyset)) = \{\emptyset, \{\emptyset\}\}$ est un ensemble à deux éléments.

Définition 17. Si E et F sont deux ensembles, la notation F^E désigne l'ensemble de toutes les applications de E dans F .

Remarque 16. Attention c'est bien l'ensemble de départ qui est en exposant, et l'ensemble d'arrivée à la base, dans cette notation. On peut par exemple noter (même si c'est assez peu employé) $\mathbb{R}^{\mathbb{N}}$ l'ensemble des suites réelles définies pour tout entier naturel.

Proposition 9. L'application $\varphi : \begin{cases} \mathcal{P}(E) & \rightarrow & \{0, 1\}^E \\ A & \mapsto & 1_A \end{cases}$ est une bijection.

Démonstration. Cette propriété explique bien le nom de fonction caractéristique puisqu'il s'agit de caractériser un sous-ensemble de E par une application de E dans $\{0, 1\}$. En particulier, elle explique pourquoi, dans le cas d'un ensemble fini E à n éléments, $\mathcal{P}(E)$ contiendra 2^n éléments. Prouvons séparément l'injectivité et la surjectivité :

- prouver l'injectivité revient, par contraposée, à prouver que $A \neq B \Rightarrow 1_A \neq 1_B$. Dire que $A \neq B$ revient à dire qu'il existe (au moins) un élément x appartenant à A mais pas à B (ou le contraire, les conséquences étant complètement symétriques). On a alors, par définition de la fonction caractéristique, $1_A(x) = 1$ et $1_B(x) = 0$, ce qui implique certainement que $1_A \neq 1_B$.
- choisissons maintenant une application f de E dans $\{0, 1\}$ quelconque et notons $A = f^{-1}(\{1\})$. Par définition, on a donc, $\forall x \in A, f(x) = 1$, et $\forall x \notin A, f(x) \neq 1$, c'est-à-dire $f(x) = 0$ puisque f ne prend que les valeurs 0 et 1. Autrement dit, $f = 1_A$, et f admet donc l'ensemble A comme antécédent par l'application φ , qui est bien surjective.

□

3.3 Relations d'équivalence.

Nous avons déjà étudié il y a quelques semaines les relations d'ordre sur un ensemble, dont le but est en gros de classer les éléments d'un ensemble. Les relations d'équivalence procèdent sur le même principe, mais en tentant de regrouper les éléments en paquets d'éléments « qui se ressemblent ».

Définition 18. Une **relation d'équivalence** R sur un ensemble E est :

- réflexive : $\forall x \in E, xRx$
- symétrique : $\forall (x, y) \in E^2, xRy \Leftrightarrow yRx$.
- transitive : $\forall (x, y, z) \in E^3, xRy$ et yRz impliquent xRz .

Exemple : La relation d'équivalence la plus simple possible sur un ensemble quelconque E est la relation d'égalité (mais elle ne présente aucun intérêt !). Quelques exemples moins triviaux :

- la relation de congruence modulo 2π sur $\mathbb{R}$ (qui consiste à identifier les réels correspondant au même point sur le cercle trigonométrique)
- la relation de parallélisme sur l'ensemble de toutes les droites du plan (en considérant deux droites confondues comme parallèles).

Définition 19. Soit $x \in E$ et R une relation d'équivalence sur E , la **classe d'équivalence** de x pour la relation R est l'ensemble $C_x = \{y \in E \mid xRy\}$.

Définition 20. Soit E un ensemble, une liste de sous-ensembles $(A_1, A_2, \dots, A_n)$ de E forme une **partition** de E si :

- $\forall (i, j) \in \{1, n\}^2, i \neq j \Rightarrow A_i \cap A_j = \emptyset$.
- $\bigcup_{i=1}^n A_i = E$.

Remarque 17. Autrement dit, tout élément $x \in E$ appartient à exactement un des sous-ensembles A_i . Cette notion de partition est essentielle en probabilités, comme nous le verrons en fin d'année.

Proposition 10. Si R est une relation sur E , l'ensemble des classes d'équivalence de R forme une partition de l'ensemble E .

Démonstration. En effet, supposons que deux classes C_x et C_y ne soient pas disjointes, il existe donc un z tel que xRz et zRy , ce qui implique par transitivité que xRy . Les classes d'équivalence de x et de y sont alors les mêmes. Par ailleurs, tout élément de E appartient certainement à une classe d'équivalence : la sienne ! □

Exemples : Pour la relation de parallélisme sur l'ensemble des droites du plan, la classe d'équivalence d'une droite donnée est constituée de toutes les droites du plan qui lui sont parallèles. On l'appelle **direction** de la droite.

La relation définie sur $\mathbb{C}$ par xRy si $|x| = |y|$ est une relation d'équivalence, et ses classes d'équivalence sont, dans le plan complexe, tous les cercles concentriques centrés en l'origine du repère.