

# *Géométrie I*

*David Cimasoni*

# Table des matières

|            |                                                               |           |
|------------|---------------------------------------------------------------|-----------|
| <b>I</b>   | <b>Groupes et actions de groupes</b>                          | <b>5</b>  |
| I.1        | Groupes et sous-groupes . . . . .                             | 5         |
| I.1.1      | Groupes . . . . .                                             | 5         |
| I.1.2      | Sous-groupes . . . . .                                        | 11        |
| I.2        | Homomorphismes de groupes . . . . .                           | 12        |
| I.2.1      | Homomorphismes . . . . .                                      | 13        |
| I.2.2      | Isomorphismes . . . . .                                       | 14        |
| I.3        | Actions de groupes . . . . .                                  | 16        |
| I.3.1      | Actions de groupes, définition et exemples . . . . .          | 16        |
| I.3.2      | Relations d'équivalence et orbites . . . . .                  | 17        |
| I.3.3      | Le stabilisateur . . . . .                                    | 19        |
| I.3.4      | La formule des orbites . . . . .                              | 20        |
| I.3.5      | La formule de Burnside . . . . .                              | 22        |
| <b>II</b>  | <b>Groupes d'isométries</b>                                   | <b>25</b> |
| II.1       | Distances et isométries . . . . .                             | 25        |
| II.1.1     | Espaces métriques . . . . .                                   | 25        |
| II.1.2     | Le groupe d'isométries d'un espace métrique . . . . .         | 27        |
| II.2       | Le groupe des isométries de l'espace $\mathbb{R}^n$ . . . . . | 28        |
| II.2.1     | Le produit semi-direct . . . . .                              | 29        |
| II.2.2     | Applications affines . . . . .                                | 30        |
| II.2.3     | Le groupe $\text{Isom}(\mathbb{R}^n)$ . . . . .               | 32        |
| II.3       | Classification des isométries . . . . .                       | 34        |
| II.3.1     | Le cas de la droite . . . . .                                 | 34        |
| II.3.2     | Classification des isométries du plan . . . . .               | 35        |
| II.3.3     | Le cas général de l'espace $\mathbb{R}^n$ . . . . .           | 39        |
| II.4       | Groupes de symétries . . . . .                                | 42        |
| II.4.1     | Calcul de groupes de symétries . . . . .                      | 42        |
| II.4.2     | Sous-groupes finis de $\text{Isom}(\mathbb{R}^n)$ . . . . .   | 47        |
| <b>III</b> | <b>Géométrie hyperbolique</b>                                 | <b>53</b> |
| III.1      | Inversions . . . . .                                          | 53        |
| III.1.1    | Définition et premières propriétés . . . . .                  | 53        |
| III.1.2    | Une application : le porisme de Steiner . . . . .             | 58        |
| III.1.3    | Projection stéréographique . . . . .                          | 59        |

|                                                        |    |
|--------------------------------------------------------|----|
| III.2 Transformations de Möbius et birapport . . . . . | 60 |
| III.2.1 Le groupe de Möbius . . . . .                  | 61 |
| III.2.2 Le birapport . . . . .                         | 64 |
| III.3 Le disque de Poincaré . . . . .                  | 66 |
| III.3.1 La distance hyperbolique . . . . .             | 66 |
| III.3.2 Le disque de Poincaré . . . . .                | 71 |
| III.4 Isométries hyperboliques . . . . .               | 72 |
| III.4.1 Le demi-plan de Poincaré . . . . .             | 72 |
| III.4.2 Le groupe $\text{Isom}(\mathbb{H})$ . . . . .  | 73 |

# Introduction

En très résumé, l'essentiel de la première partie du cours a été consacrée à l'étude de la *géométrie euclidienne* dans le plan, d'abord au moyen de méthodes classiques (chapitre I), puis au moyen de méthodes analytiques (chapitre II). Tentons de prendre un peu de recul pour discerner les concepts principaux de cette théorie.

Clairement, le concept le plus fondamental est celui de *distance*. En effet, les *notions géométriques* étudiées en géométrie euclidienne sont précisément les notions laissées invariantes par les transformations qui préservent les distances – ce qu'on appelle les *isométries*. Par exemple, les propriétés telles que “être une droite”, “être un cercle” ou “être deux points à distance  $x$  l'un de l'autre” sont des propriétés géométriques : toute isométrie envoie une droite sur une droite, un cercle sur un cercle, et deux points à distance  $x$  sur deux points à distance  $x$ . En revanche, “être une droite verticale” n'est pas une propriété géométrique, puisqu'elle n'est pas préservée par les isométries ; “être une droite verticale” n'a, pour ainsi dire, pas de sens en géométrie euclidienne.

Notons encore que l'ensemble  $G$  des isométries du plan  $X = \mathbb{R}^2$  n'est pas simplement un ensemble, mais est muni d'une structure supplémentaire. En effet, on peut toujours composer deux isométries, et le résultat sera à son tour une isométrie. De plus, l'application identité est clairement une isométrie. Finalement, l'inverse d'une isométrie est encore une isométrie. Techniquement, on dit que  $G$  est un *groupe*, qui *agit* sur l'ensemble  $X$ . (Voir le chapitre I ci-dessous.)

En résumé :

On peut comprendre la géométrie euclidienne comme l'étude des propriétés d'objets de l'ensemble  $X = \mathbb{R}^2$  (et plus généralement, de  $X = \mathbb{R}^n$ ) invariantes par l'action du groupe  $G$  des isométries de  $X$ .

Par la suite, nous avons brièvement traité de *géométrie projective* (chapitre III). Dans cette théorie, les notions fondamentales sont celles d'*alignement* (de points) et de *concourance* (de droites). Les transformations qui préservent ces notions sont ce qu'on appelle les *transformations projectives* du plan projectif  $\mathbb{P}^2$ , et plus généralement de l'espace projectif  $\mathbb{P}^n$ . Ces transformations forment à leur tour un groupe, considérablement plus grand que celui des isométries euclidiennes. En géométrie projective, les notions géométriques sont les propriétés invariantes par ce groupe de transformations. Et comme ce groupe est énorme, le nombre de notions géométriques sera plus restreint. Par exemple, les propriétés telles que “être

une droite”, “être trois points alignés” ou “être quatre points alignés de birapport  $x$  fixé” sont invariantes : ce sont des notions de géométrie projective. En revanche, les concepts tels que “distance”, “cercle” et “angle” ne sont pas invariants : ils n’ont pas de sens en géométrie projective.

De plus, puisqu’un énoncé de géométrie projective ne contient que des notions invariantes par le groupe des transformations projectives, il est parfois possible d’utiliser ces transformations pour ramener un tel énoncé donné en un énoncé plus simple. C’est le *principe de Poncelet*, que l’on a vu en première partie du cours, puis utilisé pour donner des démonstrations très simples des théorèmes de Pappus et de Desargues.

On peut résumer cette discussion par l’affirmation suivante, à la base du *programme d’Erlangen* que le jeune FELIX KLEIN développa en 1872 :

Une *géométrie*, ce n’est rien d’autre qu’un groupe  $G$  agissant sur un ensemble  $X$  ; on étudie les propriétés invariantes par cette action.

Par exemple, en géométrie euclidienne de dimension 2, on étudie les propriétés invariantes par l’action du groupe d’isométries  $G = \text{Isom}(\mathbb{R}^2)$  sur le plan  $X = \mathbb{R}^2$ . En géométrie projective, il s’agit de l’action du groupe des transformations projectives  $\text{PGL}(3, \mathbb{R})$  sur le plan projectif  $\mathbb{P}^2$ . D’autres exemples naturels en dimension 2 sont fournis par la *géométrie hyperbolique*, qui traite de l’action du groupe  $\text{PSL}(2, \mathbb{R})$  sur le plan hyperbolique  $\mathbb{H}^2$ , et la *géométrie sphérique*, où l’on s’intéresse à l’action du groupe orthogonal  $\text{O}(3)$  sur la sphère  $\mathbb{S}^2$ . Chacune de ces géométries a bien entendu son intérêt et son utilité propre.

Notons pour finir qu’une grande partie du cours de Géométrie II sera dédié à l’étude d’une autre géométrie, appelée la *topologie*. Dans cette théorie, l’ensemble  $X$  est n’importe quel ensemble sur lequel on peut définir la notion d’application continue (ce qu’on appelle un *espace topologique*), et le groupe  $G$  est formé de toutes les applications bijectives continues d’inverse continue (ce qu’on appelle un *homéomorphisme* de  $X$ ). La topologie consiste en l’étude des propriétés des espaces topologiques invariantes par homéomorphisme : connexité, compacité, et autres. Mais ceci est une autre histoire.

La seconde partie du cours de Géométrie I aura le plan suivant, que l’on espère avoir suffisamment motivé par la discussion ci-dessus. Dans un premier chapitre, nous traiterons de la théorie des groupes et des actions de groupes. Le deuxième chapitre sera consacré à l’étude du groupe des isométries euclidiennes du plan  $\mathbb{R}^2$ , et plus généralement de l’espace  $\mathbb{R}^n$ . Au chapitre trois, on donnera une introduction à la géométrie hyperbolique en dimension 2, traitée avec un minimum de méthodes analytiques.

# Chapitre I: Groupes et actions de groupes

On va réunir dans ce chapitre toutes les considérations abstraites nécessaires au développement de la théorie. Les sections I.1 et I.2 consistent en une introduction standard aux concepts de groupes et d'homomorphismes de groupes. Ils seront traités de manière quelque peu aride, mais plus de détails sont donnés au cours d'Algèbre I. La section I.3 est consacrée aux actions de groupes, notion fondamentale en géométrie comme nous l'avons vu en introduction.

## I.1 Groupes et sous-groupes

### I.1.1 Groupes

Commençons par la définition d'une structure omniprésente en mathématiques, celle de groupe.

#### Définition I.1

Soit  $G$  un ensemble muni d'une *loi de composition*, c'est-à-dire d'une application

$$G \times G \longrightarrow G, \quad (g, h) \longmapsto g \star h.$$

On dit que  $G$  est un **groupe** s'il satisfait les trois axiomes suivants :

(G1) *Associativité* :  $g \star (h \star k) = (g \star h) \star k$  pour tous  $g, h, k \in G$ .

(G2) *Existence d'un élément neutre* : il existe  $e \in G$  tel que  $e \star g = g \star e = g$  pour tout  $g \in G$ .

(G3) *Existence d'un inverse* : pour tout  $g \in G$ , il existe  $g' \in G$  tel que  $g \star g' = g' \star g = e$ .

Si de plus cette loi de composition satisfait l'égalité  $g \star h = h \star g$  pour tous  $g, h \in G$ , alors  $G$  est dit **commutatif**, ou **abélien**.

Le cardinal de  $G$ , noté  $|G|$ , est appelé son **ordre**.

Cette définition est bien entendu très abstraite. La meilleure manière de la "comprendre" est de se convaincre que de tels objets se retrouvent un peu partout en mathématiques, d'où l'utilité de les étudier de manière abstraite.

Mais avant de donner une liste de tels exemples, quelques remarques s'imposent.

**Remarques et notations.**

1. Formellement, un groupe est la donnée d'un couple  $(G, \star)$ , mais on le note simplement  $G$  lorsqu'il n'y a pas de confusion possible.

On adoptera presque systématiquement la notation  $gh$  au lieu de  $g \star h$  : c'est ce qu'on appelle la *notation multiplicative*.

Dans le cas abélien, on utilise parfois la *notation additive*  $g + h$ .

2. L'associativité signifie que lorsque l'on calcule dans un groupe, on peut "laisser tomber les parenthèses". Par exemple, étant donnés quatre éléments  $g_1, \dots, g_4$  d'un groupe  $G$ , l'associativité implique les identités

$$((g_1 g_2) g_3) g_4 = (g_1 (g_2 g_3)) g_4 = g_1 ((g_2 g_3) g_4) = g_1 (g_2 (g_3 g_4)) = (g_1 g_2) (g_3 g_4),$$

et l'on notera donc simplement cet élément par  $g_1 g_2 g_3 g_4$ .

Néanmoins, on conservera souvent des parenthèses, mais à des fins presque purement pédagogiques.

3. L'élément neutre est nécessairement unique.

En effet, si  $e_1$  et  $e_2$  sont deux éléments neutres dans un groupe  $G$ , alors

$$e_1 = e_1 \star e_2 = e_2,$$

où la première (resp. seconde) égalité découle du fait que  $e_2$  (resp.  $e_1$ ) est un élément neutre. J

Cet unique élément neutre est habituellement noté  $e$ , ou  $e_G$  en cas de confusion possible. En notation multiplicative, on le note  $1$  ou  $1_G$ , d'où les égalités  $1g = g1 = g$  pour tout  $g \in G$ . Dans le cas abélien, en notation additive, on l'écrit  $0$  ou  $0_G$ , d'où les égalités  $0+g = g+0 = g$  pour tout  $g \in G$ . Remarquons que l'ensemble vide n'est *pas* un groupe, puisqu'il ne contient pas de neutre.

4. L'inverse d'un élément  $g \in G$  fixé est unique.

En effet, si  $g', g'' \in G$  sont deux inverses de  $g \in G$ , alors on a les égalités suivantes dans  $G$  :

$$g' \stackrel{(G2)}{=} g' e \stackrel{(G3)}{=} g' (g g'') \stackrel{(G1)}{=} (g' g) g'' \stackrel{(G3)}{=} e g'' \stackrel{(G2)}{=} g'',$$

d'où l'identité  $g' = g''$ . J

En notation multiplicative, l'unique inverse de  $g \in G$  s'écrit  $g^{-1}$ , d'où les égalités  $g g^{-1} = g^{-1} g = 1$ . Dans le cas abélien, en notation additive, on l'écrit  $-g$ ; on écrit aussi  $g + (-h) =: g - h$ , d'où l'égalité  $g - g = 0$  pour tout  $g \in G$ .

5. Les règles de calcul suivantes se déduisent des définitions :

- Pour tous  $x, g, h \in G$ , si  $xg = xh$  ou  $gx = hx$ , alors  $g = h$ .

- Pour tous  $g, h \in G$ ,  $(g^{-1})^{-1} = g$  et  $(gh)^{-1} = h^{-1}g^{-1}$ .

Dans le cas abélien, en notation additive, elles se traduisent comme suit :

- Pour tous  $x, g, h \in G$ , si  $x + g = x + h$  alors  $g = h$ .
- Pour tous  $g, h \in G$ ,  $-(-g) = g$  et  $-(g + h) = -g - h$ .

6. Étant donné un élément  $g \in G$  et un entier  $n \in \mathbb{Z}$ , la  $n^{\text{ième}}$  puissance de  $g$  est l'élément  $g^n \in G$  défini comme suit (en notation multiplicative) :

$$g^n = \begin{cases} g \cdots g & (n \text{ fois}) & \text{si } n > 0, \\ e & & \text{si } n = 0, \\ g^{-1} \cdots g^{-1} & (|n| \text{ fois}) & \text{si } n < 0. \end{cases}$$

On montre facilement la troisième règle de calcul suivante :

- Pour tous  $g \in G$  et  $n, m \in \mathbb{Z}$ ,  $g^n g^m = g^{n+m}$  et  $(g^n)^m = g^{nm}$ .

Dans le cas abélien, en notation additive, la  $n^{\text{ième}}$  puissance de  $g$  se note  $n \cdot g$ . Par définition, on a donc :

$$n \cdot g = \begin{cases} g + \cdots + g & (n \text{ fois}) & \text{si } n > 0, \\ e & & \text{si } n = 0, \\ (-g) + \cdots + (-g) & (|n| \text{ fois}) & \text{si } n < 0, \end{cases}$$

et la troisième règle de calcul se traduit comme suit :

- Pour tous  $g \in G$  et  $n, m \in \mathbb{Z}$ ,  $n \cdot g + m \cdot g = (n+m) \cdot g$  et  $m \cdot (n \cdot g) = (mn) \cdot g$ .

7. Lorsqu'on veut démontrer que  $(G, \star)$  est un groupe, il faut vérifier les axiomes  $(G1)$ ,  $(G2)$  et  $(G3)$ , mais aussi que la loi  $\star$  est une *loi interne*, c'est à dire qu'elle définit bien une application  $G \times G \rightarrow G$ . En d'autres termes, il faut montrer

$$g, h \in G \Rightarrow g \star h \in G.$$

Par exemple, l'ensemble  $G = \{-1, 0, 1\}$  muni de l'addition satisfait les trois axiomes, mais n'est néanmoins *pas* un groupe !

Voici à présent une longue liste de familles d'exemples.

### Exemples (de groupes).

1. *Ensembles de nombres.*

L'ensemble

$$\mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\}$$

des nombres entiers muni de l'addition est un groupe abélien d'ordre infini. De même, les ensembles  $\mathbb{Q}$ ,  $\mathbb{R}$  et  $\mathbb{C}$  munis de l'addition sont des groupes

abéliens d'ordre infini. C'est à ces exemples qu'on doit la notation additive dans un groupe abélien.

Par contre, l'ensemble  $\mathbb{N} = \{0, 1, 2, \dots\}$  des entiers naturels muni de l'addition n'est *pas* un groupe : la loi est bien interne, il satisfait bien les axiomes (G1) et (G2), mais pas (G3), puisqu'aucun élément de  $\mathbb{N}$  n'a d'inverse à l'exception du neutre.

L'ensemble  $\mathbb{Q}^* := \mathbb{Q} \setminus \{0\}$  muni de la multiplication est un groupe abélien d'ordre infini. De même, les ensembles  $\mathbb{R}^* := \mathbb{R} \setminus \{0\}$  et  $\mathbb{C}^* := \mathbb{C} \setminus \{0\}$  sont des groupes abéliens (infinis) pour la multiplication, ainsi que  $\mathbb{R}_+^* := (0, \infty)$ .

## 2. Espaces vectoriels.

Si  $(E, +, \cdot)$  est un espace vectoriel sur  $K = \mathbb{R}$  ou  $\mathbb{C}$ , alors  $(E, +)$  est un groupe abélien : c'est exactement la signification des premiers quatre axiomes pour un espace vectoriel.

## 3. Groupes de matrices.

Pour  $n \geq 1$  fixé et  $K = \mathbb{R}$  ou  $\mathbb{C}$ , considérons l'ensemble

$$\mathrm{GL}(n, K) := \{M \in \mathcal{M}_n(K) \mid \det(M) \neq 0\}$$

des matrices carrées inversibles de taille  $n$  à coefficients dans  $K$ . Muni de la multiplication matricielle, il forme un groupe appelé le *groupe général linéaire* de degré  $n$  de  $K$ .

Notons que pour  $n = 1$ , on retrouve  $\mathrm{GL}(1, K) = K^* = K \setminus \{0\}$ , qui est abélien. Le groupe  $\mathrm{GL}(n, K)$  n'est en revanche jamais abélien pour  $n \geq 2$ . Similairement, on vérifie facilement que l'ensemble de matrices

$$\mathrm{SL}(n, K) := \{M \in \mathcal{M}_n(K) \mid \det(M) = 1\}$$

est un groupe pour la multiplication matricielle : il s'agit du *groupe spécial linéaire* de degré  $n$  de  $K$ . Pour  $n = 1$ , on trouve  $\mathrm{SL}(1, K) = \{(1)\}$ , un groupe d'ordre 1. Il s'agit (d'une des "incarnations") du *groupe trivial*.

Pour  $n \geq 1$  fixé, l'ensemble

$$\mathrm{O}(n) := \{Q \in \mathcal{M}_n(\mathbb{R}) \mid Q^T Q = I\}$$

des matrices orthogonales de taille  $n$  est un groupe pour la multiplication matricielle ; on parle du *groupe orthogonal* de degré  $n$ . Le théorème principal de la section II.4 (B) en première partie du cours se traduit alors comme suit : le choix d'une base orthonormale de  $\mathbb{R}^n$  permet d'identifier les isométries linéaires de  $\mathbb{R}^n$  avec le groupe  $\mathrm{O}(n)$ .

Finalement, l'ensemble

$$\mathrm{SO}(n) := \mathrm{SL}(n, \mathbb{R}) \cap \mathrm{O}(n)$$

forme un groupe, appelé le *groupe spécial orthogonal* de degré  $n$ . Comme ci-dessus, le choix d'une base orthonormale de  $\mathbb{R}^n$  permet d'identifier les isométries linéaires de  $\mathbb{R}^n$  qui *préservent l'orientation* avec le groupe  $\mathrm{SO}(n)$ . Par des résultats vus en première partie du cours, le groupe  $\mathrm{SO}(2)$  n'est autre que le groupe des matrices de rotation.

4. *Cercle.*

Considérons l'ensemble

$$\mathbb{S}^1 := \{z \in \mathbb{C} \mid \|z\| = 1\}$$

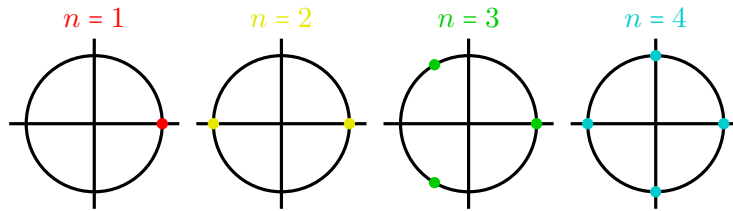
des nombres complexes de norme 1. Muni de la multiplication complexe, il forme un groupe abélien d'ordre infini.

5. *Racines de l'unité.*

Soit  $n \geq 1$  un entier ; considérons l'ensemble

$$\mu_n(\mathbb{C}) := \{z \in \mathbb{C} \mid z^n = 1\}$$

des racines  $n^{\text{ièmes}}$  de l'unité dans  $\mathbb{C}$ , illustré ci-dessous pour  $n = 1, 2, 3, 4$ .



Muni de la multiplication complexe, on vérifie que c'est un groupe abélien d'ordre  $n$ . Par exemple  $\mu_1(\mathbb{C}) = \{1\}$  est une nouvelle incarnation du groupe trivial, tandis que  $\mu_2(\mathbb{C}) = \{-1, 1\}$  est un groupe d'ordre 2.

6. *Groupes symétriques.*

Soit  $X$  un ensemble non-vidé. Alors, l'ensemble

$$S(X) := \{f: X \rightarrow X \mid f \text{ est bijective}\}$$

muni de la composition des applications est un groupe : cela découle de faits élémentaires vus en LOGIQUE ET THÉORIE DES ENSEMBLES. On notera  $\text{id}$  ou  $\text{id}_X$  son neutre, qui est l'application identité. Ce groupe est appelé le *groupe symétrique sur  $X$* .

Si  $X$  est un ensemble fini et que l'on numérote ses éléments  $\{1, 2, \dots, n\}$ , alors son groupe symétrique est noté  $S_n$ . C'est le *groupe des permutations* de  $n$  objets, d'ordre  $n! = n \cdot (n-1) \cdots 2 \cdot 1$ .

Par exemple, on a  $S_1 = \{\text{id}\}$ , une nouvelle incarnation du groupe trivial, et  $S_2$  un groupe abélien d'ordre 2, tandis que  $S_n$  n'est jamais abélien pour  $n > 2$ .

7. *Entiers modulo  $n$ .*

Soit  $n \geq 1$  un entier ; considérons l'ensemble

$$\mathbb{Z}_n := \{[0], [1], \dots, [n-1]\}$$

muni de la loi de composition  $[a] + [b] := [\ell]$ , où  $\ell \in \{0, 1, \dots, n-1\}$  est le reste de la division de  $a + b \in \mathbb{Z}$  par  $n$ . C'est un groupe abélien d'ordre  $n$ , avec élément neutre  $[0]$  et inverse  $-[a] = [n-a]$  pour tout  $a = 0, 1, \dots, n-1$ . Par exemple, la loi de composition sur  $\mathbb{Z}_2 = \{[0], [1]\}$  est donnée par :  $[0] + [0] = [1] + [1] = [0]$ ,  $[0] + [1] = [1] + [0] = [1]$ .

#### 8. *Produit direct.*

Soient  $(G_1, *)$  et  $(G_2, \bullet)$  deux groupes. Alors, on vérifie que le produit cartésien  $G_1 \times G_2$  est un groupe pour la loi de composition suivante : pour  $g_1, h_1 \in G_1$  et  $g_2, h_2 \in G_2$ , on pose

$$(g_1, g_2)(h_1, h_2) = (g_1 * h_1, g_2 \bullet h_2) \in G_1 \times G_2.$$

On parle du *produit direct* de  $G_1$  et  $G_2$ .

Plus généralement, si  $G_1, \dots, G_n$  sont des groupes, alors le produit cartésien  $G_1 \times \dots \times G_n$  est un groupe pour la loi de composition définie composante par composante.

Par exemple, le groupe abélien  $(\mathbb{R}^n, +)$  n'est autre que le produit direct de  $n$  copies de  $(\mathbb{R}, +)$ .

#### 9. *Les groupes diédraux*

Pour  $n \geq 2$  un entier fixé, notons  $P_n$  un  $n$ -gone régulier (si  $n = 2$ , il s'agit d'un segment de droite). Pour fixer les idées, disons que  $P_n$  est le polygone avec sommets en l'ensemble  $\mu_n(\mathbb{C})$  des racines  $n^{\text{ièmes}}$  de l'unité dans  $\mathbb{C}$ . Considérons l'ensemble  $D_{2n}$  de toutes les *symétries* de  $P_n$ , c'est-à-dire les isométries du plan  $\mathbb{R}^2 = \mathbb{C}$  qui préservent globalement  $P_n$  :

$$D_{2n} = \{f: \mathbb{R}^2 \rightarrow \mathbb{R}^2 \mid f \text{ est une isométrie et } f(P_n) = P_n\}.$$

Muni de la composition des isométries,  $D_{2n}$  est un groupe appelé le *groupe diédral*.

Nous allons maintenant vérifier qu'il est d'ordre  $2n$ , d'où la notation. (C'est un premier exemple de détermination d'un groupe de symétrie ; de nombreux autres exemples seront donnés en section II.4.) Plus précisément, soient  $r \in D_{2n}$  l'isométrie donnée par la rotation d'angle  $2\pi/n$  dans le sens trigonométrique autour de l'origine, et  $s \in D_{2n}$  l'isométrie donnée par la réflexion d'axe horizontal, i.e la conjugaison complexe. Nous allons vérifier que tout élément  $\alpha$  de  $D_{2n}$  s'écrit de manière unique comme  $\alpha = r^\ell s^\varepsilon$  avec  $\ell = 0, 1, \dots, n-1$  et  $\varepsilon = 0, 1$ .

<sup>r</sup> En effet, on verra en section II.4.1 qu'une telle isométrie doit forcément fixer l'origine. Par un résultat de la première partie du cours, c'est donc une isométrie *linéaire*. Par l'exemple 3 ci-dessus, le choix d'une base orthonormale de  $\mathbb{R}^2$  (disons,  $(1, i)$ ) permet d'identifier ces isométries avec le groupe  $O(2)$ . Ainsi, on s'intéresse à l'ensemble

$$D_{2n} = \{\alpha \in O(2) \mid \alpha(P_n) = P_n\}.$$

Comme nous l'avons rappelé en exemple 3 ci-dessus,  $\text{SO}(2)$  est l'ensemble de matrices de rotations. Ainsi, l'intersection  $D_{2n} \cap \text{SO}(2)$  est l'ensemble des rotations du plan autour de l'origine qui fixent  $P_n$ , donc les rotations d'angle un multiple de  $2\pi/n$ . En clair,

$$D_{2n} \cap \text{SO}(2) = \{r^\ell \mid \ell = 0, 1, \dots, n-1\},$$

soit  $n$  éléments distincts. D'autre part, considérons l'application de  $D_{2n} \cap \text{SO}(2)$  dans  $\{\alpha \in D_{2n} \mid \alpha \notin \text{SO}(2)\}$  qui associe à  $\rho$  sa composition  $\rho s$  avec  $s$ . Il s'agit clairement d'une bijection (d'inverse  $\alpha \mapsto \alpha s$ , puisque  $s^2 = 1$ ). Ainsi,

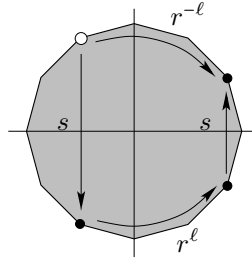
$$\{\alpha \in D_{2n} \mid \alpha \notin \text{SO}(2)\} = \{r^\ell s \mid \ell = 0, 1, \dots, n-1\},$$

à nouveau  $n$  éléments distincts. Comme  $D_{2n}$  est l'union disjointe de cet ensemble et de  $D_{2n} \cap \text{SO}(2)$ , l'affirmation est démontrée.  $\square$

La façon dont se composent les éléments de ce groupe est entièrement déterminée par les égalités suivantes, aussi appelées *relations* :

$$r^n = 1, \quad s^2 = 1 \quad \text{et} \quad sr^\ell s = r^{-\ell} \quad \text{pour tout } \ell.$$

Si les deux premières relations sont claires, la troisième est moins évidente et se comprend à l'aide de la figure suivante. (Bien-sûr, une preuve formelle peut se faire à l'aide des matrices associées à  $r$  et  $s$ .)



Il s'ensuit que le groupe diédral d'ordre  $2n$  n'est jamais abélien si  $n \geq 3$ . En effet, les relations ci-dessus impliquent les égalités  $sr = r^{-1}s = r^{n-1}s$  qui n'est pas égal à  $rs$  si  $n \geq 3$ .

### I.1.2 Sous-groupes

Passons à présent à la seconde définition de cette section.

#### Définition I.2

Soit  $G$  un groupe, et  $H \subset G$  un sous-ensemble de  $G$ . On dit que  $H$  est un **sous-groupe** de  $G$ , noté  $H < G$ , si  $H$  est un groupe pour la restriction de la loi de composition de  $G$  à  $H$ .

De manière plus concrète, cela signifie que  $H$  satisfait les trois propriétés suivantes.

1. Pour tous  $h_1, h_2$  dans  $H$ ,  $h_1 h_2$  appartient aussi à  $H$ .

2. L'élément neutre  $e_G$  appartient à  $H$ .
3. Pour tout élément  $h \in H$ , son inverse  $h^{-1}$  appartient aussi à  $H$ .

Un exercice amusant consiste à vérifier que, dans le cas d'un sous-ensemble  $H$  non-vidé, ces trois propriétés sont équivalentes à une unique condition. Comme il s'agit d'un des premiers exercices de ce type, il n'est sans doute pas inutile de le résoudre ici.

**Lemme I.1.** *Soit  $G$  un groupe. Un sous-ensemble non-vidé  $H \subset G$  est un sous-groupe de  $G$  si et seulement si, pour tous  $h_1, h_2$  dans  $H$ ,  $h_1 h_2^{-1}$  appartient aussi à  $H$ .*

*Démonstration.* Les trois propriétés ci-dessus impliquent trivialement la condition de l'énoncé ; vérifions la réciproque. Soit donc  $H \subset G$  non-vidé satisfaisant la condition de l'énoncé, que l'on notera par  $(C)$ . Comme  $H$  est non-vidé, il existe  $h \in H$ . Par la condition  $(C)$  appliquée à  $h_1 = h_2 = h \in H$ ,  $h h^{-1} = e$  appartient à  $H$ , et la seconde propriété est vérifiée. Soit à présent  $h \in H$  quelconque ; par la condition  $(C)$  appliquée à  $h_1 = e$  (élément de  $H$  par le point précédent) et  $h_2 = h \in H$ ,  $eh^{-1} = h^{-1}$  appartient aussi à  $H$ , et la troisième propriété est vérifiée. Reste la première propriété. Soient donc  $h_1, h_2 \in H$  ; par la condition  $(C)$  appliquée à  $h_1 \in H$  et  $h_2^{-1}$  (élément de  $H$  par le point précédent),  $h_1(h_2^{-1})^{-1} = h_1 h_2$  appartient aussi à  $H$ , ce qui conclut la démonstration.  $\square$

**Exemples** (de sous-groupes).

1. Pour tout groupe  $G$ , les sous-ensembles  $H = \{e\}$  et  $H = G$  sont trivialement des sous-groupes de  $G$ .
2. Clairement, les entiers forment un sous-groupe des rationnels :  $\mathbb{Z} < \mathbb{Q}$ . De même,  $\mathbb{Q} < \mathbb{R} < \mathbb{C}$ .
3. Le groupe des racines  $n^{\text{ièmes}}$  de l'unité vu en exemple 5 ci-dessus est un sous-groupe du groupe des nombres complexes de norme 1 :  $\mu_n(\mathbb{C}) < \mathbb{S}^1$ .
4. Les groupes de matrices vus en exemple 3 satisfont :

$$\mathrm{SO}(n) < \mathrm{O}(n) < \mathrm{GL}(n, \mathbb{R}) \quad \text{et} \quad \mathrm{SO}(n) < \mathrm{SL}(n, \mathbb{R}) < \mathrm{GL}(n, \mathbb{R}).$$

5. Notons pour finir que le groupe  $C_n := D_{2n} \cap \mathrm{SO}(2)$  consistant en toutes les isométries du plan qui préservent globalement  $P_n$  et qui préservent l'orientation est un sous-groupe (abélien d'ordre  $n$ ) du groupe diédral :  $C_n < D_{2n}$ .

## I.2 Homomorphismes de groupes

Un simple coup d'oeil aux exemples 6 et 7 de groupes suffit à se convaincre que  $S_2$  et  $\mathbb{Z}_2$  sont vraiment deux noms différents pour "le même groupe". Comment traduire cette intuition en un énoncé plus précis ?

Voici un principe assez général en mathématiques<sup>1</sup> : après avoir muni un ensemble d'une structure (par exemple, celle d'espace vectoriel), il est utile de

---

1. Ce principe est formalisé dans ce qu'on appelle le *langage des catégories*.

considérer les applications qui préservent cette structure (dans cet exemple, les applications linéaires). Cela permet en particulier de répondre à la question ci-dessus.

Pour la structure de groupe, ce rôle est joué par ce qu'on appelle les *homomorphismes de groupes*.

### I.2.1 Homomorphismes

#### Définition I.3

Soient  $(G, \star)$  et  $(G', \bullet)$  deux groupes. Une application  $\phi: G \rightarrow G'$  est un **homomorphisme de groupes** (ou plus simplement, un **homomorphisme**) si

$$\phi(g \star h) = \phi(g) \bullet \phi(h)$$

pour tous  $g, h \in G$ .

En clair, un homomorphisme est une application entre deux groupes qui est compatible avec les lois de composition. Il peut sembler surprenant qu'on n'exige ni compatibilité avec l'élément neutre

$$\phi(e_G) = e_{G'}, \quad (\text{I.2.1})$$

ni compatibilité avec l'inverse

$$\phi(g^{-1}) = \phi(g)^{-1} \text{ pour tout } g \in G. \quad (\text{I.2.2})$$

En réalité, ces deux propriétés sont toujours automatiquement satisfaites par un homomorphisme de groupes : cela découle facilement de la définition.

Les homomorphismes satisfont deux autres propriétés immédiates mais extrêmement importantes. Tout d'abord, l'application identité  $\text{id}_G: G \rightarrow G$  est un homomorphisme. Ensuite, si  $\phi: G \rightarrow G'$  et  $\psi: G' \rightarrow G''$  sont des homomorphismes, alors la composition  $\psi \circ \phi: G \rightarrow G''$  est à son tour un homomorphisme.<sup>2</sup>

Introduisons encore un peu de terminologie. Étant donné un homomorphisme de groupes  $\phi: G \rightarrow G'$ , son **noyau** est le sous-ensemble de  $G$  défini par

$$\ker(\phi) := \{g \in G \mid \phi(g) = e\}.$$

On vérifie que  $\ker(\phi)$  est un sous-groupe de  $G$ .

**Exemples** (d'homomorphismes de groupes).

1. Pour tous groupes  $G, G'$  l'application  $G \rightarrow G', g \mapsto e$  est un homomorphisme de groupes. On parle d'*homomorphisme trivial*. Son noyau est le groupe  $G$  tout entier.

---

2. Ces propriétés sont importantes du point de vue du langage des catégories mentionné plus haut : elles assurent en effet que les groupes et homomorphismes de groupes forment, précisément, une *catégorie*.

2. Si  $H$  est un sous-groupe d'un groupe  $G$ , alors l'inclusion  $i: H \hookrightarrow G$  est un homomorphisme. Son noyau est trivial :  $\ker(i) = \{e\}$ .
3. Le déterminant définit un homomorphisme de groupes

$$\det: \mathrm{GL}(n, \mathbb{R}) \longrightarrow \mathbb{R}^*.$$

Cela découle de l'égalité bien connue  $\det(AB) = \det(A)\det(B)$ . Par définition, son noyau est égal au groupe spécial linéaire  $\mathrm{SL}(n, \mathbb{R})$ .

4. La signature d'une permutation définit un homomorphisme

$$\mathrm{sgn}: S_n \longrightarrow \{\pm 1\} = \mu_2(\mathbb{C}),$$

puisque  $\mathrm{sgn}(\sigma\tau) = \mathrm{sgn}(\sigma)\mathrm{sgn}(\tau)$  pour toutes permutations  $\sigma, \tau \in S_n$ .

Son noyau  $A_n := \ker(\mathrm{sgn})$  est appelé le *groupe alterné* de degré  $n$ .

5. Comme  $\exp(t+s) = \exp(t)\exp(s)$ , l'application

$$\mathbb{R} \longrightarrow \mathbb{S}^1, \quad t \mapsto \exp(2\pi it)$$

est un homomorphisme. Son noyau est le sous-groupe  $\mathbb{Z}$  de  $\mathbb{R}$ .

Un des intérêts de la définition du noyau réside dans l'énoncé suivant : pour vérifier qu'un homomorphisme est injectif, il suffit de vérifier que son noyau est trivial. Plus précisément :

**Proposition I.2.** *Un homomorphisme  $\phi: G \rightarrow G'$  est injectif si et seulement si  $\ker(\phi) = \{e_G\}$ .*

*Démonstration.* Supposons  $\phi$  injectif, et soit  $g \in \ker(\phi)$ . Par définition et l'égalité (I.2.1), on a alors  $\phi(g) = e_{G'} = \phi(e_G)$ . Par injectivité de  $\phi$ ,  $g$  est égal à  $e_G$  et le noyau est donc trivial. Supposons réciproquement que  $\ker(\phi) = \{e_G\}$ , et soient  $g, h \in G$  tels que  $\phi(g) = \phi(h)$ . Comme  $\phi$  est un homomorphisme, et par l'égalité (I.2.2), il suit

$$\phi(gh^{-1}) = \phi(g)\phi(h^{-1}) = \phi(g)\phi(h)^{-1} = e_{G'}.$$

En d'autres termes,  $gh^{-1}$  est un élément de  $\ker(\phi) = \{e_G\}$ . Ainsi,  $gh^{-1} = e_G$ , d'où  $g = h$ , et  $\phi$  est injective.  $\square$

### I.2.2 Isomorphismes

La réponse à la question du début de la section se trouve dans la définition suivante.

#### Définition I.4

Un homomorphisme de groupes  $\phi: G \rightarrow G'$  qui est bijectif est appelé un **isomorphisme (de groupes)**. S'il existe un tel isomorphisme entre  $G$  et  $G'$ , on dit que  $G$  et  $G'$  sont **isomorphes**, ce qui est noté  $G \simeq G'$ .

Une définition plus naturelle<sup>3</sup> est la suivante : un homomorphisme  $\phi: G \rightarrow G'$  est un isomorphisme s'il existe un homomorphisme  $\psi: G' \rightarrow G$  tel que  $\phi \circ \psi = \text{id}_{G'}$  et  $\psi \circ \phi = \text{id}_G$ . On vérifie facilement que ces deux définitions sont équivalentes.

**Exemples** (de groupes isomorphes).

1. Tous les groupes à un élément sont évidemment isomorphes entre eux : on parle donc *du* groupe trivial. En particulier,  $S_1$ ,  $\mathbb{Z}_1$ ,  $\mu_1(\mathbb{C})$ ,  $C_1$ ,  $\text{SO}(1)$  et  $\text{SL}(1, \mathbb{R})$  sont isomorphes.
2. Tous les groupes à deux éléments sont isomorphes entre eux. Cela découle du fait que si  $G = \{e, g\}$  est un groupe, alors la loi de composition est entièrement déterminée par les axiomes : on a forcément  $eg = ge = g$  et  $e^2 = g^2 = e$ . On peut donc parler *du* groupe d'ordre deux. En particulier,  $S_2$ ,  $\mathbb{Z}_2$ ,  $\mu_2(\mathbb{C})$  et  $C_2$  sont isomorphes.  
Il en va de même pour tous les groupes d'ordre trois (exercice). En revanche, tous les groupes d'ordre 4 ne sont pas isomorphes (voir ci-dessous).
3. Pour  $n$  fixé, soit  $\phi: \mathbb{Z}_n \rightarrow \mathbb{C}^*$  l'application donnée par  $\phi([a]) = \exp(2\pi ia/n)$ . Notons d'abord que  $\phi([a])^n = 1$  pour tout  $a$  ; ainsi, nous sommes en présence d'une application  $\phi: \mathbb{Z}_n \rightarrow \mu_n(\mathbb{C})$ . De plus,

$$\phi([a]+[b]) = \exp(2\pi i(a+b)/n) = \exp(2\pi ia/n) \exp(2\pi ib/n) = \phi([a])\phi([b]),$$

et  $\phi$  est donc un homomorphisme. Finalement, si  $\phi([a]) = 1$ , alors  $[a] = 0$  ; cela signifie que le noyau de  $\phi$  est trivial, et donc, par la proposition I.2, que  $\phi: \mathbb{Z}_n \rightarrow \mu_n(\mathbb{C})$  est injectif. Comme ces deux groupes ont ordre  $n$ ,  $\phi$  est un homomorphisme bijectif, i.e. un isomorphisme. Les groupes  $\mathbb{Z}_n$  et  $\mu_n(\mathbb{C})$  sont donc isomorphes.

On montre similairement que  $\mathbb{Z}_n$  et  $C_n$  sont isomorphes, d'où :  $\mathbb{Z}_n \simeq \mu_n(\mathbb{C}) \simeq C_n$ . On parle du *groupe cyclique d'ordre  $n$* .

4. Les groupes  $\text{SO}(2)$  et  $\mathbb{S}^1$  sont isomorphes (exercice).
5. Les groupes  $S_3$  et  $D_6$  sont isomorphes (exercice).
6. Les groupes  $D_4$  et  $\mathbb{Z}_2 \times \mathbb{Z}_2$  sont isomorphes (exercice).

Pour démontrer que deux groupes sont isomorphes, il suffit d'expliciter un isomorphisme entre les deux. En revanche, pour démontrer que deux groupes ne sont pas isomorphes, il faut exhiber une propriété de groupes laissée invariante par les isomorphismes, satisfaite par un des groupes mais pas par l'autre. Une première telle propriété est l'ordre d'un groupe : par exemple, les groupes  $\mathbb{Z}_2$  et  $\mathbb{Z}_3$  ne sont pas isomorphes, puisqu'ils ne sont pas de même ordre. Une seconde propriété invariante par isomorphismes est la commutativité : par exemple, les groupes  $\mathbb{Z}_6$  et  $D_6$ , bien que de même ordre, ne sont pas isomorphes, puisque le premier est commutatif mais pas le second.

Finissons avec un dernier exercice : les groupes  $\mathbb{Z}_4$  et  $\mathbb{Z}_2 \times \mathbb{Z}_2$ , bien que tous deux commutatifs d'ordre 4, ne sont pas isomorphes. Pourquoi ?

3. Cette définition est plus naturelle du point de vue des catégories.

## I.3 Actions de groupes

### I.3.1 Actions de groupes, définition et exemples

Comme nous l'avons vu en introduction, le concept suivant est d'une importance capitale en géométrie.

#### Définition I.5

Soient  $G$  un groupe et  $X$  un ensemble non-vidé. Une **action** de  $G$  sur  $X$  est une application

$$G \times X \longrightarrow X, \quad (g, x) \longmapsto g \cdot x$$

qui satisfait les deux propriétés suivantes.

- (i)  $e \cdot x = x$  pour tout  $x \in X$ .
- (ii)  $g \cdot (h \cdot x) = (gh) \cdot x$  pour tous  $g, h \in G$  et  $x \in X$ .

En clair, la première propriété assure que l'élément neutre agit de manière triviale, tandis que la seconde postule la compatibilité de l'action avec la loi de composition du groupe : si l'on fait agir un élément  $h$  du groupe sur  $x$ , puis un élément  $g$  du groupe sur  $h \cdot x$ , on obtient le même résultat qu'en faisant agir directement l'élément  $gh$  sur  $x$ . Ces deux axiomes sont somme toute assez naturels, mais ici encore, la seule manière de "comprendre" une telle définition est de se convaincre que les actions de groupes sont omniprésentes en mathématiques.

Avant de donner une liste de tels exemples, commençons par une remarque :

Une action de  $G$  sur  $X$  est équivalente à un homomorphisme  $G \rightarrow S(X)$ .

Plus précisément, une action de  $G$  sur  $X$  induit un homomorphisme  $\phi: G \rightarrow S(X)$  donné par  $\phi(g)(x) = g \cdot x$ , et cette construction définit une bijection entre les actions de  $G$  sur  $X$  et les homomorphismes de  $G$  dans  $S(X)$ . Cet énoncé découle des définitions : aucune idée originale n'est requise pour la preuve. Il est vivement recommandé de tenter *seul* d'en donner une démonstration. C'est un excellent test du niveau d'assimilation des méthodes abstraites et formelles de ce chapitre. On ne saurait trop répéter qu'une maîtrise totale de ces méthodes est requise pour une bonne compréhension du cours !

#### Terminologie.

1. Une action est dite **fidèle** si seul l'élément neutre agit de manière triviale. En d'autres termes : si  $g \in G$  est tel que  $g \cdot x = x$  pour tout  $x \in X$ , alors  $g = e$ . Par la proposition I.2, cela correspond au fait que l'homomorphisme  $G \rightarrow S(X)$  associé est injectif.
2. Une action est dite **transitive** si pour tous  $x, y \in X$ , il existe un  $g \in G$  tel que  $g \cdot x = y$ .

Voici à présent une liste d'exemples.

**Exemples** (d'actions de groupes).

1. L'action *triviale* de  $G$  sur  $X$  est donnée par  $g \cdot x = x$  pour tous  $g \in G$  et  $x \in X$ . Elle correspond à l'homomorphisme trivial  $G \rightarrow S(X)$ . Elle n'est jamais fidèle (sauf si  $G$  est le groupe trivial) et jamais transitive (sauf si  $X$  a un seul élément).
2. Le groupe symétrique  $S_n$  agit sur l'ensemble  $\{1, 2, \dots, n\}$  par permutations. Formellement,  $\sigma \in S_n$  agit sur  $x \in \{1, \dots, n\}$  par  $\sigma \cdot x = \sigma(x)$ . Cela correspond à l'homomorphisme identité  $S_n \rightarrow S(\{1, \dots, n\}) = S_n$ . Cette action est toujours fidèle et transitive. Plus généralement et par le même raisonnement,  $S(X)$  agit fidèlement et transitivement sur  $X$ .
3. Le groupe général linéaire  $GL(n, \mathbb{R})$  agit sur l'ensemble  $X = \mathbb{R}^n$  par multiplication d'un vecteur par une matrice. Cette action est fidèle : si  $M \cdot x = x$  pour tout  $x \in \mathbb{R}^n$ , alors  $M$  est la matrice identité. En revanche, cette action n'est pas transitive (sauf si  $n = 0$ ) puisque  $M \cdot 0 = 0$  pour toute matrice  $M$ .
4. Si  $V$  est un espace vectoriel réel, alors le groupe  $\mathbb{R}^* = \mathbb{R} \setminus \{0\}$  agit sur l'ensemble  $V$  par multiplication scalaire. En effet,  $1 \cdot v = v$  pour tout  $v \in V$  et  $\lambda \cdot (\mu \cdot v) = (\lambda\mu) \cdot v$  pour tous  $\lambda, \mu \in \mathbb{R}^*$  et  $v \in V$  : ce sont deux des axiomes des espaces vectoriels. Cette action est fidèle, mais n'est pas transitive. En particulier,  $\mathbb{R}^*$  agit sur l'ensemble  $X = \mathbb{R}^n \setminus \{0\}$  ; on y reviendra.
5. Le groupe diédral  $D_{2n}$  agit sur l'ensemble  $X = \mu_n(\mathbb{C})$  des sommets du  $n$ -gone régulier. Cette action est fidèle (sauf si  $n = 2$ ), et transitive.
6. Le groupe  $\mathbb{S}^1$  agit sur l'ensemble  $\mathbb{C}$  par multiplication complexe. L'action est fidèle, mais pas transitive.

### I.3.2 Relations d'équivalence et orbites

Rappelons qu'une *relation* sur un ensemble  $X$  est un sous-ensemble  $\mathcal{R} \subset X \times X$ . On note habituellement l'assertion  $(x, y) \in \mathcal{R}$  par  $x\mathcal{R}y$ . Au risque de se répéter, la notion suivante est elle aussi absolument fondamentale en mathématiques.

#### Définition I.6

Une relation  $\mathcal{R}$  sur un ensemble  $X$  est une **relation d'équivalence** si :

- $\mathcal{R}$  est *réflexive* : pour tout  $x \in X$ ,  $x\mathcal{R}x$ .
- $\mathcal{R}$  est *symétrique* : pour tous  $x, y \in X$ , si  $x\mathcal{R}y$  alors  $y\mathcal{R}x$ .
- $\mathcal{R}$  est *transitive* : pour tous  $x, y, z \in X$ , si  $x\mathcal{R}y$  et  $y\mathcal{R}z$ , alors  $x\mathcal{R}z$ .

Pour alléger la notation, une relation d'équivalence  $\mathcal{R}$  est souvent notée par le symbole  $\sim$ . Étant donné  $x \in X$ , on note habituellement  $[x]$  le sous-ensemble de  $X$  défini par

$$[x] := \{y \in X \mid x \sim y\}.$$

On parle de la **classe d'équivalence** de  $x$  modulo la relation d'équivalence  $\sim$ . Notons que par réflexivité,  $x$  appartient à sa classe d'équivalence. De plus, par symétrie et transitivité, deux classes d'équivalence  $[x]$  et  $[y]$  sont soit identiques (si  $x \sim y$ ), soit disjointes (sinon). Ainsi, les classes d'équivalence forment une *partition* de  $X$  : en d'autres termes, l'ensemble  $X$  est l'union disjointe de ces classes. Plus précisément, si  $S \subset X$  contient exactement un élément dans chaque classe d'équivalence de  $X$ , alors  $X = \bigsqcup_{x \in S} [x]$ . On dit que  $S$  est un *système de représentants* de ces classes d'équivalence.

Finalement, on notera  $X/\sim$  l'ensemble de ces classes d'équivalence.

**Exemples** (de relations d'équivalence).

1. Étant donné un ensemble quelconque  $X$ , la relation d'égalité définie par  $x \sim y \iff x = y$  est une relation d'équivalence sur  $X$ . Chaque élément  $x \in X$  est seul dans sa classe, et  $X/\sim = X$ .
2. "Être de la même parité que" est une relation d'équivalence sur l'ensemble des entiers  $\mathbb{Z}$ , et  $\mathbb{Z}/\sim = \{[0], [1]\} = \mathbb{Z}_2$ . Plus généralement, pour  $n$  fixé, la relation

$$x \sim y \iff x - y \text{ est divisible par } n$$

est une relation d'équivalence sur  $\mathbb{Z}$ , et  $\mathbb{Z}/\sim = \{[0], \dots, [n-1]\} = \mathbb{Z}_n$ .

3. La relation  $\leq$  n'est pas une relation d'équivalence sur  $\mathbb{Z}$  : elle est réflexive, transitive, mais pas symétrique. Comme vous l'avez vu en Analyse I, il s'agit de ce qu'on appelle une *relation d'ordre*.
4. La relation sur le plan  $\mathbb{R}^2$  définie par  $x \mathcal{R} y \iff d(x, y) < 1$  n'est pas une relation d'équivalence : elle est réflexive, symétrique, mais pas transitive.

Le lien entre les notions d'action de groupes et de relation d'équivalence est donné par la proposition suivante.

**Proposition I.3.** *Si un groupe  $G$  agit sur un ensemble  $X$ , alors la relation sur  $X$  définie par*

$$x \sim y \iff \text{il existe } g \in G \text{ tel que } g \cdot x = y$$

*est une relation d'équivalence.*

*Démonstration.* Par la première propriété d'une action de groupe,  $e \cdot x = x$  pour tout  $x \in X$  ; la relation est donc réflexive. Soient à présent  $x, y \in X$  tels que  $x \sim y$  ; par définition, il existe donc  $g \in G$  tel que  $g \cdot x = y$ . Soit  $g^{-1} \in G$  son inverse ; par définition d'une action de groupe, on a

$$g^{-1} \cdot y = g^{-1} \cdot (g \cdot x) = (g^{-1}g) \cdot x = e \cdot x = x.$$

En particulier, on a bien  $y \sim x$  et la relation est symétrique. Finalement, soient  $x, y, z \in X$  tels que  $x \sim y$  et  $y \sim z$  ; il existe donc  $g, h \in G$  tels que  $g \cdot x = y$  et  $h \cdot y = z$ . L'action de l'élément  $hg \in G$  sur  $x$  est alors

$$(hg) \cdot x = h \cdot (g \cdot x) = h \cdot y = z.$$

Ainsi, on a bien  $x \sim z$ , et la relation est transitive. □

La classe d'équivalence de  $x \in X$  modulo cette relation d'équivalence est appelée **l'orbite** de  $x$  par l'action de  $G$ , et notée  $G \cdot x$ . Par définition,

$$G \cdot x = \{g \cdot x \mid g \in G\}.$$

L'ensemble des orbites est noté  $X/G$ . Notons qu'une action est transitive si et seulement si elle définit une unique orbite.

La proposition I.3 et la discussion la précédant impliquent immédiatement le résultat suivant.

**Corollaire I.4.** *Soient  $G$  un groupe agissant sur un ensemble  $X$  et  $S \subset X$  un système de représentants des orbites de cette action. Alors,  $X = \bigsqcup_{x \in S} G \cdot x$ .  $\square$*

Reprenons par le menu les exemples d'actions de groupes vues ci-dessus, et tentons de comprendre les orbites correspondantes.

**Exemples (d'orbites).**

1. L'action triviale d'un groupe  $G$  induit la relation d'égalité sur  $X$ , et l'on a donc  $X/G = X$ .
2. Comme l'action de  $G = S(X)$  sur  $X$  est transitive, il y a une unique orbite et l'ensemble  $X/G$  est réduit à un point.
3. On peut vérifier que pour tous  $x, y \in \mathbb{R}^n$  non-nuls, il existe une matrice  $M$  inversible telle que  $M \cdot x = y$ . Par conséquent, l'action de  $\text{GL}(n, \mathbb{R})$  sur  $\mathbb{R}^n$  définit deux orbites : l'une ne contient que l'origine, et l'autre contient tous les vecteurs non-nuls de  $\mathbb{R}^n$ .
4. Deux éléments  $x, y$  de  $\mathbb{R}^n \setminus \{0\}$  sont dans la même orbite par l'action de  $\mathbb{R}^*$  si et seulement s'ils sont sur la même droite passant par l'origine. Ainsi, l'ensemble des orbites est l'ensemble des droites de  $\mathbb{R}^n \setminus \{0\}$  passant par l'origine. Cet ensemble, noté  $\mathbb{P}^{n-1}$ , est appelé *l'espace projectif réel de dimension  $n - 1$* .
5. L'action de  $D_{2n}$  sur les sommets du  $n$ -gone régulier étant transitive, elle ne définit qu'une seule orbite.
6. Les orbites de l'action de  $\mathbb{S}^1$  sur  $\mathbb{C}$  sont les cercles centrés en l'origine de rayon  $r \in [0, \infty)$ . Ainsi,  $\mathbb{C}/\mathbb{S}^1 = [0, \infty)$ .

### I.3.3 Le stabilisateur

Introduisons encore un peu de terminologie. Pour  $x \in X$ , on appelle **stabilisateur** de  $x$  le sous-ensemble de  $G$  donné par

$$\text{Stab}(x) := \{g \in G \mid g \cdot x = x\}.$$

La notation malheureuse  $G_x$  est aussi très souvent utilisée. Notons qu'une action est fidèle si et seulement si l'intersection de tous les stabilisateurs se réduit à l'élément neutre :  $\bigcap_{x \in X} \text{Stab}(x) = \{e\}$ .

**Proposition I.5.** *Le stabilisateur de  $x$  est un sous-groupe de  $G$ . De plus, si  $x$  et  $y$  sont dans la même orbite, alors  $\text{Stab}(x)$  et  $\text{Stab}(y)$  sont des sous-groupes isomorphes de  $G$ .*

*Démonstration.* La première assertion est laissée en exercices. Soient donc  $x, y \in X$  tels qu'il existe  $g \in G$  avec  $g \cdot x = y$ . Rappelons qu'alors,  $g^{-1} \cdot y = x$ . Considérons l'application  $\phi_g: G \rightarrow G$  définie par  $\phi_g(h) = ghg^{-1}$ . Les égalités suivantes dans  $G$  montrent qu'il s'agit d'un homomorphisme :

$$\phi_g(h_1 h_2) = g(h_1 h_2)g^{-1} = gh_1(g^{-1}h_2g^{-1}) = (gh_1g^{-1})(gh_2g^{-1}) = \phi_g(h_1)\phi_g(h_2).$$

De plus, si  $h$  est dans le stabilisateur de  $x$  (i.e. si  $h \cdot x = x$ ), alors  $\phi_g(h)$  est dans le stabilisateur de  $y$ ; cela découle des égalités suivantes dans  $X$  :

$$\phi_g(h) \cdot y = (ghg^{-1}) \cdot y = (gh) \cdot (g^{-1} \cdot y) = (gh) \cdot x = g \cdot (h \cdot x) = g \cdot x = y.$$

Ainsi,  $\phi_g$  définit un homomorphisme  $\text{Stab}(x) \rightarrow \text{Stab}(y)$ . De même,  $\phi_{g^{-1}}$  définit un homomorphisme  $\text{Stab}(y) \rightarrow \text{Stab}(x)$ . L'égalité

$$(\phi_g \circ \phi_{g^{-1}})(h) = \phi_g(\phi_{g^{-1}}(h)) = g(g^{-1}hg)g^{-1} = (gg^{-1})h(gg^{-1}) = h$$

montre que  $\phi_g \circ \phi_{g^{-1}} = \text{id}$  et l'on vérifie similairement que  $\phi_{g^{-1}} \circ \phi_g = \text{id}$ . Ainsi,  $\phi_g: \text{Stab}(x) \rightarrow \text{Stab}(y)$  est un isomorphisme, ce qui conclut la preuve.  $\square$

### I.3.4 La formule des orbites

Le théorème suivant constituera notre outil principal pour la détermination des groupes de symétries d'objets géométriques au chapitre suivant.

#### Théorème I.6: Formule des orbites

Si un groupe fini  $G$  agit sur un ensemble  $X$ , alors pour tout  $x \in X$ , on a

$$|G| = |G \cdot x| \cdot |\text{Stab}(x)|.$$

*Démonstration.* Soit donc  $x \in X$  fixé. Considérons l'application  $\vartheta: G \rightarrow G \cdot x$  donnée par  $\vartheta(g) = g \cdot x$ . Comme c'est une application bien définie, il suit

$$G = \bigsqcup_{y \in G \cdot x} \vartheta^{-1}(y).$$

Pour tout  $y \in G \cdot x$  fixé, il existe  $h \in G$  tel que  $y = h \cdot x$ . Cela permet de définir une application  $\vartheta^{-1}(y) \rightarrow G$  par  $g \mapsto h^{-1}g$ . Mais si  $g$  est un élément de  $\vartheta^{-1}(y)$  (i.e. si  $g \cdot x = y$ ), alors  $h^{-1}g$  est un élément du stabilisateur de  $x$ ; cela découle des égalités

$$(h^{-1}g) \cdot x = h^{-1} \cdot (g \cdot x) = h^{-1} \cdot y = h^{-1} \cdot (h \cdot x) = (h^{-1}h) \cdot x = x.$$

Ainsi, l'assignation  $g \mapsto h^{-1}g$  définit une application  $\vartheta^{-1}(y) \rightarrow \text{Stab}(x)$ , qui est clairement bijective (d'inverse  $g \mapsto hg$ ). Par conséquent, pour tout  $y \in G \cdot x$ ,  $\vartheta^{-1}(y)$  est en bijection avec  $\text{Stab}(x)$ . Il suit

$$|G| = \sum_{y \in G \cdot x} |\vartheta^{-1}(y)| = \sum_{y \in G \cdot x} |\text{Stab}(x)| = |G \cdot x| \cdot |\text{Stab}(x)|,$$

et le théorème est démontré.  $\square$

L'énoncé suivant est une conséquence du corollaire I.4 et du théorème I.6.

**Corollaire I.7.** *Soient  $G$  un groupe fini agissant sur un ensemble fini  $X$ , et  $S \subset X$  un système de représentants des orbites de cette action. Alors,*

$$|X| = \sum_{x \in S} \frac{|G|}{|\text{Stab}(x)|}. \quad \square$$

Avant de donner de véritables applications de la formule des orbites, tentons de la comprendre sur quelques exemples simples.

### Exemples.

1. Dans le cas de l'action triviale,  $G \cdot x = \{x\}$  et  $\text{Stab}(x) = G$  pour tout  $x \in X$ . La formule des orbites se lit alors simplement  $|G| = 1 \cdot |G|$ .
2. Considérons l'action de  $G = S_n$  sur  $X = \{1, \dots, n\}$  et fixons  $x = n \in X$ . Comme cette action est transitive,  $G \cdot x = X$ . De plus,

$$\text{Stab}(x) = \{\sigma \in S_n \mid \sigma(n) = n\} = S(\{1, \dots, n-1\}) = S_{n-1}.$$

Ainsi, le théorème I.6 implique l'égalité  $|S_n| = n \cdot |S_{n-1}|$ , et récursivement  $|S_n| = n \cdot (n-1) \cdots 2 \cdot 1 = n!$ .

3. Le groupe diédral  $D_{2n}$  agit transitivement sur l'ensemble  $X = \mu_n(\mathbb{C})$  des sommets du  $n$ -gone régulier. Pour  $x = 1 \in X$ , on a donc  $|G \cdot x| = |X| = n$ . De plus,  $\text{Stab}(x) = \{\text{id}, s\}$ . Ainsi, on retrouve le fait que l'ordre de  $D_{2n}$  est égal à  $2n$ .

Comme nous l'avons mentionné plus haut, la formule des orbites trouvera de nombreuses applications géométriques au chapitre II. Mais elle permet aussi d'obtenir des résultats fondamentaux en théorie des groupes. En voici un exemple.

### Le théorème de Lagrange sur les groupes.

Soit  $H$  un sous-groupe d'un groupe fini  $G$ . Alors, le groupe  $H$  agit sur l'ensemble  $X = G$  par la loi de composition dans  $G$  :  $h \cdot x = hx$  pour  $h \in H$  et  $x \in G$ . De plus, pour tout  $x \in G$ ,

$$\text{Stab}(x) = \{h \in H \mid hx = x\} = \{e\}.$$

Le nombre d'orbites de cette action, noté  $[G : H]$ , est appelé *l'indice* de  $H$  dans  $G$ . Par le corollaire I.7 et l'équation ci-dessus,

$$|G| = \sum_{x \in S} \frac{|H|}{|\text{Stab}(x)|} = \sum_{x \in S} |H| = [G : H] \cdot |H|.$$

On a donc démontré :

**Théorème I.8: Théorème de Lagrange**

Si  $H$  est un sous-groupe d'un groupe fini  $G$ , alors on a  $|G| = [G : H] \cdot |H|$ .  $\square$

**I.3.5 La formule de Burnside**

Le dernier résultat de ce chapitre, une magnifique formule probablement due à CAUCHY et à FROBENIUS mais qui porte le nom de BURNSIDE, nécessite encore un peu de terminologie.

Étant donnée une action d'un groupe  $G$  sur un ensemble  $X$ , l'ensemble des **points fixes** de  $g \in G$  est le sous-ensemble  $X^g \subset X$  défini par

$$X^g := \{x \in X \mid g \cdot x = x\}.$$

La formule de Burnside dit essentiellement que le nombre d'orbites d'une action est égal à la moyenne du nombre de points fixes des éléments du groupe. Plus précisément :

**Théorème I.9: Formule de Burnside**

Si un groupe fini  $G$  agit sur un ensemble fini  $X$ , alors

$$|X/G| = \frac{1}{|G|} \sum_{g \in G} |X^g|.$$

*Démonstration.* La preuve que nous proposons repose sur ce qu'on appelle un *argument de double comptage*. Ce type d'argument, un grand classique en théorie combinatoire, consiste à compter le nombre d'éléments d'un ensemble donné de deux manières différentes, puis à comparer les deux résultats. L'ensemble à considérer dans notre cas est le sous-ensemble  $Y \subset G \times X$  défini par

$$Y := \{(g, x) \in G \times X \mid g \cdot x = x\},$$

dont on va compter les éléments “verticalement” et “horizontalement”, comme illustré en figure I.1. Commençons par évaluer le cardinal de  $Y$  “verticalement” : l'ensemble  $G \times X$  se décompose en l'union disjointe de toutes les “tranches verticales”  $\{g\} \times X$  pour  $g \in G$ . Ainsi,

$$Y = Y \cap (G \times X) = Y \cap \left( \bigsqcup_{g \in G} \{g\} \times X \right) = \bigsqcup_{g \in G} Y \cap (\{g\} \times X) = \bigsqcup_{g \in G} \{g\} \times X^g,$$

où la dernière égalité découle des définitions de  $Y$  et de  $X^g$ . Par conséquent, le cardinal de  $Y$  est donné par

$$|Y| = \sum_{g \in G} |X^g|. \quad (\star)$$

Faisons à présent le même raisonnement, mais “horizontalement” cette fois-ci : on obtient

$$Y = Y \cap (G \times X) = Y \cap \left( \bigsqcup_{x \in X} G \times \{x\} \right) = \bigsqcup_{x \in X} Y \cap (G \times \{x\}) = \bigsqcup_{x \in X} \text{Stab}(x) \times \{x\}.$$

Notons  $N = |X/G|$  le nombre d'orbites, et fixons un système  $S = \{x_1, \dots, x_N\}$  de représentants des orbites. Par l'égalité ci-dessus, le corollaire I.4, la proposition I.5 et le théorème I.6, il suit

$$|Y| = \sum_{x \in X} |\text{Stab}(x)| = \sum_{i=1}^N \sum_{x \in G \cdot x_i} |\text{Stab}(x)| = \sum_{i=1}^N |G \cdot x_i| \cdot |\text{Stab}(x_i)| = \sum_{i=1}^N |G| = N \cdot |G|.$$

La comparaison de l'équation  $(*)$  et de l'égalité ci-dessus donne le résultat.  $\square$

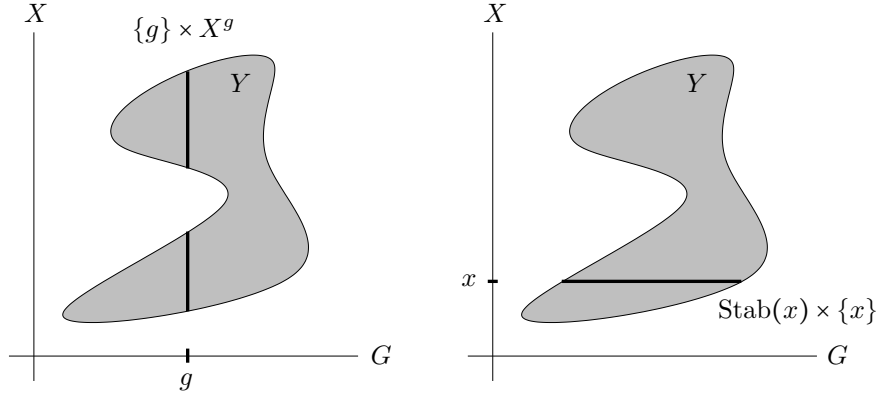


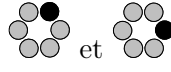
FIGURE I.1 – L'argument de double comptage.

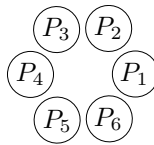
### Application à l'énumération d'objets symétriques.

La formule de Burnside est extrêmement utile pour les problèmes d'énumération d'objets symétriques. En voici un exemple. Fixons un entier  $k \geq 1$ .

Combien de colliers à six perles peut-on fabriquer avec  $k$  types de perles à disposition ?

La difficulté réside dans le fait qu'un collier à six perles possède des symétries.

Par exemple, les colliers représentés par les dessins  et  ne doivent pas être comptés comme deux colliers différents, mais bien comme le même collier. Pour résoudre ce problème, numérotons de  $P_1$  à  $P_6$  les six perles apparaissant sur chaque dessin de collier, par exemple comme suit :



L'ensemble de tous les dessins possibles de colliers à six perles de  $k$  types possibles est donné par

$$X = \{x: \{P_1, \dots, P_6\} \rightarrow \{1, \dots, k\}\}.$$

Notons que les six perles sont disposées selon les sommets d'un hexagone régulier. Ainsi, le groupe diédral  $G = D_{12}$  agit sur ces perles, et du coup, sur

l'ensemble  $X$  ; de manière plus formelle, l'action de  $g \in D_{12}$  sur  $x \in X$  est donnée<sup>4</sup> par  $(g \cdot x)(P_i) = x(g^{-1}(P_i))$ . De plus, deux dessins de colliers (i.e. deux éléments de  $X$ ) décrivent le même collier si et seulement si ces éléments de  $X$  sont dans la même orbite pour l'action du groupe diédral. Ainsi, compter le nombre de colliers différents, c'est compter le nombre  $|X/G|$  d'orbites de cette action.

Par la formule de Burnside, il ne reste plus qu'à évaluer  $|X^g|$  pour chacun des 12 éléments du groupe diédral  $D_{12} = \{r^\ell s^\varepsilon \mid \ell = 0, \dots, 5 \text{ et } \varepsilon = 0, 1\}$ . Allons-y :

- Pour  $g = e$ , on a évidemment  $X^e = X$ , d'où  $|X^e| = k^6$ .
- L'élément  $g = r$  est la rotation d'une perle dans le sens trigonométrique. Ainsi,  $X^r$  est l'ensemble des dessins de colliers invariants par rotation d'une perle ; ce sont les dessins "constants", i.e. dont toutes les perles sont du même type. On obtient donc  $|X^g| = k$  pour  $g = r$ , et pour  $g = r^5$  par le même raisonnement.
- Pour  $g = r^2$  et  $g = r^4$ ,  $X^g$  est l'ensemble des dessins de colliers invariants par rotation de deux perles ; ce sont les dessins de la forme . Ils sont au nombre de  $|X^g| = k^2$ , puisque deux types de perles peuvent être utilisées.
- Pour  $g = r^3$ ,  $X^g$  est l'ensemble des dessins invariants par rotation de trois perles, i.e. des dessins de la forme . On obtient donc  $|X^g| = k^3$ .
- L'éléments  $g = s$  est la réflexion par l'axe horizontal. Les points fixes  $X^g$  sont donc les dessins de la forme , qui sont au nombre de  $k^4$ . De même,  $|X^g| = k^4$  pour  $g = r^2s$  et  $g = r^4s$ .
- Finalement,  $g = rs$  est la réflexion par l'axe de symétrie<sup>5</sup> passant entre les perles  $P_1$  et  $P_2$  et entre les perles  $P_4$  et  $P_5$ . Les  $k^3$  points fixes correspondants sont les dessins de la forme . De même,  $|X^g| = k^3$  pour  $g = r^3s$  et  $g = r^5s$ .

En additionnant ces 12 contributions, on obtient donc la réponse suivante :

Avec  $k$  types de perles à disposition, on peut fabriquer exactement

$$\frac{1}{12}(2k + 2k^2 + 4k^3 + 3k^4 + k^6)$$

colliers à six perles.

Par exemple, on peut fabriquer 1 collier avec 1 type de perles, 13 avec 2 types, 92 avec 3, 430 avec 4, 1505 avec 5, 4291 avec 6 et 10528 colliers avec 7 types de perles.

4. la présence de l'inverse est nécessaire pour avoir une action, mais il ne joue pas de rôle dans cette discussion puisque  $X^{g^{-1}} = X^g$  pour tout  $g \in G$ .

5. rappelons que dans le groupe diédral, les éléments se composent comme des applications ou des matrices, i.e. de droite à gauche.

## Chapitre II: Groupes d'isométries

Dans ce chapitre, nous allons appliquer les notions algébriques abstraites vues précédemment à des questions géométriques. Dans un premier temps, nous verrons que l'ensemble des isométries (d'un espace métrique  $X$ ) est un groupe qui agit naturellement sur  $X$ . Nous calculerons ensuite ce groupe pour l'espace  $\mathbb{R}^n$  (Section 2.2), puis nous donnerons une classification de ces isométries dans le cas du plan. Finalement, en section 2.4, nous utiliserons les résultats du chapitre précédent pour déterminer les groupes de symétries de certains objets géométriques.

### II.1 Distances et isométries

#### II.1.1 Espaces métriques

Dans un premier temps, nous allons considérer les espaces métriques généraux. En voici la définition.

##### Définition II.1

Soit  $X$  un ensemble. Une application  $d: X \times X \rightarrow [0, \infty)$  est appelée une **distance** (ou une **métrique**) sur  $X$  si elle satisfait les trois propriétés suivantes pour tous  $x, y, z$  dans  $X$  :

- *Identité des indiscernables* :  $d(x, y) = 0$  si et seulement si  $x = y$ .
- *Symétrie* :  $d(x, y) = d(y, x)$ .
- *Inégalité du triangle* :  $d(x, z) \leq d(x, y) + d(y, z)$ .

Un ensemble  $X$  muni d'une distance est appelé un **espace métrique**.

##### Remarques et terminologie.

1. Formellement, un espace métrique est donc la donnée d'une paire  $(X, d)$ , mais on le note souvent  $X$  lorsqu'il n'y a pas de confusion possible.
2. Si  $Y$  est un sous-ensemble de  $X$ , une métrique  $d$  sur  $X$  définit automatiquement une métrique sur  $Y$  : on parle de la *métrique induite*.
3. Étant donné un élément  $x$  d'un espace métrique  $(X, d)$  et un nombre réel  $r \geq 0$ , on appelle la *boule* (fermée) de rayon  $r$  centrée en  $x$  l'ensemble

$$B_d(x, r) := \{y \in X \mid d(x, y) \leq r\}.$$

Les trois axiomes d'une métrique formalisent les propriétés évidentes de la distance standard, dite *euclidienne*, dans le plan ou l'espace. Il est important de se rendre compte que cette définition recouvre néanmoins des objets extrêmement variés, qui apparaissent dans de nombreux domaines des mathématiques. Dans le cadre de ce cours, nous nous bornerons à une liste d'exemples assez peu exotiques. Vous en trouverez d'autres en exercices.

**Exemples** (d'espaces métriques).

1. L'ensemble  $X = \mathbb{R}$  admet la métrique  $d$  donnée par  $d(x, y) = |x - y|$ . Les boules correspondantes sont les intervalles :  $B_d(x, r) = [x - r, x + r]$ .
2. Sur le plan  $X = \mathbb{R}^2$ , la métrique la plus naturelle est la *métrique euclidienne* définie par la longueur du segment entre les deux points ; on parle aussi de *métrique standard*. L'inégalité du triangle, intuitivement évidente, n'est autre que la proposition I.20 des Éléments d'Euclide. Le théorème de Pythagore nous permet de calculer cette distance en coordonnées : si les points  $x$  et  $y$  ont coordonnées  $(x_1, x_2)$  et  $(y_1, y_2)$ , respectivement, alors

$$d(x, y) = \sqrt{(x_1 - y_1)^2 + (x_2 - y_2)^2}.$$

Notons que la preuve analytique de l'inégalité du triangle pour cette métrique repose sur *l'inégalité de Cauchy-Schwarz*, que vous avez vue en analyse. Les boules correspondantes sont bien évidemment les boules "standards".

3. Sur le plan, la formule

$$d(x, y) = |x_1 - y_1| + |x_2 - y_2|$$

définit une métrique parfois appelée la *métrique new-yorkaise*, pour des raisons évidentes à quiconque se déplace à Manhattan. Les boules correspondantes sont des carrés pivotés d'un angle de  $\pi/4$  (voir figure II.1).

4. D'une manière plus générale, étant donné un réel  $p \geq 1$ , on définit

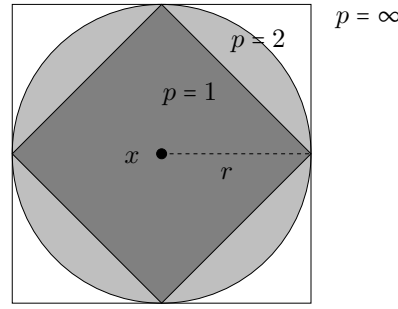
$$d_p(x, y) = \sqrt[p]{|x_1 - y_1|^p + |x_2 - y_2|^p}.$$

Observez que les cas  $p = 1$  et  $p = 2$  correspondent aux métriques new-yorkaises et euclidiennes, respectivement. Si les deux premières propriétés sont – comme c'est souvent le cas – évidentes, la preuve de l'inégalité du triangle repose sur *l'inégalité de Minkowski*. À la limite, lorsque  $p = \infty$ , on obtient la distance

$$d_\infty(x, y) = \max\{|x_1 - y_1|, |x_2 - y_2|\}.$$

Les boules correspondantes sont illustrées en figure II.1.

Notons que cette famille d'exemples se généralise de manière immédiate à l'espace  $\mathbb{R}^n$ .

FIGURE II.1 – Les boules  $B_{d_p}(x, r)$  pour  $p = 1$ ,  $p = 2$ , et  $p = \infty$ .

5. On peut munir un ensemble non-vide  $X$  quelconque de la *métrie discrète* définie par

$$d(x, y) = \begin{cases} 0 & \text{si } x = y; \\ 1 & \text{sinon.} \end{cases}$$

La boule correspondante  $B_d(x, r)$  est réduite à  $\{x\}$  si  $r < 1$ , et couvre  $X$  tout entier si  $r \geq 1$ .

### II.1.2 Le groupe d'isométries d'un espace métrique

La notion d'isométrie de  $\mathbb{R}^n$  vue en première partie du cours se généralise de manière évidente à un espace métrique quelconque.

#### Définition II.2

Soit  $(X, d)$  un espace métrique. Une **isométrie** de  $X$  est une application  $f: X \rightarrow X$  bijective qui préserve les distances, i.e. telle que

$$d(f(x), f(y)) = d(x, y) \quad \text{pour tous } x, y \in X.$$

On notera  $\text{Isom}(X)$  l'ensemble des isométries de  $X$ .

Notons que cette définition est légèrement redondante. En effet une application  $f: X \rightarrow X$  qui préserve les distances sera automatiquement injective par identité des indiscernables. Dans  $\mathbb{R}^n$ , toute application préservant la distance euclidienne est automatiquement bijective ; c'est une conséquence des résultats vus en première partie du cours. En revanche, il existe des espaces métriques  $(X, d)$  et des applications  $f: X \rightarrow X$  qui préservent la distance  $d$  mais ne sont pas surjectives. L'exemple le plus simple est sans doute  $X = [0, \infty)$  muni de la métrique usuelle  $d(x, y) = |x - y|$ , et  $f: [0, \infty) \rightarrow [0, \infty)$  donnée par  $f(x) = x + 1$ .

Le lien entre la notion géométrique d'isométrie et celle plus algébrique d'action de groupe est donnée par la proposition suivante.

**Proposition II.1.** *Soit  $(X, d)$  un espace métrique. L'ensemble  $\text{Isom}(X)$  des isométries de  $X$  est un groupe pour la loi de composition donnée par la composition des applications, et ce groupe agit de manière fidèle sur l'ensemble  $X$ .*

*Démonstration.* Il y a plusieurs points à vérifier, tous à peu près triviaux. Tout d'abord, la composition de deux isométries est une isométrie : la composition de deux applications bijectives est en effet bijective, et si  $f, g: X \rightarrow X$  sont deux applications qui préservent les distances, alors leur composition  $f \circ g$  fait de même ; cela découle des égalités

$$d((f \circ g)(x), (f \circ g)(y)) = d(f(g(x)), f(g(y))) = d(g(x), g(y)) = d(x, y).$$

Bien évidemment, la composition des applications est associative, et l'élément neutre pour cette loi de composition, l'identité  $\text{id}_X$  sur  $X$ , est trivialement une isométrie. De plus, toute isométrie  $f$  étant par définition bijective, elle admet une application inverse  $f^{-1}: X \rightarrow X$ , et cet inverse préserve à son tour les distances :

$$d(f^{-1}(x), f^{-1}(y)) = d(f(f^{-1}(x)), f(f^{-1}(y))) = d(x, y).$$

Ainsi,  $\text{Isom}(X)$  est un groupe. L'action d'une isométrie  $f$  sur un élément  $x \in X$  est tout simplement donnée par  $f \cdot x := f(x)$ . Cela définit bel et bien une action du groupe  $\text{Isom}(X)$  sur l'ensemble  $X$ , puisque

$$e \cdot x = \text{id}_X \cdot x = \text{id}_X(x) = x \quad \text{et} \quad f \cdot (g \cdot x) = f(g(x)) = (f \circ g)(x) = (f \circ g) \cdot x$$

pour tout  $x \in X$ . Finalement, cette action est fidèle car la seule application  $f: X \rightarrow X$  telle que  $f(x) = x$  pour tout  $x \in X$  est l'application  $f = \text{id}_X$ .  $\square$

### Exemples.

1. Soit  $X$  un ensemble quelconque. Toute application injective  $f: X \rightarrow X$  préserve la métrique discrète. Ainsi, pour cette métrique,

$$\text{Isom}(X) = \{f: X \rightarrow X \mid f \text{ est bijective}\} = S(X),$$

le groupe symétrique sur  $X$ .

(Notons qu'en général,  $\text{Isom}(X)$  est un sous-groupe de  $S(X)$ .)

2. Soit  $X = \{0, 1, 3\} \subset \mathbb{R}$  muni de la métrique induite par la métrique standard sur  $\mathbb{R}$ . Dans cet exemple,  $\text{Isom}(X)$  n'est autre que le groupe trivial  $\{\text{id}\}$ . En effet, aucune permutation de cet ensemble ne préserve la distance, si ce n'est l'identité.

## II.2 Le groupe des isométries de l'espace $\mathbb{R}^n$

Nous allons dès à présent et pour le reste du chapitre nous concentrer sur une famille d'exemples plus naturels du point de vue géométrique : les espaces  $\mathbb{R}^n$  munis de la métrique euclidienne. La détermination du groupe  $\text{Isom}(\mathbb{R}^n)$  va nécessiter l'introduction d'un dernier concept algébrique abstrait, celui de *produit semi-direct*.

## II.2.1 Le produit semi-direct

**Définition II.3**

Soient  $G$  un groupe, et  $H_1, H_2$  deux sous-groupes de  $G$ . On dit que  $G$  est le **produit semi-direct** de  $H_1$  et  $H_2$ , noté  $G = H_1 \rtimes H_2$ , si les deux conditions suivantes sont satisfaites :

- (i) Tout  $g \in G$  s'écrit de manière unique  $g = h_1 h_2$  avec  $h_1 \in H_1$  et  $h_2 \in H_2$ .
- (ii) Pour tous  $h_1 \in H_1$  et  $h_2 \in H_2$ , on a  $h_2 h_1 h_2^{-1} \in H_1$ .

**Remarques.** 1. La première condition signifie que l'application

$$\phi: H_1 \times H_2 \longrightarrow G, \quad (h_1, h_2) \longmapsto h_1 h_2$$

est bijective. Ainsi, cette condition signifie que, *comme ensemble*,  $G$  est simplement donné par le produit cartésien  $H_1 \times H_2$ .

- 2. Cette première condition est équivalente aux égalités  $H_1 H_2 = G$  (tout élément de  $G$  est produit d'éléments de  $H_1$  et  $H_2$ ) et  $H_1 \cap H_2 = \{e\}$  (cette écriture est unique).
- 3. Dans la seconde condition, l'élément  $h_2 h_1 h_2^{-1}$  est appelé le *conjugué* de  $h_1$  par  $h_2$ . Ainsi, le groupe  $H_2$  agit par conjugaison sur  $H_1$ . Cette seconde condition est équivalente (en présence de la première) au fait que  $H_1$  est un *sous-groupe normal* de  $G$ , noté habituellement par  $H_1 \triangleleft G$ .
- 4. La loi de composition dans  $G$  (et donc, la structure de groupe sur  $G$ ) est entièrement déterminée par les lois de composition dans les groupes  $H_1$  et  $H_2$ , et par l'action par conjugaison de  $H_2$  sur  $H_1$ . En effet, étant donnés deux éléments quelconques de  $G$ , on peut les écrire de manière unique comme  $h_1 h_2$  et  $h'_1 h'_2$ ; leur produit est alors donné par

$$(h_1 h_2)(h'_1 h'_2) = h_1 h_2 h'_1 (h_2^{-1} h_2) h'_2 = h_1 \underbrace{(h_2 h'_1 h_2^{-1})}_{\in H_1 \text{ par (ii)}} (h_2 h'_2) = h''_1 h''_2,$$

avec  $h''_1 = h_1 (h_2 h'_1 h_2^{-1}) \in H_1$  et  $h''_2 = h_2 h'_2 \in H_2$ .

- 5. Un cas particulièrement simple et éclairant est le suivant. Supposons qu'on ait  $G = H_1 \rtimes H_2$  avec  $H_2$  agissant par conjugaison de manière triviale sur  $H_1$ . (Cela signifie que  $h_1 h_2 = h_2 h_1$  pour tous  $h_1 \in H_1$  et  $h_2 \in H_2$ .) Dans ce cas,  $G$  n'est autre que le produit direct de  $H_1$  et  $H_2$  :  $G = H_1 \times H_2$ , cette fois non seulement comme ensemble, mais *comme groupe*. En effet, l'application  $\phi: H_1 \times H_2 \rightarrow G$  définie ci-dessus est alors non seulement bijective par la première condition, mais aussi un homomorphisme de groupes par les égalités

$$\begin{aligned} \phi((h_1, h_2)(h'_1, h'_2)) &= \phi(h_1 h'_1, h_2 h'_2) = h_1 h'_1 h_2 h'_2 = h_1 h_2 h'_1 h'_2 \\ &= \phi(h_1, h_2) \phi(h'_1, h'_2). \end{aligned}$$

Ainsi, il faut comprendre le produit semi-direct comme une généralisation de la notion de produit direct.

**Exemples** (de produits semi-directs).

1. Considérons les deux sous-groupes  $C_n = D_{2n} \cap \text{SO}(2) = \{r^\ell \mid \ell = 0, \dots, n-1\}$  et  $C_2 = \{s^\varepsilon \mid \varepsilon = 0, 1\}$  du groupe diédral  $D_{2n}$ . Alors,  $D_{2n}$  est le produit semi-direct de  $C_n$  et  $C_2$  :

$$D_{2n} = C_n \rtimes C_2.$$

En effet, on a vu que tout élément  $g \in D_{2n}$  s'écrit de manière unique comme  $g = r^\ell s^\varepsilon$ , et donc, comme composition d'un élément de  $C_n$  et d'un élément de  $C_2$ . De plus, l'égalité  $sr^\ell s^{-1} = r^{-\ell}$  implique que la seconde condition est aussi satisfaite.

2. Fixons un entier  $n \geq 1$ , et considérons les groupes  $\text{O}(n)$  et  $\text{SO}(n)$ . Soit  $R \in \text{O}(n)$  une matrice de déterminant  $-1$  telle que  $R^2 = \text{id}$ , par exemple la matrice représentant l'application  $(x_1, x_2, \dots, x_n) \mapsto (-x_1, x_2, \dots, x_n)$ . Alors,

$$\text{O}(n) = \text{SO}(n) \rtimes C_2,$$

où  $C_2 = \{\text{id}, R\}$  (exercice).

3. Le groupe symétrique  $S_n$  est le produit semi-direct du sous-groupe  $A_n$  des permutations paires et de n'importe quel sous-groupe de la forme  $\{\text{id}, \tau\}$ , où  $\tau$  est une transposition (exercice).
4. Le groupe général linéaire satisfait  $\text{GL}(n, \mathbb{R}) = \text{SL}(n, \mathbb{R}) \rtimes \mathbb{R}^*$  ; c'est aussi un exercice.

Un autre exemple, plus géométrique, est donné par l'ensemble des applications affines, auxquelles nous consacrons le paragraphe suivant.

## II.2.2 Applications affines

Étant donné un élément  $v \in \mathbb{R}^n$ , on notera  $\tau_v: \mathbb{R}^n \rightarrow \mathbb{R}^n$  la *translation* par le vecteur  $v$  définie par  $\tau_v(x) = x + v$ . L'ensemble de toutes ces translations forme un groupe pour la composition des applications. La correspondance  $v \leftrightarrow \tau_v$  nous montre qu'il s'agit tout simplement du groupe abélien  $(\mathbb{R}^n, +)$ .

Rappelons la définition d'une application affine vue en première partie du cours.

### Définition II.4

Une application  $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$  est dite **affine** si l'on peut l'écrire comme la composition  $f = \tau_v \circ \alpha$  d'une translation et d'une application linéaire.

La terminologie suivante nous sera aussi utile : un **sous-espace affine** de  $\mathbb{R}^n$  est un sous-ensemble  $E$  de  $\mathbb{R}^n$  de la forme  $E = \tau_v(W)$  avec  $W$  un sous-espace vectoriel de  $\mathbb{R}^n$ . En d'autres termes, on a

$$E = \{x + v \mid x \in W\}$$

pour un  $v \in \mathbb{R}^n$  et un sous-espace vectoriel  $W \subset \mathbb{R}^n$  fixés. La **dimension** de  $E$  est bien-entendu définie comme la dimension de l'espace vectoriel correspondant. Ainsi, les sous-espaces affines de dimension 0 de  $\mathbb{R}^n$  sont tout simplement les points de  $\mathbb{R}^n$ . De même, les sous-espaces affines de dimension 1 de  $\mathbb{R}^n$  sont les droites dans  $\mathbb{R}^n$ .

Ces sous-espaces vont jouer un rôle important dans la compréhension des isométries de  $\mathbb{R}^n$  pour la raison suivante.

**Proposition II.2.** *Étant donnée une application affine  $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$ , son ensemble de points fixes  $\text{Fix}(f) := \{x \in \mathbb{R}^n \mid f(x) = x\}$  est soit vide, soit un sous-espace affine de  $\mathbb{R}^n$ .*

*Démonstration.* Soit donc  $f = \tau_v \circ \alpha$  une application affine. Un élément  $x \in \mathbb{R}^n$  est un point fixe de  $f$  si et seulement si  $x = f(x) = \tau_v(\alpha(x)) = \alpha(x) + v$ , ce que l'on peut écrire  $\beta(x) = v$  avec  $\beta$  l'application linéaire donnée par  $\beta = \text{id} - \alpha$ . Ainsi,

$$\text{Fix}(f) = \{x \in \mathbb{R}^n \mid \beta(x) = v\} = \beta^{-1}(v).$$

Or, pour  $\beta: \mathbb{R}^n \rightarrow \mathbb{R}^n$  linéaire et  $v \in \mathbb{R}^n$ ,  $\beta^{-1}(v)$  est soit vide, soit un sous-espace affine de  $\mathbb{R}^n$ . En effet, en supposant cet ensemble non-vide et en choisissant un élément  $x_0 \in \beta^{-1}(v)$ , il suit

$$\beta^{-1}(v) = \{x \mid \beta(x) = \beta(x_0)\} = \{(x - x_0) + x_0 \mid \beta(x - x_0) = 0\} = \{y + x_0 \mid y \in \ker(\beta)\}.$$

On obtient donc l'égalité  $\text{Fix}(f) = \beta^{-1}(v) = \tau_{x_0}(W)$ , avec  $W$  le sous-espace vectoriel donné par le noyau de  $\beta$ .  $\square$

Une application affine n'est pas bijective en général ; l'ensemble de toutes les applications affines n'a donc aucune chance d'être un groupe. En revanche, si l'on se restreint aux applications affines bijectives de  $\mathbb{R}^n$ , on est bien en présence d'un groupe. Plus précisément :

**Proposition II.3.** *L'ensemble  $\text{Aff}(\mathbb{R}^n)$  des applications affines inversibles de  $\mathbb{R}^n$  est un groupe pour la composition des applications. De plus,*

$$\text{Aff}(\mathbb{R}^n) = \mathbb{R}^n \rtimes \text{GL}(n, \mathbb{R}),$$

où  $\mathbb{R}^n$  est le sous-groupe des translations,  $\text{GL}(n, \mathbb{R})$  celui des applications linéaires inversibles, et où  $\alpha \circ \tau_v \circ \alpha^{-1} = \tau_{\alpha(v)}$  pour tout  $\alpha \in \text{GL}(n, \mathbb{R})$  et  $v \in \mathbb{R}^n$ .

*Démonstration.* Par définition, tout élément  $f \in \text{Aff}(\mathbb{R}^n)$  s'écrit comme la composition  $f = \tau_v \circ \alpha$  d'un élément  $\tau_v$  de  $\mathbb{R}^n$  et d'un élément  $\alpha \in \text{GL}(n, \mathbb{R})$ . De plus, cette écriture est unique. En effet, puisque  $\alpha$  est linéaire,  $\alpha(0) = 0$ , d'où

$$f(0) = (\tau_v \circ \alpha)(0) = \tau_v(\alpha(0)) = \tau_v(0) = 0 + v = v.$$

Ainsi,  $\tau_v$  est uniquement déterminé par  $f$ , et du coup,  $\alpha = \tau_{-v} \circ f$  aussi. De plus, la conjugaison d'une translation  $\tau_v$  par un élément  $\alpha \in \text{GL}(n, \mathbb{R})$  satisfait

$$\alpha(\tau_v(\alpha^{-1}(x))) = \alpha(\alpha^{-1}(x) + v) = \alpha(\alpha^{-1}(x)) + \alpha(v) = x + \alpha(v) = \tau_{\alpha(v)}(x).$$

En d'autres termes,  $\alpha \circ \tau_v \circ \alpha^{-1} = \tau_{\alpha(v)}$ . Il découle de cette égalité que la composition de deux applications affines satisfait

$$(\tau_v \circ \alpha) \circ (\tau_{v'} \circ \alpha') = \tau_v \circ \overbrace{(\alpha \circ \tau_{v'} \circ \alpha^{-1})}^{=\tau_{\alpha(v')}} \circ \alpha \circ \alpha' = \tau_{v+\alpha(v')} \circ (\alpha \circ \alpha'),$$

qui est bien une application affine. Il suit que l'inverse d'une application affine (invertible) est l'application affine donnée par

$$(\tau_v \circ \alpha)^{-1} = \tau_{-\alpha^{-1}(v)} \circ \alpha^{-1}.$$

Comme la composition des applications (affines) est associative et  $\text{id} = \tau_0 \circ \text{id}$  est une application affine, on conclut que  $\text{Aff}(\mathbb{R}^n)$  est un groupe, le produit semi-direct  $\mathbb{R}^n \rtimes \text{GL}(n, \mathbb{R})$ .  $\square$

Insistons sur une conséquence de cette preuve : la composition de deux applications affines est donnée par la formule

$$(\tau_v \circ \alpha) \circ (\tau_{v'} \circ \alpha') = \tau_{v+\alpha(v')} \circ (\alpha \circ \alpha'). \quad (*)$$

Cette formule est évidemment très utile pour calculer des compositions d'applications affines, en particulier d'isométries.

### II.2.3 Le groupe $\text{Isom}(\mathbb{R}^n)$

Nous allons à présent utiliser les résultats de la première partie du cours pour déterminer la structure du groupe des isométries de  $\mathbb{R}^n$ . Rappelons les résultats en question sous la forme d'un seul énoncé.

#### Théorème II.4

Toute isométrie de  $\mathbb{R}^n$  est une application affine. De plus, une application linéaire  $\alpha: \mathbb{R}^n \rightarrow \mathbb{R}^n$  est une isométrie si et seulement si sa matrice dans une base orthonormale de  $\mathbb{R}^n$  est un élément de  $\text{O}(n)$ .  $\square$

Ainsi, il est possible d'associer à toute isométrie  $f$  de  $\mathbb{R}^n$  une unique application linéaire  $\alpha: \mathbb{R}^n \rightarrow \mathbb{R}^n$  dont la matrice dans une base orthonormale est une matrice orthogonale. En particulier, le déterminant de cette matrice sera  $+1$  ou  $-1$ . On dit que l'isométrie  $f$  de  $\mathbb{R}^n$  *préserve l'orientation* dans le premier cas, et qu'elle *renverse l'orientation* dans le second. On note

$$\text{Isom}^+(\mathbb{R}^n) = \{f \in \text{Isom}(\mathbb{R}^n) \mid f \text{ préserve l'orientation}\},$$

et l'on vérifie qu'il s'agit d'un sous-groupe de  $\text{Isom}(\mathbb{R}^n)$ , par exemple au moyen de la formule  $(*)$  ci-dessus.

Tout est maintenant en place pour la détermination de ces groupes.

**Théorème II.5: Isométries de  $\mathbb{R}^n$**

Le groupe  $\text{Isom}(\mathbb{R}^n)$  des isométries de  $\mathbb{R}^n$  est le sous-groupe de  $\text{Aff}(\mathbb{R}^n)$  donné par le produit semi-direct

$$\text{Isom}(\mathbb{R}^n) = \mathbb{R}^n \rtimes \text{O}(n),$$

où  $\mathbb{R}^n$  est le sous-groupe des translations,  $\text{O}(n)$  le groupe orthogonal de degré  $n$ , et où  $\alpha \circ \tau_v \circ \alpha^{-1} = \tau_{\alpha(v)}$  pour tout  $\alpha \in \text{O}(n)$  et  $v \in \mathbb{R}^n$ .

De même, le sous-groupe  $\text{Isom}^+(\mathbb{R}^n)$  des isométries de  $\mathbb{R}^n$  préservant l'orientation est le produit semi-direct

$$\text{Isom}^+(\mathbb{R}^n) = \mathbb{R}^n \rtimes \text{SO}(n).$$

*Démonstration.* Par le premier point du théorème II.4, toute isométrie  $f$  de  $\mathbb{R}^n$  s'écrit  $f = \tau_v \circ \alpha$  avec  $\alpha$  une isométrie linéaire de  $\mathbb{R}^n$ . Par le second point de ce même théorème, le choix d'une base orthonormale de  $\mathbb{R}^n$  permet d'identifier ces isométries linéaires avec le groupe orthogonal  $\text{O}(n)$ . Ainsi, toute isométrie  $f$  de  $\mathbb{R}^n$  s'écrit comme la composition d'un élément de  $\mathbb{R}^n$  et d'un élément de  $\text{O}(n)$ . Cette écriture est unique, par le même argument que dans la preuve de la proposition II.3. Finalement, l'égalité  $\alpha \circ \tau_v \circ \alpha^{-1} = \tau_{\alpha(v)}$ , démontrée dans cette même preuve, s'applique bien évidemment au cas particulier où  $\alpha$  est une isométrie linéaire. Cela démontre que  $\text{Isom}(\mathbb{R}^n)$  est bien égal au produit semi-direct  $\mathbb{R}^n \rtimes \text{O}(n)$ .

Pour le sous-groupe  $\text{Isom}^+(\mathbb{R}^n)$ , la preuve est exactement la même. Il s'agit juste de remarquer que – par définition – si l'isométrie  $f$  préserve l'orientation, alors la matrice orthogonale correspondante sera de déterminant  $+1$ , et donc un élément du sous-groupe  $\text{SO}(n)$  de  $\text{O}(n)$ .  $\square$

Concluons cette section par une conséquence facile de la proposition II.2.

**Proposition II.6.** *Si une isométrie  $f \in \text{Isom}(\mathbb{R}^n)$  fixe un ensemble non-vide de points  $F$ , alors  $f$  fixe le sous-espace affine engendré par  $F$ , c'est-à-dire, le plus petit sous-espace affine de  $\mathbb{R}^n$  contenant  $F$ .*

*Démonstration.* Comme  $f$  est une isométrie, c'est une application affine. Par la proposition II.2, son ensemble de points fixes  $\text{Fix}(f)$  est un sous-espace affine de  $\mathbb{R}^n$ . Ainsi  $\text{Fix}(f)$  est un sous-espace affine, contenant  $F$  par hypothèse. Par minimalité,  $\text{Fix}(f)$  contient donc le sous-espace affine engendré par  $F$ .  $\square$

Cette proposition signifie en particulier que si  $f$  fixe deux points  $x, y$  distincts, alors  $f$  fixe toute la droite  $\overline{xy}$ . De même, si  $f$  fixe trois points non-alignés  $x, y, z$ , alors  $f$  fixe le plan engendré par ces trois points.

**Corollaire II.7.** *Si  $f, g \in \text{Isom}(\mathbb{R}^n)$  coïncident sur un ensemble non-vide de points, alors  $f$  et  $g$  coïncident sur le sous-espace affine engendré par cet ensemble de points.*

*Démonstration.* On applique simplement la proposition II.6 à l'isométrie  $g^{-1} \circ f$ , dont les points fixes sont les points de coïncidence de  $f$  et  $g$ .  $\square$

En particulier, si  $f, g \in \text{Isom}(\mathbb{R}^2)$  coïncident sur trois points non-alignés, alors  $f = g$ . On utilisera ce résultat sous peu.

## II.3 Classification des isométries

Déterminer un groupe d'isométries est une chose, mais *classifier* ces isométries en est une autre. Dans cette section, on va donner une classification des isométries de la droite  $\mathbb{R}$ , puis du plan  $\mathbb{R}^2$ . Dans un dernier paragraphe, on expliquera ce qui peut être étendu au cas général de l'espace  $\mathbb{R}^n$ .

### II.3.1 Le cas de la droite

Pour s'échauffer, commençons par le cas facile des isométries de la droite  $\mathbb{R}$ . Étant donné un point  $x \in \mathbb{R}$ , notons  $\sigma_x$  la réflexion par le point  $x$ . Analytiquement, cette réflexion est donnée par la formule  $\sigma_x(y) = 2x - y$ .

#### Théorème II.8: Isométries de la droite

Toute isométrie de  $\mathbb{R}$  est composition d'au plus 2 réflexions. De plus, les éléments de  $\text{Isom}(\mathbb{R})$  sont précisément :

- l'identité (composition de 0 réflexion),
- les réflexions (composition de 1 réflexion), et
- les translations (composition de 2 réflexions).

*Démonstration.* Le point de départ de cette démonstration est l'observation faite ci-dessus : si  $f \in \text{Isom}(\mathbb{R})$  fixe deux points distincts, alors  $f = \text{id}$ .

Supposons donc que  $f$  possède un unique point fixe  $x$ . Nous allons maintenant vérifier que dans ce cas,  $f$  est forcément la réflexion  $\sigma_x$ . Soit en effet  $y \neq x$ , et notons  $d := d(x, y) > 0$ . Comme  $x$  est un point fixe et  $f$  est une isométrie,

$$d(x, f(y)) = d(f(x), f(y)) = d(x, y) = d.$$

Ainsi, le point  $f(y)$  est à distance  $d = d(x, y)$  de  $x$ . Or, l'ensemble des points à distance  $d$  de  $x$  est exactement  $\{y, \sigma_x(y)\}$ . Mais  $f(y)$  ne peut pas être égal à  $y$ , puisque  $x$  est l'unique point fixe de  $f$ . Ainsi,  $f(y) = \sigma_x(y)$ . Comme cette égalité vaut pour tout point  $y \neq x$  (et pour  $x$ ),  $f$  est bien la réflexion  $\sigma_x$ .

Supposons à présent que  $f$  ne possède aucun point fixe. Soit  $x \in \mathbb{R}$  quelconque, posons  $y = \frac{x+f(x)}{2}$  le point médian entre  $x$  et  $f(x)$ , et considérons la nouvelle isométrie  $g := \sigma_y \circ f$ . Par définition,  $g(x) = \sigma_y(f(x)) = x$ . Ainsi,  $g$  possède (au moins) un point fixe. Par le début de la preuve,  $g$  est soit l'identité, soit la réflexion  $\sigma_x$ . Le premier cas est exclu, puisqu'il impliquerait que  $f$  est égal à  $\sigma_y$  qui a un point fixe. Ainsi,  $g = \sigma_x$ , d'où  $f = \sigma_y \circ \sigma_x$ . Il ne reste plus qu'à vérifier que

la composition de deux réflexions (distinctes) est une translation (non-triviale). Cela découle des égalités

$$\sigma_y(\sigma_x(z)) = \sigma_y(2x - z) = 2y - (2x - z) = z + 2(y - x) = \tau_{2v}(z),$$

où  $v = y - x \in \mathbb{R}$ . □

Notons qu'il est également possible de donner une démonstration (plus simple) de cet énoncé en se basant sur le théorème II.5. C'est un exercice.

### II.3.2 Classification des isométries du plan

Passons à présent au cas du plan  $\mathbb{R}^2$ . Étant donnée une droite  $\ell \subset \mathbb{R}^2$ , on notera  $\sigma_\ell \in \text{Isom}(\mathbb{R}^2)$  la réflexion d'axe  $\ell$ . Elle est illustrée en figure II.2.

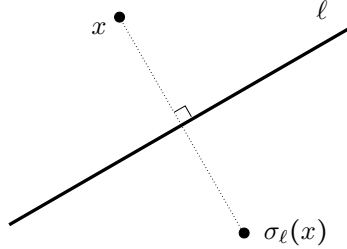


FIGURE II.2 – La réflexion  $\sigma_\ell$  d'axe  $\ell$ .

La première partie du théorème II.8 se généralise parfaitement.

#### Théorème II.9: Isométries du plan

Toute isométrie du plan est composition d'au plus 3 réflexions.

*Démonstration.* On va procéder par étapes, en se basant sur la proposition II.6. L'idée principale est qu'il est toujours possible d'augmenter la dimension de l'ensemble des points fixes d'une isométrie en la précomposant avec une réflexion bien choisie. (Notons que cette idée a déjà été utilisée dans la preuve du théorème II.8.)

(a) Commençons par le cas d'une isométrie  $f$  sans point fixe ; nous allons vérifier que l'on peut écrire  $f = \sigma_\ell \circ g$  avec  $\text{Fix}(g)$  non-vide. En effet, soit  $x \in \mathbb{R}^2$  quelconque, et soit  $\ell$  la médiatrice du segment  $[x, f(x)]$ . Par construction,  $\sigma_\ell(x)$  est égal à  $f(x)$ . Ainsi,  $g = \sigma_\ell \circ f \in \text{Isom}(\mathbb{R}^2)$  satisfait  $g(x) = \sigma_\ell(f(x)) = x$ . Nous avons donc bien que  $f = \sigma_\ell \circ g$  avec  $x \in \text{Fix}(g)$ .

(b) Supposons à présent que  $f \in \text{Isom}(\mathbb{R}^2)$  a un unique point fixe  $x$ , et vérifions que l'on peut écrire  $f = \sigma_\ell \circ g$  avec  $\text{Fix}(g)$  contenant une droite. En effet, soit  $y$  quelconque mais différent de  $x$ , et soit  $\ell$  la médiatrice du segment  $[y, f(y)]$ . Comme ci-dessus, l'isométrie  $g = \sigma_\ell \circ f$  satisfait  $g(y) = y$  par construction. De plus, puisque  $x$  est un point fixe de  $f$  qui est une isométrie,

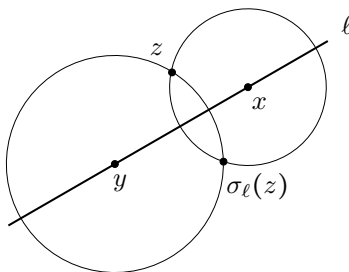
$$d(x, f(y)) = d(f(x), f(y)) = d(x, y).$$

Ainsi,  $x$  est un élément de l'ensemble des points à égale distance de  $y$  et de  $f(y)$ . Cet ensemble n'est autre que la droite  $\ell$ . Il suit

$$g(x) = \sigma_\ell(f(x)) = \sigma_\ell(x) = x.$$

L'isométrie  $g$  telle que  $f = \sigma_\ell \circ g$  fixe donc deux points distincts  $x, y$ . Par la proposition II.6, elle fixe toute la droite  $\overline{xy}$  engendrée par ces deux points.

(c) Supposons finalement que  $f \in \text{Isom}(\mathbb{R}^2)$  fixe une droite  $\ell$ ; nous allons vérifier qu'alors, soit  $f = \sigma_\ell$ , soit  $f = \text{id}$ . En effet, prenons deux points distincts  $x, y \in \ell$ , ainsi que  $z \notin \ell$ . Comme plus haut, le fait que  $x, y$  sont des points fixes de  $f$  qui est une isométrie implique  $d(x, f(z)) = d(x, z)$  et  $d(y, f(z)) = d(y, z)$ . Cela signifie que  $f(z)$  se situe à l'intersection des deux cercles illustrés ci-dessous. En d'autres termes,  $f(z)$  appartient à l'ensemble  $\{z, \sigma_\ell(z)\}$ .



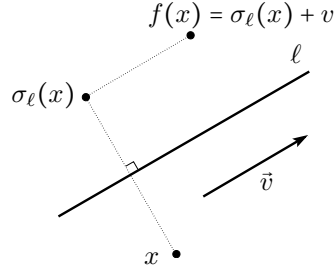
Si  $f(z) = z$ , alors  $f$  fixe trois points  $x, y, z$  non-alignés. Par la proposition II.6,  $f$  est l'identité. Si  $f(z) = \sigma_\ell(z)$ , alors  $f$  et  $\sigma_\ell$  coïncident sur les trois points  $x, y, z$  non-alignés. Par le corollaire II.7, ces deux isométries coïncident partout.

Les points (a), (b) et (c) ci-dessus peuvent se résumer comme suit : si  $f \in \text{Isom}(\mathbb{R}^2)$  est telle que  $\text{Fix}(f)$  a dimension  $i < 2$ , alors on peut trouver une droite  $\ell$  telle que  $f = \sigma_\ell \circ g$  avec  $\text{Fix}(g)$  de dimension strictement supérieure à  $i$ . (On adopte ici la convention que la dimension de l'ensemble vide vaut  $-1$ .) Ainsi, dans le pire des cas,  $f$  n'a pas de point fixe et il faut la précomposer avec 3 réflexions pour trouver une isométrie avec ensemble de points fixes de dimension 2, c'est-à-dire avant de trouver l'identité. Cela démontre le théorème.  $\square$

Ce résultat peut maintenant être utilisé pour classifier les isométries du plan. Voici une liste de telles isométries, dont les premiers exemples ont déjà été vus.

- La *translation* par un vecteur  $v \in \mathbb{R}^2$ , notée  $\tau_v$  ;
- La *réflexion* d'axe  $\ell$ , notée  $\sigma_\ell$ , illustrée en figure II.2.
- Pour un point  $O \in \mathbb{R}^2$  et un réel  $\alpha \geq 0$ , on notera  $\rho_O^\alpha$  la *rotation* d'angle  $\alpha$  autour du point  $O$  dans le sens trigonométrique.
- Une *réflexion glissée* est une réflexion composée avec une translation parallèle à l'axe de réflexion :  $f = \tau_v \circ \sigma_\ell$  avec  $\vec{v}$  parallèle à  $\ell$ . Cette isométrie est illustrée en figure II.3.

Notons qu'il est bien entendu possible de donner des formules analytiques explicites pour ces isométries, mais nous allons utiliser une approche synthétique.

FIGURE II.3 – Une réflexion glissée  $f = \tau_v \circ \sigma_\ell$ .**Théorème II.10: Classification des isométries du plan**

Les isométries du plan sont précisément :

- l'identité (composition de 0 réflexion),
- les réflexions (composition de 1 réflexion),
- les translations (composition de 2 réflexions d'axes parallèles),
- les rotations (composition de 2 réflexions d'axes non-parallèles), et
- les réflexions glissées (composition de 3 réflexions).

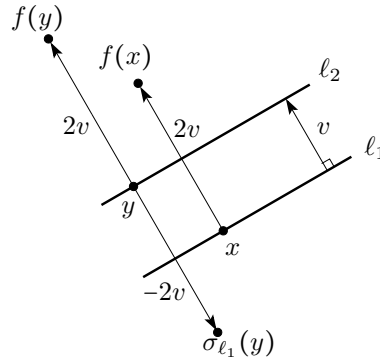
*Démonstration.* Par le théorème II.9, on sait qu'un élément  $f \in \text{Isom}(\mathbb{R}^2)$  est composition de 0, 1, 2 ou 3 réflexions. Trivialement, si  $f$  est composition de 0 (resp. 1) réflexion, alors  $f$  est l'identité (resp. une réflexion).

Supposons à présent que  $f = \sigma_{\ell_2} \circ \sigma_{\ell_1}$  avec  $\ell_1, \ell_2$  deux droites (distinctes) du plan. Il faut distinguer deux cas, selon si ces droites sont parallèles ou non. Nous allons maintenant démontrer que dans le premier cas,  $f$  n'est autre que la translation  $\tau_{2v}$ , où  $v \in \mathbb{R}^2$  est le vecteur normal à  $\ell_1$  tel que  $\ell_1 + v = \ell_2$ . Notons qu'en vertu du corollaire II.7, il suffit de vérifier que  $f = \sigma_{\ell_2} \circ \sigma_{\ell_1}$  et  $\tau_{2v}$  coïncident sur 3 points non-alignés. Ainsi, il suffit de vérifier que  $f$  et  $\tau_{2v}$  coïncident sur  $\ell_1 \cup \ell_2$ . Cela découle des égalités suivantes, illustrées plus bas : pour  $x \in \ell_1$ ,

$$f(x) = \sigma_{\ell_2}(\sigma_{\ell_1}(x)) = \sigma_{\ell_2}(x) = x + 2v = \tau_{2v}(x),$$

tandis que pour  $y \in \ell_2$ ,

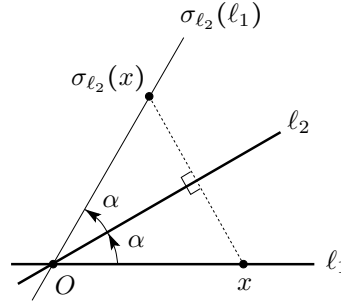
$$f(y) = \sigma_{\ell_2}(\sigma_{\ell_1}(y)) = \sigma_{\ell_2}(y - 2v) = y + 2v = \tau_{2v}(y).$$



Supposons à présent que  $f = \sigma_{\ell_2} \circ \sigma_{\ell_1}$  avec  $\ell_1, \ell_2$  non-parallèles. Nous allons montrer que dans ce cas,  $f$  n'est autre que la rotation  $\rho_O^{2\alpha}$ , où  $O$  est l'intersection de  $\ell_1$  et de  $\ell_2$  et  $\alpha > 0$  l'angle (orienté dans le sens trigonométrique) entre  $\ell_1$  et  $\ell_2$ . Comme avant, il suffit de vérifier que  $f$  et  $\rho_O^{2\alpha}$  coïncident sur  $\ell_1 \cup \ell_2$ . Pour  $x \in \ell_1$ , cela découle des égalités suivantes :

$$f(x) = \sigma_{\ell_2}(\sigma_{\ell_1}(x)) = \sigma_{\ell_2}(x) = \rho_O^{2\alpha}(x),$$

où la dernière égalité est illustrée ci-dessous.



Pour  $y \in \ell_2$ , cela découle des égalités suivantes qui se vérifient de manière similaire :

$$f(y) = \sigma_{\ell_2}(\sigma_{\ell_1}(y)) = \sigma_{\ell_2}(\rho_O^{-2\alpha}(y)) = \rho_O^{2\alpha}(y).$$

Ainsi, il ne reste plus qu'à démontrer que si  $f$  est composition d'exactly 3 réflexions, alors  $f$  est une réflexion glissée. Soit donc  $f = \sigma_{\ell_3} \circ \sigma_{\ell_2} \circ \sigma_{\ell_1}$  avec  $\ell_1, \ell_2, \ell_3$  trois droites du plan. Nous allons d'abord montrer que  $f$  peut s'écrire  $f = \tau_v \circ \sigma_\ell$ . En effet, c'est clairement le cas si  $\ell_2$  et  $\ell_3$  sont parallèles. Si ces deux droites ne sont pas parallèles, alors  $f = \rho_O^{2\alpha} \circ \sigma_{\ell_1}$  par le point précédent. Mais la rotation  $\rho_O^{2\alpha} = \sigma_{\ell_3} \circ \sigma_{\ell_2}$  peut aussi s'écrire  $\rho_O^{2\alpha} = \sigma_{\ell'_3} \circ \sigma_{\ell'_2}$ , où  $\ell'_2$  et  $\ell'_3$  sont obtenues à partir de  $\ell_2$  et  $\ell_3$  via une rotation centrée en  $O \in \ell_2 \cap \ell_3$  d'angle tel que  $\ell'_2$  est parallèle à  $\ell_1$ . On obtient alors

$$f = \rho_O^{2\alpha} \circ \sigma_{\ell_1} = \sigma_{\ell'_3} \circ \sigma_{\ell'_2} \circ \sigma_{\ell_1} = \sigma_\ell \circ \tau_{v'},$$

où  $\ell = \ell'_3$  et  $v' = \frac{1}{2}(\ell'_2 - \ell_1)$ . De plus, on peut vérifier que  $\sigma_\ell \circ \tau_{v'} = \tau_v \circ \sigma_\ell$  avec  $v = \sigma_\ell(v' + p) - p \in \ell$ . Ainsi,  $f$  est bien de la forme  $\tau_v \circ \sigma_\ell$ , comme annoncé.

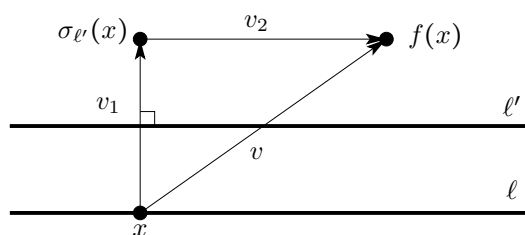
La dernière étape consiste à vérifier que toute application de la forme  $f = \tau_v \circ \sigma_\ell$  est une réflexion glissée. Notons  $v_1$  (resp.  $v_2$ ) la composante de  $v$  normale à  $\ell$  (resp. parallèle à  $\ell$ ), et soit  $\ell' = \ell + \frac{1}{2}v_1$ , comme illustré ci-dessous. Nous allons montrer que  $f$  n'est autre que la réflexion glissée  $\tau_{v_2} \circ \sigma_{\ell'}$ , ce qui terminera la démonstration. Comme avant, il suffit de vérifier que  $f = \tau_v \circ \sigma_\ell$  et  $\tau_{v_2} \circ \sigma_{\ell'}$  coïncident sur  $\ell \cup \ell'$ . Pour  $x \in \ell$ , cela découle des égalités suivantes :

$$f(x) = \tau_v(\sigma_\ell(x)) = \tau_v(x) = x + v = \tau_{v_2}(x + v_1) = \tau_{v_2}(\sigma_{\ell'}(x)).$$

Le dernière égalité est illustrée ci-dessous. Pour  $y \in \ell'$ , cela découle des égalités suivantes qui se vérifient de manière similaire :

$$f(y) = \tau_v(\sigma_\ell(y)) = \sigma_\ell(y) + v = \tau_{v_2}(y) = \tau_{v_2}(\sigma_{\ell'}(y)).$$

Cela conclut la démonstration. □



Notons la conséquence suivante de ce théorème : les isométries linéaires du plan sont l'identité, les réflexions par les droites contenant l'origine, et les rotations autour de l'origine. En particulier, on retrouve le fait connu que les éléments de  $\text{SO}(2)$  sont les rotations autour de l'origine, mais également le fait nouveau que les éléments de  $\text{O}(2) \setminus \text{SO}(2)$  sont exactement les réflexions par les droites contenant l'origine.

### II.3.3 Le cas général de l'espace $\mathbb{R}^n$

Dans les deux paragraphes précédents, nous avons classifié les isométries de la droite  $\mathbb{R}$  et du plan  $\mathbb{R}^2$ . Dans quelle mesure est-il possible d'étendre ces résultats au cas général de  $\mathbb{R}^n$  avec  $n$  fixé mais arbitraire ? Il s'avère que la première partie de la démonstration (dans le cas du plan, le théorème II.9) se généralise parfaitement. C'est le sujet de ce dernier paragraphe, qui ne fait pas partie du champ de l'examen.

On aura besoin de la terminologie suivante : un sous-espace affine  $H \subset \mathbb{R}^n$  de dimension  $n - 1$  est appelé un *hyperplan*. Étant donné un hyperplan  $H$ , on notera  $\sigma_H$  la *réflexion* d'hyperplan  $H$  définie comme suit : l'image de  $x \in \mathbb{R}^n$  est l'unique point  $\sigma_H(x) \in \mathbb{R}^n$  tel que le vecteur  $x - \sigma_H(x)$  est orthogonal à  $H$  et le segment  $[x, \sigma_H(x)]$  est bissecté par  $H$ . Le fait que  $\sigma_H$  est une isométrie de  $\mathbb{R}^n$  se démontre facilement.

Dans les cas  $n = 1$  et  $n = 2$ , on reconnaît les réflexions considérées aux paragraphes précédents. Ils sont illustrés en figure II.4, ainsi que le cas  $n = 3$ .

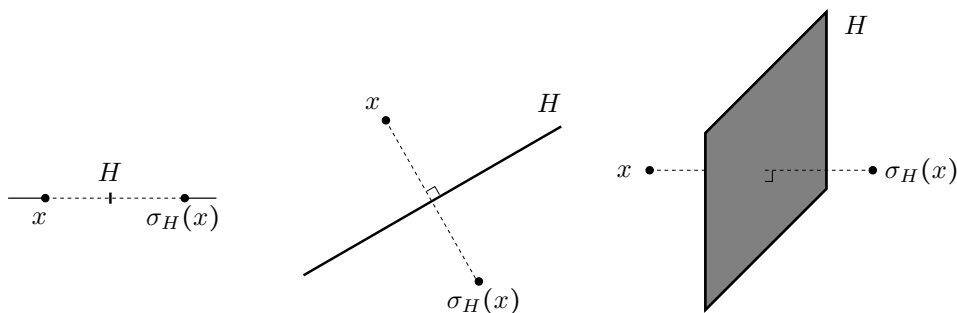


FIGURE II.4 – Réflexions hyperplanes en dimension  $n = 1$ ,  $n = 2$  et  $n = 3$ .

Voici le résultat principal de ce paragraphe.

**Théorème II.11: Isométries de  $\mathbb{R}^n$** 

Toute isométrie de  $\mathbb{R}^n$  est composition d'au plus  $n+1$  réflexions hyperplanes. De plus, une isométrie de  $\mathbb{R}^n$  préserve l'orientation si et seulement si elle est composition d'un nombre pair de réflexions.

Cette démonstration étant relativement conséquente, nous allons la diviser en plusieurs parties. Plus précisément, nous allons d'abord démontrer deux lemmes.

**Lemme II.12.** *Étant donnés  $x_0$  et  $y_0$  deux points distincts de  $\mathbb{R}^n$ , soit  $H \subset \mathbb{R}^n$  défini par*

$$H = \{x \in \mathbb{R}^n \mid d(x, x_0) = d(x, y_0)\}.$$

*Alors,  $H$  est l'hyperplan orthogonal au vecteur  $v = y_0 - x_0$  passant par le milieu du segment  $[x_0, y_0]$ .*

*Démonstration.* En utilisant la formule

$$d(x, x_0)^2 = \langle x - x_0, x - x_0 \rangle = \|x\|^2 + \|x_0\|^2 - 2\langle x, x_0 \rangle,$$

on voit que  $x$  est un élément de  $H$  si et seulement si  $x$  satisfait l'équation

$$\langle x, \overbrace{y_0 - x_0}^v \rangle = \frac{1}{2}(\|y_0\|^2 - \|x_0\|^2) =: \lambda.$$

En d'autres termes,  $H = \beta^{-1}(\lambda)$  où  $\beta: \mathbb{R}^n \rightarrow \mathbb{R}$  est l'application linéaire (surjective) donnée par  $\beta(x) = \langle x, v \rangle$ . Par l'argument de la preuve de la proposition II.2,  $H$  est un translaté de  $\ker(\beta)$ , le sous-espace vectoriel de dimension  $n-1$  donné par tous les vecteurs orthogonaux à  $v$ . Par définition,  $H$  passe par le milieu du segment  $[x_0, y_0]$ .  $\square$

Pour le lemme suivant nous avons besoin d'une convention : posons que la dimension de l'ensemble vide est égale à  $-1$ .

**Lemme II.13.** *Si  $f \in \text{Isom}(\mathbb{R}^n)$  satisfait  $\dim \text{Fix}(f) \geq n-i$ , alors on peut écrire  $f$  comme composition d'au plus  $i$  réflexions hyperplanes.*

*Démonstration.* La preuve se fait par récurrence sur  $i \geq 0$ . Le cas  $i = 0$  est une tautologie :  $f \in \text{Isom}(\mathbb{R}^n)$  satisfait  $\dim \text{Fix}(f) \geq n$  si et seulement si  $f$  est l'identité, auquel cas  $f$  est composition de  $i = 0$  réflexion.

Supposons donc l'énoncé démontré jusqu'à  $i \geq 0$ , et soit  $f \in \text{Isom}(\mathbb{R}^n)$  avec  $\dim \text{Fix}(f) = n - (i+1)$ . Nous devons vérifier que  $f$  est composition d'au plus  $i+1$  réflexions. Comme  $i \geq 0$ ,  $\dim \text{Fix}(f) = n - (i+1) < n$ ; ainsi, il existe un point  $x_0$  tel que  $f(x_0) \neq x_0$ . Soit

$$H := \{x \in \mathbb{R}^n \mid d(x, x_0) = d(x, f(x_0))\} \subset \mathbb{R}^n$$

l'hyperplan médian entre  $x_0$  et  $f(x_0)$ . Par le lemme II.12, il s'agit aussi de l'hyperplan orthogonal au vecteur  $v = f(x_0) - x_0$  passant par le milieu du segment  $[x_0, f(x_0)]$ . En particulier, si l'on note  $g := \sigma_H \circ f \in \text{Isom}(\mathbb{R}^n)$ , on a

$g(x_0) = \sigma_H(f(x_0)) = x_0$  par définition de  $\sigma_H$ . Ainsi,  $x_0$  est un élément de  $\text{Fix}(g) \setminus \text{Fix}(f)$ . De plus,  $\text{Fix}(g)$  contient  $\text{Fix}(f)$ . En effet, si  $x$  est un point fixe de  $f$ , alors

$$d(x, x_0) = d(f(x), f(x_0)) = d(x, f(x_0)).$$

Par définition de  $H$ , cela signifie que  $x$  est élément de  $H$ . Il suit que  $\sigma_H(x) = x$ , et donc, que

$$g(x) = \sigma_H(f(x)) = \sigma_H(x) = x.$$

Ainsi,  $x$  est bien un point fixe de  $g$ . Mais puisque  $\text{Fix}(f) \subset \text{Fix}(g)$  et  $x_0$  est un élément de  $\text{Fix}(g) \setminus \text{Fix}(f)$ , la dimension de  $\text{Fix}(g)$  est strictement supérieure à celle de  $\text{Fix}(f)$ . En d'autres termes,  $\dim \text{Fix}(g) \geq n - i$ . Par hypothèse de récurrence,  $g$  est composition d'au plus  $i$  réflexions. On en conclut que  $f = \sigma_H \circ g$  est composition d'au plus  $i + 1$  réflexions, ce qu'il fallait démontrer.  $\square$

*Démonstration du théorème II.11.* Notons tout d'abord que le lemme II.13 implique immédiatement la première partie de l'énoncé, puisque toute isométrie  $f \in \text{Isom}(\mathbb{R}^n)$  satisfait  $\dim \text{Fix}(f) \geq -1 = n - (n + 1)$ .

Pour la seconde assertion, considérons l'homomorphisme  $\varepsilon: \text{Isom}(\mathbb{R}^n) \rightarrow \{\pm 1\}$  défini par la composition des homomorphismes suivants

$$\text{Isom}(\mathbb{R}^n) = \mathbb{R}^n \rtimes \text{O}(n) \longrightarrow \text{O}(n) \xrightarrow{\det} \{\pm 1\}.$$

En d'autres termes,  $\varepsilon$  est défini par  $\varepsilon(\tau_v \circ \alpha) = \det(\alpha)$ . Par définition, son noyau est donné par  $\ker(\varepsilon) = \text{Isom}^+(\mathbb{R}^n) = \mathbb{R}^n \rtimes \text{SO}(n)$ , le sous-groupe des isométries préservant l'orientation. Comme  $\varepsilon$  est un homomorphisme, il ne reste plus qu'à vérifier que  $\varepsilon(\sigma_H) = -1$  pour tout hyperplan  $H \subset \mathbb{R}^n$ . Pour ce faire, remarquons d'abord l'égalité

$$\tau_v \circ \sigma_H \circ \tau_v^{-1} = \sigma_{\tau_v(H)},$$

que l'on démontre facilement. Comme  $\varepsilon$  est un homomorphisme,

$$\varepsilon(\sigma_{\tau_v(H)}) = \varepsilon(\tau_v \circ \sigma_H \circ \tau_v^{-1}) = \varepsilon(\tau_v) \varepsilon(\sigma_H) \varepsilon(\tau_v)^{-1} = \varepsilon(\sigma_H).$$

Ainsi, on peut supposer sans restreindre la généralité que  $H$  contient l'origine, i.e. que  $\sigma_H$  est linéaire, auquel cas  $\varepsilon(\sigma_H) = \det(\sigma_H)$ . Choisissons une base orthonormale  $e_1, \dots, e_{n-1}$  de  $H \subset \mathbb{R}^n$ , et complétons-la en une base orthonormale de  $\mathbb{R}^n$  avec  $e_n$ , un vecteur normal à  $H$ . En utilisant la matrice de  $\sigma_H$  dans cette base, on obtient

$$\varepsilon(\sigma_H) = \det(\sigma_H) = \begin{vmatrix} 1 & & & \\ & \ddots & & \\ & & 1 & \\ & & & -1 \end{vmatrix} = -1,$$

ce qui conclut la démonstration.  $\square$

Pour ce qui est de la classification des isométries de  $\mathbb{R}^n$ , la situation est la suivante : si l'on se donne un  $n$  fixé, il est théoriquement possible de classer les isométries en considérant toutes les combinaisons possibles d'au plus  $n + 1$  hyperplans dans  $\mathbb{R}^n$ . Dans le cas  $n = 1$ , c'est extrêmement facile, on l'a vu dans la preuve du théorème II.8. Dans le cas  $n = 2$ , cela demande un peu de travail : c'était la preuve du théorème II.10. Pour  $n = 3$ , c'est encore relativement facile. Mais plus  $n$  devient grand, plus le nombre de configurations possibles d'hyperplans grandit, et il est illusoire d'espérer donner une classification complète des isométries de l'espace  $\mathbb{R}^n$  pour un  $n$  élevé.

## II.4 Groupes de symétries

S'il n'y avait qu'une seule phrase à se rappeler de ces deux premiers chapitres, cela serait sans doute celle-ci, due à M.A. Armstrong :

*Numbers measure size, groups measure symmetry.*

La formalisation de cette idée se fait au moyen de la notion de *groupes de symétries*. Le but de cette dernière section est d'utiliser les résultats du chapitre I pour déterminer les groupes de symétries d'objets géométriques, en particulier des polyèdres réguliers. On déterminera ensuite entièrement tous les groupes finis qui peuvent être réalisés comme groupes de symétries d'objets dans  $\mathbb{R}^n$  pour  $n = 1, 2, 3$ .

### II.4.1 Calcul de groupes de symétries

Voici la définition principale de cette section.

#### Définition II.5

Le **groupe de symétries** d'un sous-ensemble  $Y \subset \mathbb{R}^n$  est

$$\text{Sym}(Y) = \{f \in \text{Isom}(\mathbb{R}^n) \mid f(Y) = Y\}.$$

De même, on définit le **groupe de symétries propres** de  $Y$  comme

$$\text{Sym}^+(Y) = \{f \in \text{Isom}^+(\mathbb{R}^n) \mid f(Y) = Y\}.$$

On vérifie immédiatement qu'il s'agit bien de groupes, qui agissent naturellement sur  $Y$ . Par définition, on a de plus

$$\text{Sym}^+(Y) < \text{Sym}(Y) < \text{Isom}(\mathbb{R}^n) \quad \text{et} \quad \text{Sym}^+(Y) < \text{Isom}^+(\mathbb{R}^n).$$

Comme le dit la phrase de Armstrong, ces groupes permettent de mesurer combien un objet géométrique est *symétrique*. Donnons quelques exemples faciles, avant de passer à des exemples plus intéressants.

**Exemples** (Premiers exemples de groupes de symétries).

1. Un objet “pris au hasard”, tel le sous-ensemble du plan  $Y = \infty$ , n'est pas du tout symétrique :  $\text{Sym}^+(Y) = \text{Sym}(Y) = \{\text{id}\}$ .
2. On vérifie facilement que le sous-ensemble du plan donné par la “croix”  $Y = \dagger$  satisfait  $\text{Sym}^+(Y) = \{\text{id}\} < C_2 = \text{Sym}(Y)$ .
3. Soit  $Y = \mathbb{S}^1$ , le sous-ensemble du plan donné par les points à distance 1 de l'origine. On vérifie que  $\text{Sym}(\mathbb{S}^1) = \text{O}(2)$  et  $\text{Sym}^+(\mathbb{S}^1) = \text{SO}(2)$ . Ainsi, le cercle  $\mathbb{S}^1$  possède un nombre infini de symétries.

Le calcul de groupes de symétries plus compliqués nécessite un résultat intermédiaire, et la terminologie suivante. Le *barycentre* (ou *centre de masse*) d'un sous-ensemble fini  $X \subset \mathbb{R}^n$  est l'élément de  $\mathbb{R}^n$  défini par

$$B_X = \frac{1}{|X|} \sum_{x \in X} x.$$

Voici le résultat en question.

**Lemme II.14.** *Pour tout sous-ensemble  $X \subset \mathbb{R}^n$  fini et toute application affine bijective  $f \in \text{Aff}(\mathbb{R}^n)$ , on a  $f(B_X) = B_{f(X)}$ .*

*Démonstration.* Par définition,  $f$  peut s'écrire  $\tau_v \circ \alpha$  avec  $v \in \mathbb{R}^n$  et  $\alpha \in \text{GL}(n, \mathbb{R})$ . Ainsi, il suffit de vérifier l'égalité  $f(B_X) = B_{f(X)}$  pour  $f = \tau_v$  et  $f = \alpha \in \text{GL}(n, \mathbb{R})$ . Cela découle des égalités

$$B_{\tau_v(X)} = \frac{1}{|\tau_v(X)|} \sum_{y \in \tau_v(X)} y = \frac{1}{|X|} \sum_{x \in X} (x + v) = B_X + \frac{1}{|X|} \sum_{x \in X} v = B_X + v = \tau_v(B_X)$$

et

$$B_{\alpha(X)} = \frac{1}{|\alpha(X)|} \sum_{y \in \alpha(X)} y = \frac{1}{|X|} \sum_{x \in X} \alpha(x) = \alpha\left(\frac{1}{|X|} \sum_{x \in X} x\right) = \alpha(B_X). \quad \square$$

Nous pouvons maintenant commencer à calculer les groupes de symétries d'objets plus intéressants. Commençons par revisiter un exemple déjà traité (de manière quelque peu lacunaire) en section I.1.

### Le $n$ -gone régulier dans $\mathbb{R}^2$

Pour  $n \geq 2$ , soit  $P_n \subset \mathbb{R}^2$  le  $n$ -gone régulier avec sommets  $X = \mu_n(\mathbb{C})$ . (Dans le cas  $n = 2$ , il s'agit du segment  $P_2 = [-1, 1]$ ). Comme  $B_X = 0$  et tout élément  $f \in \text{Sym}(P_n)$  satisfait  $f(X) = X$ , le lemme II.14 implique que tout élément de  $\text{Sym}(P_n)$  fixe l'origine. Notons que cela comble une lacune de l'exemple 9 de la section I.1.

Comptons maintenant l'ordre de ce groupe au moyen de la formule des orbites. Comme  $G = \text{Sym}(P_n)$  agit transitivement sur  $X$ , on a  $|G \cdot x| = |X| = n$  pour  $x = 1 \in X$ . De plus, les éléments de  $\text{Stab}(x)$  sont des isométries qui fixent l'origine et  $x = 1$ . Par la proposition II.6, ces éléments fixent la droite réelle. Par cette même proposition, si  $f \in \text{Stab}(x)$  fixe un point hors de la droite réelle, alors

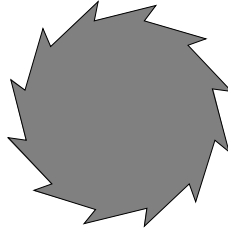
$f = \text{id}$ . Il suit que  $\text{Stab}(x) = \{\text{id}, s\}$ , où  $s = \sigma_{\mathbb{R}}$  la conjugaison complexe. Par la formule des orbites, l'ordre de  $\text{Sym}(P_n)$  est donc  $2n$ . Par les mêmes arguments, l'ordre de  $\text{Sym}^+(P_n)$  est  $n$ .

Mais comme  $\text{Sym}^+(P_n)$  contient  $C_n = \{r^\ell\}$  engendré par  $r = \rho_0^{2\pi/n}$ , qui est d'ordre  $n$ , on a donc  $\text{Sym}^+(P_n) = C_n$ . Finalement, si l'on note  $C_2 = \{\text{id}, s\}$ , on vérifie comme en section II.2.1 que  $\text{Sym}(P_n) = C_n \rtimes C_2$ , où l'action de  $C_2$  sur  $C_n$  est donnée par  $sr^\ell s = r^{-\ell}$ . En résumé,

$$\text{Sym}^+(P_n) = C_n \quad \text{et} \quad \text{Sym}(P_n) = C_n \rtimes C_2 = D_{2n}.$$

### La scie circulaire dans $\mathbb{R}^2$

Voici une variation du calcul précédent : pour  $n \geq 2$ , soit  $Y_n \subset \mathbb{R}^2$  obtenu en recollant  $n$  triangles identiques mais non-isocèles aux côtés du  $n$ -gone régulier  $P_n$ . En voici un exemple qui explique la terminologie.



On calcule facilement  $\text{Sym}^+(Y_n) = \text{Sym}(Y_n) = C_n$  : c'est un exercice.

### Le cube dans $\mathbb{R}^3$

Soit  $C$  un cube dans  $\mathbb{R}^3$  ; tentons de déterminer les groupes  $G = \text{Sym}(C)$  et  $G^+ = \text{Sym}^+(C)$  en suivant la même méthode que ci-dessus.

Soit  $X$  l'ensemble des 8 sommets du cube. Le groupe  $G$  agit sur  $X$  de manière transitive : par exemple, les rotations autour des axes passants par le centre de faces opposées du cube suffisent à envoyer un sommet quelconque sur n'importe quel autre sommet. Par la formule des orbites appliquée à un  $x \in X$  quelconque mais fixé,

$$|G| = |G \cdot x| |\text{Stab}(x)| = 8 |\text{Stab}(x)|.$$

Il s'agit à présent de calculer l'ordre du groupe  $\text{Stab}(x)$ . Nous allons à nouveau appliquer la formule des orbites, cette fois-ci au groupe  $\text{Stab}(x)$  agissant sur l'ensemble  $A_x$  des arêtes du cube adjacentes à notre sommet fixé  $x$ . Clairement,  $\text{Stab}(x)$  agit transitivement sur l'ensemble  $A_x$  qui compte 3 éléments, d'où

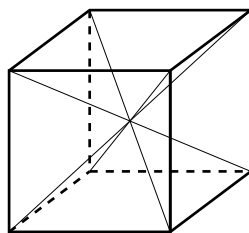
$$|\text{Stab}(x)| = |\text{Stab}(x) \cdot a| |\text{Stab}(a)| = 3 |\text{Stab}(a)|$$

pour une arête  $a$  adjacente à  $x$ . Reste à déterminer l'ordre de

$$\text{Stab}(a) = \{f \in \text{Stab}(x) \mid f(a) = a\} = \{f \in G \mid f(x) = x, f(a) = a\}.$$

Mais rappelons que par le lemme II.14, tout élément de  $G = \text{Sym}(C)$  fixe le barycentre  $B$  de  $X$ . Ainsi, un élément  $f$  de  $\text{Stab}(a)$  fixe le sommet  $x$  et le barycentre  $B$ , et envoie l'arête  $a$  sur elle-même. Par la proposition II.6,  $f$  fixe le plan  $\pi$  engendré par l'arête  $a$  et le barycentre. Par les arguments de la section II.3.2, une telle isométrie est soit l'identité, soit la réflexion  $\sigma_\pi$ . Ainsi,  $\text{Stab}(a)$  compte exactement 2 éléments. Il suit des égalités ci-dessus l'ordre de  $G = \text{Sym}(C)$  est donné par  $|G| = 8 \cdot 3 \cdot 2 = 48$ .

L'ordre de  $G^+$  se calcule exactement de la même manière. La seule différence est que, puisque  $\sigma_\pi$  ne préserve pas l'orientation,  $\text{Stab}(a)$  est trivial. Nous obtenons donc  $|G^+| = 8 \cdot 3 \cdot 1 = 24$ .



Il s'agit maintenant de déterminer ces groupes. Commençons par  $G^+$ . Considérons l'action de  $G^+$  sur l'ensemble des 4 diagonales du cube illustrées ci-dessus. Comme on l'a vu en section I.3, cela définit un homomorphisme de groupes

$$\phi: G^+ \longrightarrow S_4.$$

Soit  $f \in \ker(\phi)$ ; cela signifie que  $f$  est une isométrie de  $\mathbb{R}^3$  qui préserve l'orientation, fixe le barycentre  $B$  (il s'agit donc d'une rotation autour d'un axe passant par  $B$ ), envoie le cube sur lui-même, et chacune des 4 diagonales du cube sur elle-même. Une telle isométrie est forcément l'identité. Ainsi, le noyau de  $\phi$  est trivial, et  $\phi$  est donc injectif. Comme  $|G^+| = 24 = |S_4|$ ,  $\phi$  est un homomorphisme bijectif, donc un isomorphisme. Le groupe des symétries propres du cube est donc isomorphe au groupe symétrique  $S_4$ .

Pour déterminer le groupe  $G$  de toutes les symétries du cube, considérons les sous-groupes  $G^+$  et  $C_2 = \{\text{id}, c\}$ , où  $c \in G$  est la symétrie centrale de centre  $B$ . (Si le cube est centré en  $B = 0$ ,  $c$  est simplement donné par la multiplication par  $-1$ .) Par Lagrange et les calculs ci-dessus, on a  $[G : G^+] = \frac{|G|}{|G^+|} = \frac{48}{24} = 2$ . De plus, comme  $c$  est composition de 3 réflexions, il ne préserve pas l'orientation, et l'on a donc  $c \in G \setminus G^+$ . Il suit des arguments de la section II.2.1 que  $G$  est le produit semi-direct  $G^+ \rtimes C_2$ . De plus, l'élément  $c$  commute avec tous les éléments de  $G^+$  (penser à la multiplication par  $-1$  qui commute avec les rotations autour d'axes passant par l'origine) : il s'agit donc d'un produit direct  $G = G^+ \times C_2$ .

En conclusion, on a donc

$$\text{Sym}^+(C) \simeq S_4 \quad \text{et} \quad \text{Sym}(C) \simeq S_4 \times C_2.$$

On verra en exercice comment utiliser le résultat  $|\text{Sym}^+(C)| = 24$  pour compter le nombre de coloriations des faces d'un cube avec  $k$  couleurs. Notons qu'il faut bel et bien se restreindre au groupe des symétries *propres*, puisque les éléments

de  $\text{Sym}(C)$  qui renversent l'orientation ne correspondent pas à des transformations "physiques" du cube  $C \subset \mathbb{R}^3$ .

### Le tétraèdre dans $\mathbb{R}^3$

Le cas du tétraèdre  $T$  dans  $\mathbb{R}^3$  est bien plus facile, et sera vu en exercice. On démontre que

$$\text{Sym}^+(T) \simeq A_4 \quad \text{et} \quad \text{Sym}(T) \simeq S_4.$$

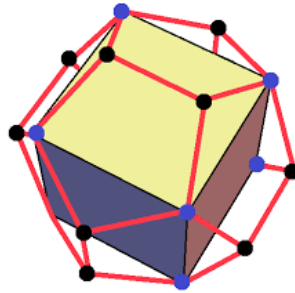
On verra aussi en exercice comment utiliser le résultat  $|\text{Sym}^+(T)| = 12$  pour compter le nombre de coloriage des faces d'un tétraèdre avec  $k$  couleurs.

### Le dodécaèdre dans $\mathbb{R}^3$

La détermination du groupe de symétries du dodécaèdre  $D \subset \mathbb{R}^3$  suit exactement les mêmes lignes, mais est plus difficile à visualiser sans modèle du dodécaèdre en main. Nous allons nous contenter d'en donner les étapes principales en laissant de côté quelques détails.

Comme  $G := \text{Sym}(D)$  agit transitivement sur l'ensemble  $X$  des 20 sommets de  $D$ , la formule des orbites implique que  $|G| = 20|\text{Stab}(x)|$  pour tout  $x \in X$ . De même,  $\text{Stab}(x)$  agit transitivement sur l'ensemble  $A_x$  des 3 arêtes de  $D$  adjacentes à  $x$ , d'où  $|\text{Stab}(x)| = 3|\text{Stab}(a)|$  pour tout  $a \in A_x$ . Exactement comme dans le cas du cube,  $\text{Stab}(a) = \{\text{id}, \sigma_\pi\}$ , où  $\pi$  est le plan engendré par  $a$  et le barycentre de  $X$ . Ainsi, on obtient  $|G| = 120$ . Similairement, on vérifie que l'ordre de  $G^+ = \text{Sym}^+(D)$  est  $|G^+| = 60$ .

Il s'agit maintenant de comprendre que le dodécaèdre  $D$  possède exactement 5 cubes inscrits, qui correspondent aux 5 diagonales d'une face fixée de  $D$ . Un de ces cubes est illustré ci-dessous.



L'action de  $\text{Sym}^+(D)$  sur cet ensemble à 5 éléments définit un homomorphisme

$$\phi: \text{Sym}^+(D) \longrightarrow S_5$$

dont le noyau est trivial. On vérifie ensuite que l'image de cet homomorphisme contient tous les 3-cycles de  $S_5$ . Comme ces 3-cycles engendrent  $A_5$ , l'image de  $\phi$  contient  $A_5$ . Mais l'ordre de l'image de  $\phi$  est égal à  $|\text{Sym}^+(D)| = 60 = |A_5|$ . Ainsi,  $\phi$  définit un isomorphisme entre  $\text{Sym}^+(D)$  et  $A_5$ . Exactement comme pour le cube, on vérifie enfin que  $G = G^+ \times C_2$ , où  $C_2 = \{\text{id}, c\}$ . On a donc

$$\text{Sym}^+(D) \simeq A_5 \quad \text{et} \quad \text{Sym}(D) \simeq A_5 \times C_2.$$

### L'octaèdre et l'icosaèdre : dualité

Qu'en est-il des deux solides platoniciens restants, à savoir l'octaèdre et l'icosaèdre ? En fait, on obtient gratuitement leurs groupes de symétries par l'observation suivante.

Étant donné un polyèdre  $P$ , le *polyèdre dual*  $P^*$  est construit comme suit : on place un sommet au milieu de chaque face de  $P$  ; on relie par une arête deux sommets correspondants à deux faces adjacentes, ce qui donne naissance à une face de  $P^*$  pour chaque sommet de  $P$ . Remarquons que  $(P^*)^* = P$ . On vérifie que  $\text{Sym}(P) = \text{Sym}(P^*)$ , puisqu'une isométrie de  $\mathbb{R}^3$  préserve globalement  $P$  si et seulement si elle préserve globalement  $P^*$ .

Dans les cas des cinq solides platoniciens, on remarque que le tétraèdre  $T$  est auto-dual,  $T^* = T$ , que le dual du cube n'est autre que l'octaèdre,  $C^* = O$  et que le dual du dodécaèdre est l'icosaèdre,  $D^* = I$ . Ces dualités sont illustrées<sup>1</sup> en figure II.5.



FIGURE II.5 – Les dualités  $T^* = T$ ,  $C^* = O$ ,  $O^* = C$ ,  $D^* = I$  et  $I^* = D$ .

On a donc démontré :

#### **Théorème II.15: Symétries des solides platoniciens**

Les groupes de symétries et de symétries propres du tétraèdre  $T$ , du cube  $C$ , de l'octaèdre  $O$ , du dodécaèdre  $D$  et de l'icosaèdre  $I$  dans  $\mathbb{R}^3$  sont

$$\begin{aligned} \text{Sym}^+(T) &\simeq A_4 < S_4 \simeq \text{Sym}(T), \\ \text{Sym}^+(C) &\simeq \text{Sym}^+(O) \simeq S_4 < S_4 \times C_2 \simeq \text{Sym}(O) \simeq \text{Sym}(C), \\ \text{Sym}^+(D) &\simeq \text{Sym}^+(I) \simeq A_5 < A_5 \times C_2 \simeq \text{Sym}(I) \simeq \text{Sym}(D). \quad \square \end{aligned}$$

### II.4.2 Sous-groupes finis de $\text{Isom}(\mathbb{R}^n)$

Maintenant que nous avons déterminé les groupes de symétries de certains objets géométriques, la prochaine question naturelle est la suivante : quels groupes peuvent être réalisés comme le groupe de symétries d'un objet dans  $\mathbb{R}^n$  ? En général, c'est une question très difficile, mais il est possible d'y apporter une réponse complète pour  $n$  petit dans le cas des groupes *finis*. Comme les groupes

1. illustrations prises sur la page web de Thomas Banchoff

de symétries (resp. de symétries propres) sont des sous-groupes de  $\text{Isom}(\mathbb{R}^n)$  (resp.  $\text{Isom}^+(\mathbb{R}^n)$ ), cela revient à étudier les sous-groupes finis de  $\text{Isom}(\mathbb{R}^n)$  et de  $\text{Isom}^+(\mathbb{R}^n)$ .

On commencera par traiter le cas général de  $\mathbb{R}^n$ , puis on se restreindra aux cas de dimension 1, 2 et 3. Voici un premier résultat valable en toute dimension.

**Proposition II.16.** *Pour tout sous-groupe fini  $H < \text{Isom}(\mathbb{R}^n)$ , il existe  $v \in \mathbb{R}^n$  tel que  $f(v) = v$  pour tout  $f \in H$ .*

*Démonstration.* Considérons l'action de  $H$  sur  $\mathbb{R}^n$  et fixons un élément  $x \in \mathbb{R}^n$  quelconque. Comme  $H$  est fini, l'orbite correspondante  $H \cdot x = \{f(x) \mid f \in H\}$  est un sous-ensemble fini de  $\mathbb{R}^n$ , et l'on peut donc considérer son barycentre  $v := B_{H \cdot x} \in \mathbb{R}^n$ . Pour tout  $f \in H$ , on a

$$f(H \cdot x) = (fH) \cdot x = H \cdot x.$$

Le lemme II.14 implique alors

$$f(v) = f(B_{H \cdot x}) = B_{f(H \cdot x)} = B_{H \cdot x} = v$$

pour tout  $f \in H$ . Le barycentre de l'orbite de  $x$  fournit donc bien le point fixe recherché.  $\square$

Notons que cette construction est totalement explicite, et il est instructif de tenter de calculer un point fixe sur des exemples. Notons également que la contraposée de cet énoncé est aussi intéressante : si  $H$  est un sous-groupe de  $\text{Isom}(\mathbb{R}^n)$  qui n'a pas de point fixe, alors  $H$  est infini.

À l'aide de la proposition II.16, on montre maintenant le résultat suivant.

**Proposition II.17.** *Tout sous-groupe fini de  $\text{Isom}(\mathbb{R}^n)$  (resp.  $\text{Isom}^+(\mathbb{R}^n)$ ) est isomorphe à un sous-groupe fini de  $\text{O}(n)$  (resp.  $\text{SO}(n)$ ).*

*Démonstration.* Soit  $H < \text{Isom}(\mathbb{R}^n)$  un sous-groupe fini. Par la proposition II.16, il existe  $v \in \mathbb{R}^n$  tel que  $f(v) = v$  pour tout  $f \in H$ . Soit  $\phi: H \rightarrow \text{Isom}(\mathbb{R}^n)$  l'application définie par  $\phi(f) = \tau_{-v} \circ f \circ \tau_v$ . Il s'agit clairement d'un homomorphisme, par l'argument habituel

$$\phi(f_1 \circ f_2) = \tau_{-v} \circ f_1 \circ f_2 \circ \tau_v = (\tau_{-v} \circ f_1 \circ \tau_v) \circ (\tau_{-v} \circ f_2 \circ \tau_v) = \phi(f_1) \circ \phi(f_2).$$

De plus,  $\phi$  est injectif puisque son noyau est trivial. Finalement, pour tout  $f \in H$ ,

$$\phi(f)(0) = \tau_{-v}(f(\tau_v(0))) = \tau_{-v}(f(v)) = \tau_{-v}(v) = v - v = 0.$$

Ainsi,  $\phi(f)$  est une isométrie qui fixe l'origine, i.e. une isométrie linéaire, donc un élément de  $\text{O}(n)$ . En résumé, on a un homomorphisme injectif  $\phi: H \rightarrow \text{O}(n)$ , qui induit donc un isomorphisme  $H \simeq \phi(H) < \text{O}(n)$ .

Finalement, si  $f$  préserve l'orientation, alors  $\phi(f)$  fait clairement de même. Ainsi, si  $H$  est un sous-groupe de  $\text{Isom}^+(\mathbb{R}^n)$ ,  $\phi$  définit un isomorphisme entre  $H$  et  $\phi(H) < \text{SO}(n)$ .  $\square$

En conclusion de cette proposition, classifier les groupes finis de symétries (resp. de symétries propres) d'objets de  $\mathbb{R}^n$  revient à classifier les sous-groupes finis de  $\text{O}(n)$  (resp.  $\text{SO}(n)$ ).

### Le cas de la droite $\mathbb{R}$

Le cas  $n = 1$  est complètement trivial :  $O(1) = \{\pm 1\}$  et  $SO(1) = \{1\}$ , reflétant le fait que les sous-ensembles de  $\mathbb{R}$  qui ont un nombre fini de symétries ont au plus une symétrie non-triviale, qui n'est pas une symétrie propre.

### Le cas du plan $\mathbb{R}^2$

Passons donc au cas du plan. Rappelons que dans ce cas, on a vu que certains objets ont des groupes triviaux  $\text{Sym}^+(Y) = \text{Sym}(Y) = \{\text{id}\} =: C_1$ , la croix  $Y = \dagger$  satisfait  $\text{Sym}^+(Y) = C_1 < \{\text{id}, \sigma\} = \text{Sym}(Y)$ , un groupe<sup>2</sup> qu'on notera  $D_2$ , tandis que pour  $n \geq 2$ , la scie circulaire à  $n$  dents  $Y_n$  satisfait  $\text{Sym}^+(Y_n) = \text{Sym}(Y_n) = C_n$ , et le  $n$ -gone régulier  $P_n$  admet  $\text{Sym}^+(P_n) = C_n < D_{2n} = \text{Sym}(P_n)$ .

Nous allons maintenant démontrer que ce sont les seuls cas possibles ! Plus précisément :

#### Théorème II.18

Soit  $Y \subset \mathbb{R}^2$  un objet du plan n'admettant qu'un nombre fini  $n$  de symétries propres. Alors,  $\text{Sym}^+(Y)$  est cyclique d'ordre  $n$ . De plus,  $\text{Sym}(Y)$  est soit cyclique d'ordre  $n$ , soit isomorphe au groupe  $D_{2n}$ .

*Démonstration.* Soit donc  $Y \subset \mathbb{R}^2$  avec  $\text{Sym}^+(Y)$  d'ordre  $n$ . Par la proposition II.17,  $\text{Sym}^+(Y)$  est isomorphe à un sous-groupe  $H^+$  d'ordre  $n$  de  $SO(2)$  qui, rappelons-le, consiste en toutes les rotations  $\{\rho_0^\alpha\}_{\alpha \in [0, 2\pi)}$  autour de l'origine. Pour démontrer la première assertion, il suffit donc de vérifier que  $H^+$  est cyclique. Comme c'est trivialement le cas si  $H^+$  est d'ordre 1, supposons  $n = |H^+| \geq 2$ . Comme  $H^+$  est fini et non-trivial, il existe  $\alpha_0 \in (0, 2\pi)$  le plus petit angle de rotation parmi tous les éléments non-triviaux de  $H^+$ . Nous allons maintenant vérifier que  $H^+$  est un groupe cyclique engendré par  $\rho_0^{\alpha_0}$ . En effet, considérons  $\rho_0^\alpha \in H^+$  quelconque, avec  $0 \leq \alpha < 2\pi$ . Par division de  $\alpha$  par  $\alpha_0$ , il existe  $k \in \{0, 1, 2, \dots\}$  et  $\beta \in [0, \alpha_0)$  tel que  $\alpha = k\alpha_0 + \beta$ . Il suit que

$$\rho_0^\beta = \rho_0^{\alpha - k\alpha_0} = \rho_0^\alpha \circ (\rho_0^{\alpha_0})^{-k}$$

est un élément de  $H^+$ , puisque  $\rho_0^{\alpha_0}$  et  $\rho_0^\alpha$  sont éléments du sous-groupe  $H^+$ . Mais par minimalité de  $\alpha_0$ , on a nécessairement  $\beta = 0$ , i.e.  $\rho_0^\alpha = (\rho_0^{\alpha_0})^k$ . Cela signifie que tout élément de  $H^+$  est de la forme une puissance de  $\rho_0^{\alpha_0}$ , i.e. que  $H^+$  est le groupe cyclique engendré par  $\rho_0^{\alpha_0}$ . Il suit aussi que  $\alpha_0 = \frac{2\pi}{n}$ , puisque  $H^+$  est d'ordre  $n$ .

Passons à la seconde partie de l'énoncé. Par la proposition II.17, il existe un homomorphisme injectif  $\phi: \text{Sym}(Y) \rightarrow O(2)$ , ce qui implique que  $\text{Sym}(Y)$  est isomorphe à  $H = \phi(\text{Sym}(Y))$  et  $\text{Sym}^+(Y)$  à  $H^+ = \phi(\text{Sym}^+(Y)) = H \cap SO(2)$ . Par la discussion ci-dessus,  $H^+ = H \cap SO(2)$  est cyclique d'ordre  $n$ , engendré

2. bien entendu, ce groupe est isomorphe au groupe cyclique  $C_2$  d'ordre 2, mais il apparaît ici naturellement comme un produit (semi-)direct  $C_1 \rtimes C_2$ , d'où cette notation

par la rotation d'angle  $\frac{2\pi}{n}$  autour de l'origine. Si  $H$  est inclus dans  $\text{SO}(2)$ , alors  $\text{Sym}(Y) \simeq H = H^+ \simeq C_n$ . Supposons donc que  $H$  n'est pas inclus dans  $\text{SO}(2)$ . Comme les éléments de  $\text{O}(2) \setminus \text{SO}(2)$  sont précisément les réflexions par les droites contenant l'origine (voir la remarque après la preuve du théorème II.10), cela signifie que  $H$  contient une telle réflexion  $\sigma$ . Nous allons maintenant vérifier que  $H$  est le produit semi-direct  $H = H^+ \rtimes C_2$  avec  $C_2 = \{\text{id}, \sigma\}$ , ce qui donnera

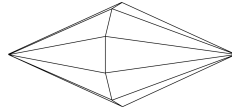
$$H = H^+ \rtimes C_2 \simeq C_n \rtimes C_2 = D_{2n}$$

et conclura la preuve. En effet, un élément quelconque  $h \in H$  s'écrit  $h = (h\sigma^\varepsilon)\sigma^\varepsilon$  avec  $\varepsilon = 0$  si  $h \in \text{SO}(2)$  et  $\varepsilon = 1$  sinon ; comme  $h\sigma^\varepsilon$  est un élément de  $H^+$ , cela montre que tout  $h \in H$  s'écrit comme produit d'un élément de  $H^+$  et d'un élément de  $C_2$ . Cette écriture est unique puisque  $\varepsilon$  est entièrement déterminé par  $h$ . De plus, tout élément  $\rho$  de  $H^+$  est une rotation autour de l'origine, et  $\sigma$  est une réflexion par une droite contenant l'origine : par l'illustration en fin de paragraphe I.1.1, on a l'égalité  $\sigma\rho\sigma^{-1} = \rho^{-1}$ . Cela démontre que  $H = H^+ \rtimes C_2$  et conclut la démonstration.  $\square$

### Le cas de l'espace $\mathbb{R}^3$

Qu'en est-il des groupes de symétries d'objets dans l'espace  $\mathbb{R}^3$  ? Dans ce cas, il est naturel de se restreindre aux groupes de symétries propres, puisqu'ils décrivent les façons dont on peut déplacer un objet physique dans l'espace à trois dimensions.

Rappelons que l'on a déjà vu apparaître les groupes suivants comme groupes de symétries propres d'objets tridimensionnels :  $A_4$  (pour le tétraèdre),  $S_4$  (pour le cube), et  $A_5$  (pour le dodécaèdre). De plus, on vérifie facilement que le groupe diédral  $D_{2n}$  ( $n \geq 2$ ) est réalisé comme le groupe des symétries propres d'une double pyramide sur un  $n$ -gone régulier, comme illustré ci-dessous. (Un collier à  $n$  perles en est un autre exemple, voir la fin de la section I.3.5.) Finalement, si l'on considère une simple pyramide sur ce même  $n$ -gone régulier, on obtient le groupe cyclique  $C_n$ .



Et ce sont les seuls cas possibles !

#### Théorème II.19

Soit  $Y \subset \mathbb{R}^3$  un objet de l'espace n'admettant qu'un nombre fini de symétries propres. Alors,  $\text{Sym}^+(Y)$  est isomorphe à un unique groupe parmi la liste suivante :  $C_n$  ( $n \geq 1$ ),  $D_{2n}$  ( $n \geq 2$ ),  $A_4$ ,  $S_4$  et  $A_5$ .

La preuve complète est absolument splendide et utilise l'intégralité des résultats du chapitre I. Néanmoins, par manque de temps, nous n'allons pas la

discuter en cours. Cette démonstration n'est donc pas exigible à l'examen, mais est donnée ici de manière incomplète, et "pour votre culture générale".

*Idée de la démonstration.* Soit donc  $Y \subset \mathbb{R}^3$  n'admettant qu'un nombre fini de symétries propres. Par la proposition II.17,  $\text{Sym}^+(Y)$  est isomorphe à un sous-groupe fini  $H$  du groupe  $\text{SO}(3)$  qui, rappelons-le, consiste en toutes les rotations autour d'axes passant par l'origine. (Nous avons vu une preuve de ce fait en première partie du cours à l'aide des valeurs propres.) Notons que  $\text{SO}(3)$  (et donc  $H$ ) préserve globalement la sphère  $\mathbb{S}^2 = \{x \in \mathbb{R}^3 \mid \|x\| = 1\}$ . Ainsi,  $H$  agit sur la sphère. L'idée est de supposer  $H$  non-trivial et de considérer l'ensemble

$$X = \{x \in \mathbb{S}^2 \mid \text{Stab}(x) \neq \{\text{id}\}\}$$

des points de la sphère tels que le stabilisateur pour l'action de  $H$  sur  $\mathbb{S}^2$  est non-trivial. Géométriquement, cela correspond à l'ensemble des points  $x \in \mathbb{S}^2$  tels que  $H$  contient une rotation d'axe  $\overline{0x}$ . Notons que  $x$  est élément de  $X$  si et seulement si le point antipodal  $-x$  est élément de  $X$ . Il découle de la proposition I.5 que l'action de  $H$  sur  $\mathbb{S}^2$  se restreint à une action de  $H$  sur  $X$ . Comme tout élément non-trivial  $g \in H$  est une rotation,  $g$  fixe exactement deux points de la sphère. Il suit que le nombre de points fixes  $|X^g|$  vaut 2 pour tout  $g \in H \setminus \{\text{id}\}$ . En notant  $N$  le nombre d'orbites pour l'action de  $H$  sur  $X$ ,  $\{x_1, \dots, x_N\}$  un système de représentants de ces orbites, et  $\nu_i := |\text{Stab}(x_i)| \geq 2$  pour  $i = 1, \dots, N$ , il suit de la formule de Burnside et du corollaire I.7 :

$$2(|H| - 1) = \sum_{g \in H \setminus \{\text{id}\}} |X^g| = \sum_{g \in H} |X^g| - |X| = |H|N - \sum_{i=1}^N \frac{|H|}{|\text{Stab}(x_i)|} = |H| \sum_{i=1}^N \left(1 - \frac{1}{\nu_i}\right).$$

On obtient donc l'équation suivante :

$$2 - \frac{2}{|H|} = \sum_{i=1}^N \left(1 - \frac{1}{\nu_i}\right) \quad \text{avec} \quad \nu_1 \geq \nu_2 \geq \dots \geq \nu_N \geq 2.$$

On vérifie facilement que les seules valeurs possibles pour  $N$  sont  $N = 2$  et  $N = 3$ .

Le cas  $N = 2$  mène aux solutions  $\nu_1 = \nu_2 = |H|$ . Ainsi,  $x_1$  et  $x_2$  sont fixés par tout le groupe  $H = \text{Stab}(x_1) = \text{Stab}(x_2)$ , i.e.  $x_1$  et  $x_2$  sont seuls dans leur orbite et  $X = \{x_1, x_2\}$ . L'ensemble  $X$  consiste donc en deux points antipodaux  $\{x, -x\}$ . Comme les isométries de  $H$  fixent  $x$  et 0, elles fixent la droite  $\ell = \overline{0x}$ ; tous les éléments de  $H$  sont donc des rotations de la forme  $\rho_\ell^\alpha$  avec  $\alpha \in [0, 2\pi)$ . On se retrouve exactement dans la situation de la preuve du théorème II.18, et l'on obtient que  $H$  est un groupe cyclique.

Le cas  $N = 3$  est plus riche. Il correspond à l'équation

$$1 + \frac{2}{|H|} = \frac{1}{\nu_1} + \frac{1}{\nu_2} + \frac{1}{\nu_3} \quad \text{avec} \quad \nu_1 \geq \nu_2 \geq \nu_3 \geq 2.$$

On vérifie que  $\nu_3$  est forcément égal à 2, tandis que  $\nu_2$  peut valoir 2 ou 3. Le cas  $\nu_2 = 2$  mène à  $(\nu_1, \nu_2, \nu_3) = (n, 2, 2)$  avec  $n \geq 2$ , d'où  $H \simeq D_{2n}$  d'une manière similaire à la preuve de la seconde partie du théorème II.18. Pour  $\nu_2 = 3$ , on obtient exactement 3 solutions à l'équation ci-dessus :  $(\nu_1, \nu_2, \nu_3)$  peut valoir  $(3, 3, 2)$ ,  $(4, 3, 2)$  ou  $(5, 3, 2)$ . Les ordres de  $H$  correspondant sont  $|H| = 12, 24$  et  $60$ , et l'on vérifie que  $H$  est alors isomorphe à  $A_4$ ,  $S_4$  et  $A_5$ , respectivement.  $\square$

## Chapitre III: Géométrie hyperbolique

Le but de ce chapitre est de donner un premier aperçu de la *géométrie hyperbolique* en dimension 2. C'est un exemple de *géométrie non-euclidienne* : une géométrie qui satisfait aux postulats I à IV d'Euclide, mais pas au postulat V. Cela démontre que le cinquième postulat ne peut pas être démontré à partir des autres, mettant ainsi un terme à deux mille ans de tentatives infructueuses. Mais la géométrie hyperbolique n'est pas une construction exotique destinée à démontrer l'indépendance du cinquième postulat ! C'est un outil mathématique très utile qui, à l'instar de la géométrie euclidienne, apparaît naturellement dans bien des domaines (étude des variétés, théorie de la relativité, théorie des nombres,...).

Ce chapitre est organisé comme suit. La section III.1 sera dédiée aux *inversions* : ce sont les transformations du plan qui jouent en géométrie hyperbolique le rôle des réflexions en géométrie euclidienne. En section III.2, on étudiera le *groupe de Möbius* engendré par ces transformations, et une quantité invariante par l'action de ce groupe : le *birapport*. On utilisera ensuite ces résultats en section III.3 pour définir la *distance hyperbolique* sur le disque unité ; l'espace métrique obtenu, appelé *le disque de Poincaré*, est un des modèles possibles de la géométrie hyperbolique. Finalement, en section III.4, on introduira un autre modèle appelé *le demi-plan de Poincaré*, qui nous permettra de déterminer plus facilement le groupe des isométries du plan hyperbolique.

La relative originalité de notre exposition réside dans le fait que nous n'utilisons presque aucun outil de géométrie analytique, la quasi-totalité de nos preuves reposant sur des arguments synthétiques.

### III.1 Inversions

Dans cette section, nous allons à nouveau faire de la géométrie synthétique dans le plan, comme au temps (regretté ?) de la première partie du cours. Le but est de rassembler tous les résultats de ce type dont nous aurons besoin pour la suite du chapitre. En cours de route, on verra aussi quelques applications magnifiques, dont le célèbre *porisme de Steiner*.

#### III.1.1 Définition et premières propriétés

Comme on l'a vu au chapitre précédent, les réflexions jouent un rôle prépondérant en géométrie euclidienne. Notons que si l'on fixe une droite  $\ell$  du plan, alors

la réflexion correspondante  $\sigma_\ell: \mathbb{R}^2 \rightarrow \mathbb{R}^2$  satisfait les trois propriétés suivantes :

- (i)  $\sigma_\ell(P) = P$  pour tout  $P \in \ell$ .
- (ii)  $\sigma_\ell$  échange les deux côtés de  $\ell$ .
- (iii)  $\sigma_\ell$  laisse globalement invariants les cercles et les droites orthogonaux à  $\ell$ .

De plus, on vérifie facilement que  $\sigma_\ell$  est uniquement déterminée par ces trois propriétés (exercice).

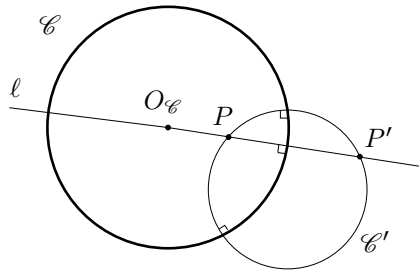
La question naturelle à laquelle on se propose de répondre est la suivante : si l'on remplace la droite  $\ell$  par un cercle  $\mathcal{C}$ , existe-t-il une transformation analogue ? La réponse est positive.

La construction et l'étude de cette transformation reposent sur la notion de *puissance d'un point* vue en première partie (Euclide III.35 et III.36). Rappelons donc qu'étant donné un cercle  $\mathcal{C}$ , un point  $A$  du plan, et une droite passant par  $A$  et coupant  $\mathcal{C}$  en deux points  $P, P'$ , le produit des distances  $AP \cdot AP'$  ne dépend que de  $A$  et de  $\mathcal{C}$ , mais pas du choix de la droite : c'est la *puissance du point  $A$  par rapport au cercle  $\mathcal{C}$* .

**Proposition III.1.** *Étant donné un cercle  $\mathcal{C}$  de centre  $O_\mathcal{C}$  dans le plan  $\mathbb{R}^2 = \mathbb{C}$ , il existe une unique application  $i_\mathcal{C}: \mathbb{C} \setminus \{O_\mathcal{C}\} \rightarrow \mathbb{C} \setminus \{O_\mathcal{C}\}$  qui satisfait les trois propriétés suivantes :*

- (i)  $i_\mathcal{C}(P) = P$  pour tout  $P \in \mathcal{C}$ .
- (ii)  $i_\mathcal{C}$  échange l'intérieur et l'extérieur de  $\mathcal{C}$ .
- (iii)  $i_\mathcal{C}$  laisse globalement invariants les cercles et les droites orthogonaux à  $\mathcal{C}$ .

*Démonstration.* Vérifions dans un premier temps l'unicité d'une application  $i_\mathcal{C}$  satisfaisant les trois propriétés ci-dessus. Commençons par un point  $P \neq O_\mathcal{C}$  à l'intérieur de  $\mathcal{C}$ . Soient  $\ell$  la droite  $\overline{O_\mathcal{C}P}$  et  $\mathcal{C}'$  un cercle passant par  $P$  orthogonal à  $\mathcal{C}$ , comme illustré ci-dessous (on vérifie facilement qu'un tel cercle existe, et intersecte  $\ell$  en deux points).



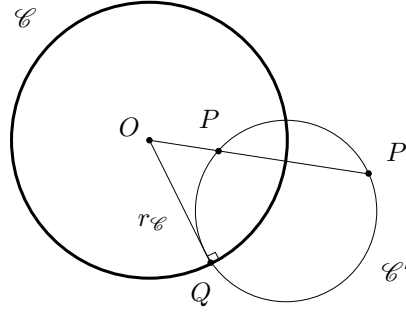
Par la propriété (iii),

$$i_\mathcal{C}(P) \in i_\mathcal{C}(\ell) \cap i_\mathcal{C}(\mathcal{C}') = \ell \cap \mathcal{C}' = \{P, P'\}.$$

Par la propriété (ii), on a forcément  $i_\mathcal{C}(P) = P'$ . Ainsi, l'image d'un point à l'intérieur de  $\mathcal{C}$  est entièrement déterminée par ces trois propriétés. La même

preuve fonctionne parfaitement pour un point à l'extérieur de  $\mathcal{C}$  : il faut juste échanger les rôles de  $P$  et de  $P'$ . Finalement, si  $P$  est un point de  $\mathcal{C}$ , alors  $i_{\mathcal{C}}(P) = P$  par la propriété (i).

Reste à démontrer que l'application  $i_{\mathcal{C}}$  construite ci-dessus est bien définie, i.e. à vérifier que  $i_{\mathcal{C}}(P)$  ne dépend pas du choix du cercle  $\mathcal{C}'$  orthogonal à  $\mathcal{C}$  passant par  $P$ , et qu'elle satisfait les propriétés (i) à (iii). Soient  $r_{\mathcal{C}}$  le rayon du cercle  $\mathcal{C}$ , et  $Q$  un point d'intersection de  $\mathcal{C}'$  avec  $\mathcal{C}$ . Par la proposition III.36 d'Euclide, la puissance du point  $O = O_{\mathcal{C}}$  par rapport au cercle  $\mathcal{C}'$  est donnée par  $r_{\mathcal{C}}^2 = (OQ)^2 = OP \cdot OP'$ . (Voir ci-dessous.) Ainsi,  $P' = i_{\mathcal{C}}(P)$  est le point sur le rayon  $\overline{OP}$  à distance  $OP'$  telle que  $OP \cdot OP' = r_{\mathcal{C}}^2$ , ce qui est indépendant du choix de  $\mathcal{C}'$ . Par construction, cette application bien définie satisfait les propriétés (i) à (iii).  $\square$



Soulignons que nous avons démontré au passage que l'application  $i_{\mathcal{C}}$  envoie un point  $P \neq O_{\mathcal{C}}$  sur l'unique point  $P'$  sur le rayon  $\overline{O_{\mathcal{C}}P}$  tel que  $O_{\mathcal{C}}P \cdot O_{\mathcal{C}}P' = r_{\mathcal{C}}^2$ .

**Terminologie.** L'application  $i_{\mathcal{C}}: \mathbb{C} \setminus \{O_{\mathcal{C}}\} \rightarrow \mathbb{C} \setminus \{O_{\mathcal{C}}\}$  est appelée **l'inversion** de cercle  $\mathcal{C}$ .

#### Remarques.

1. Une inversion composée avec elle-même donne l'identité : c'est ce qu'on appelle une *involution*. Ce fait se vérifie immédiatement au moyen de la propriété énoncée ci-dessus.
2. Plus généralement, si  $\mathcal{C}$  et  $\mathcal{C}'$  sont deux cercles de même centre  $O$  et de rayon  $r$  et  $r'$ , alors la composition  $i_{\mathcal{C}'} \circ i_{\mathcal{C}}$  est simplement l'homothétie de centre  $O$  et de facteur  $(r'/r)^2$ .
3. Si l'on considère  $\mathcal{C} = \mathbb{S}^1$  le cercle de rayon 1 centré en  $0 \in \mathbb{C}$ , alors l'inversion correspondante  $i_{\mathbb{S}^1}: \mathbb{C}^* \rightarrow \mathbb{C}^*$  est donnée par la formule  $i_{\mathbb{S}^1}(z) = \frac{1}{\bar{z}}$ , ou en coordonnées polaires, par  $i_{\mathbb{S}^1}(re^{i\varphi}) = \frac{1}{r}e^{i\varphi}$ . En effet,  $i_{\mathbb{S}^1}(z)$  se situe sur la droite  $\overline{Oz}$  (ce qui signifie que l'argument de  $i_{\mathbb{S}^1}(z)$  est égal à l'argument de  $z$ ), et sa distance à l'origine multipliée par la distance à l'origine de  $z$  vaut 1 (ce qui signifie que le module de  $i_{\mathbb{S}^1}(z)$  est l'inverse du module de  $z$ ).

4. Plus généralement, si  $\mathcal{C}$  est le cercle de rayon  $r$  centré en  $a \in \mathbb{C}$ , alors l'inversion de cercle  $\mathcal{C}$  est donnée par la formule  $i_{\mathcal{C}}(z) = \frac{r^2}{\bar{z}-\bar{a}} + a$ . Cela se vérifie en remarquant que  $i_{\mathcal{C}} = \varphi_{\mathcal{C}}^{-1} \circ i_{\mathbb{S}^1} \circ \varphi_{\mathcal{C}}$ , où  $\varphi_{\mathcal{C}}$  est la composition de la translation  $z \mapsto z - a$  et de l'homothétie  $z \mapsto r^{-1}z$ . (Les détails sont laissés en exercice.)

Voici quelques propriétés moins immédiates. Elles sont illustrées en figure III.1.

**Proposition III.2.** *L'inversion  $i_{\mathcal{C}}: \mathbb{C} \setminus \{O_{\mathcal{C}}\} \rightarrow \mathbb{C} \setminus \{O_{\mathcal{C}}\}$  envoie*

- (i) *les droites passant par  $O_{\mathcal{C}}$  sur elles-mêmes ;*
- (ii) *les cercles passant par  $O_{\mathcal{C}}$  sur des droites ne passant pas par  $O_{\mathcal{C}}$  ;*
- (iii) *les droites ne passant pas par  $O_{\mathcal{C}}$  sur des cercles passant par  $O_{\mathcal{C}}$  ;*
- (iv) *les cercles ne passant pas par  $O_{\mathcal{C}}$  sur des cercles ne passant pas par  $O_{\mathcal{C}}$ .*

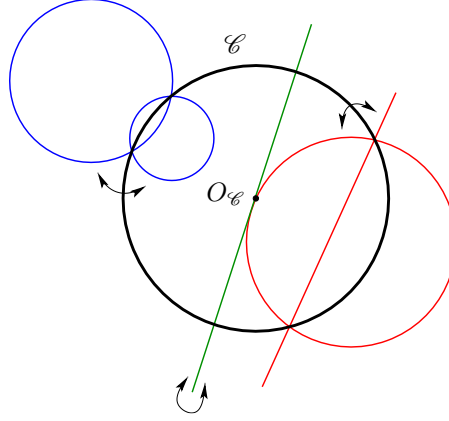


FIGURE III.1 – Propriétés de l'inversion  $i_{\mathcal{C}}$  énoncées en proposition III.2.

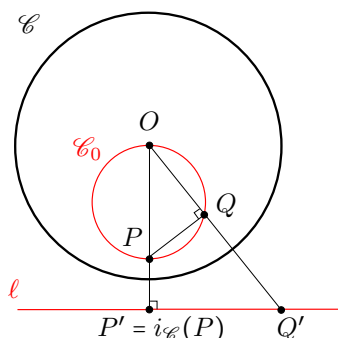
*Démonstration.* Les droites passant par  $O = O_{\mathcal{C}}$  étant orthogonales à  $\mathcal{C}$ , elles sont globalement préservées par définition de  $i_{\mathcal{C}}$ . Cela démontre le premier point.

Puisque  $i_{\mathcal{C}}$  est une involution, les points (ii) et (iii) sont équivalents. Vérifions donc (ii). Soit  $\mathcal{C}_0$  un cercle passant par  $O$ . Soient  $\overline{OP}$  un diamètre de  $\mathcal{C}_0$ , et  $\ell$  la perpendiculaire à  $\overline{OP}$  passant par  $P' = i_{\mathcal{C}}(P)$ . Je prétends que  $i_{\mathcal{C}}$  envoie  $\mathcal{C}_0 \setminus \{O\}$  sur  $\ell$ , ce qui démontre (ii). En effet, soit  $Q \in \mathcal{C}_0 \setminus \{O, P\}$ , et soit  $Q' = \overline{OQ} \cap \ell$  comme illustré ci-dessous. Les triangles  $OQP$  et  $OP'Q'$  sont semblables ; par Thalès, on a donc

$$\frac{OP}{OQ} = \frac{OQ'}{OP'}, \quad \text{d'où} \quad OQ \cdot OQ' = OP \cdot OP' = r_{\mathcal{C}}^2,$$

ce qui montre que  $i_{\mathcal{C}}(Q) = Q'$ .

Passons enfin au point (iv). Soit donc  $\mathcal{C}'$  un cercle ne passant pas par le centre  $O$  de  $\mathcal{C}$ . Considérons le cercle  $\mathcal{C}''$  centré en  $O$  de rayon  $r$  tel que  $r^2$  est



la puissance du point  $O$  par rapport au cercle  $\mathcal{C}'$ . Par construction, on a alors  $i_{\mathcal{C}''}(\mathcal{C}') = \mathcal{C}'$ . Il suit

$$i_{\mathcal{C}}(\mathcal{C}') = i_{\mathcal{C}}(i_{\mathcal{C}''}(\mathcal{C}')) = (i_{\mathcal{C}} \circ i_{\mathcal{C}''})(\mathcal{C}'),$$

qui par la remarque 2 ci-dessus est l'image du cercle  $\mathcal{C}'$  par une homothétie de centre  $O$ . Il s'agit donc bien à nouveau d'un cercle ne passant pas par  $O$ . Cela termine la démonstration du point (iv), et de la proposition.  $\square$

La prochaine étape consiste à vérifier que les inversions sont des applications *conformes*, c'est-à-dire qu'elles conservent les angles. Comme d'habitude, il est possible de démontrer cet énoncé en utilisant des arguments synthétiques ou des méthodes analytiques. Malheureusement, les preuves synthétiques sont relativement pénibles. Nous allons donc faire une entorse au principe énoncé en début de chapitre, et proposer une preuve analytique.

Une *courbe orientée* est une application  $\gamma: [0, 1] \rightarrow \mathbb{C}$  continûment différentiable avec  $\gamma'(t)$  partout non-nul. La *tangente* à la courbe  $\gamma$  au point  $P = \gamma(t_0)$  est la droite orientée d'équation paramétrique  $\gamma(t_0) + \lambda \gamma'(t_0)$ , avec  $\lambda \in \mathbb{R}$ . Étant données deux courbes orientées  $\gamma$  et  $\delta$  se coupant en un point  $P = \gamma(t_0) = \delta(t_1)$ , l'*angle orienté* entre  $\gamma$  et  $\delta$  au point  $P$  est défini comme l'angle orienté entre la tangente à  $\gamma$  en  $P$  et la tangente à  $\delta$  en  $P$  (voir figure III.2). En d'autres termes, il s'agit de l'angle orienté entre les vecteurs  $\gamma'(t_0)$  et  $\delta'(t_1)$ . Finalement, on dira qu'une application continûment différentiable  $f: \Omega \rightarrow \mathbb{C}$  (définie sur un sous-ensemble  $\Omega \subset \mathbb{C}$ ) *préserve les angles orientés* (resp. *renverse les angles orientés*) si elle satisfait la condition suivante : si  $\gamma$  et  $\delta$  sont deux courbes orientées qui s'intersectent en  $P \in \Omega$  en l'angle orienté  $\alpha$ , alors les courbes orientées  $f(\gamma)$  et  $f(\delta)$  s'intersectent en  $f(P)$  en l'angle orienté  $\alpha$  (resp.  $-\alpha$ ).

**Lemme III.3.** Soit  $f: \Omega \rightarrow \mathbb{R}^2$ ,  $(x, y) \mapsto \begin{pmatrix} f_1(x, y) \\ f_2(x, y) \end{pmatrix}$  une application continûment différentiable. Alors,  $f$  préserve les angles orientés si et seulement si, pour tout  $(x, y) \in \Omega$ ,

$$f'(x, y) = \begin{pmatrix} \frac{\partial f_1}{\partial x}(x, y) & \frac{\partial f_1}{\partial y}(x, y) \\ \frac{\partial f_2}{\partial x}(x, y) & \frac{\partial f_2}{\partial y}(x, y) \end{pmatrix}$$

est une matrice de la forme  $\begin{pmatrix} \lambda_1 & -\lambda_2 \\ \lambda_2 & \lambda_1 \end{pmatrix}$  avec  $\lambda_1^2 + \lambda_2^2 > 0$ .

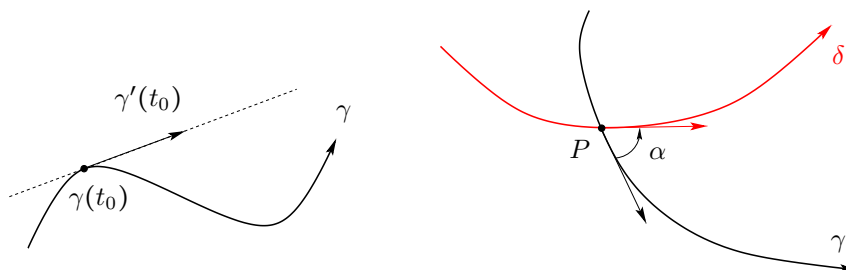


FIGURE III.2 – À gauche : la tangente à  $\gamma$  au point  $P = \gamma(t_0)$ . À droite : l'angle orienté  $\alpha$  entre  $\gamma$  et  $\delta$  au point  $P$ .

*Démonstration.* L'application  $f$  envoie la courbe orientée  $\gamma$  sur la courbe orientée  $f \circ \gamma$ . Ainsi, la tangente à  $f \circ \gamma$  au point  $f(P) = (f \circ \gamma)(t_0)$  est la droite d'équation paramétrique  $(f \circ \gamma)(t_0) + \lambda(f \circ \gamma)'(t_0)$ . Par la règle de dérivation en chaîne,

$$(f \circ \gamma)'(t_0) = f'(\gamma(t_0))(\gamma'(t_0)) = \begin{pmatrix} \frac{\partial f_1}{\partial x}(\gamma(t_0)) & \frac{\partial f_1}{\partial y}(\gamma(t_0)) \\ \frac{\partial f_2}{\partial x}(\gamma(t_0)) & \frac{\partial f_2}{\partial y}(\gamma(t_0)) \end{pmatrix} \begin{pmatrix} \gamma'_1(t_0) \\ \gamma'_2(t_0) \end{pmatrix}.$$

L'angle entre deux courbes orientées  $\gamma$  et  $\delta$  en  $P = \gamma(t_0) = \delta(t_1)$  étant défini comme l'angle entre les vecteurs  $\gamma'(t_0)$  et  $\delta'(t_1)$ , il reste à montrer qu'une application linéaire de  $\mathbb{R}^2$  dans  $\mathbb{R}^2$  conserve les angles orientés entre vecteurs si et seulement si sa matrice (dans une base orthonormale) est de la forme  $\begin{pmatrix} \lambda_1 & -\lambda_2 \\ \lambda_2 & \lambda_1 \end{pmatrix}$  avec  $\lambda_1^2 + \lambda_2^2 > 0$ . (Notons qu'il s'agit de la composition d'une rotation et d'une homothétie.) Cette vérification est laissée en exercice.  $\square$

**Proposition III.4.** *L'inversion  $i_{\mathcal{C}}: \mathbb{C} \setminus \{O_{\mathcal{C}}\} \rightarrow \mathbb{C} \setminus \{O_{\mathcal{C}}\}$  renverse les angles orientés.*

*Démonstration.* Comme on l'a déjà vu,  $i_{\mathcal{C}} = \varphi_{\mathcal{C}}^{-1} \circ i_{\mathbb{S}^1} \circ \varphi_{\mathcal{C}}$ , où  $\varphi_{\mathcal{C}}$  est la composition d'une translation et d'une homothétie. Clairement, ces transformations préservent les angles ; ainsi, il suffit de vérifier que  $i_{\mathbb{S}^1}$  renverse les angles orientés. Mais cette inversion est la composition de l'application  $f: \mathbb{C}^* \rightarrow \mathbb{C}^*$  donnée par  $f(z) = \frac{1}{z} = \frac{\bar{z}}{|z|^2}$  et de la conjugaison complexe. La conjugaison complexe renverse les angles orientés ; il reste donc à vérifier que  $f$  préserve les angles orientés. En notations réelles, cette application est donnée par  $f(x, y) = (\frac{x}{x^2+y^2}, \frac{-y}{x^2+y^2})$ , et l'on vérifie facilement que  $f'(x, y)$  satisfait les équations du lemme III.3.  $\square$

### III.1.2 Une application : le porisme de Steiner

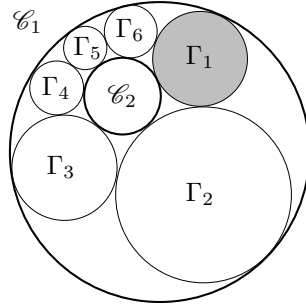
Voici une application absolument magnifique : le *porisme*<sup>1</sup> de Steiner, un des rares résultats mathématiques non-triviaux dont l'énoncé est plus compliqué que la preuve. Il est illustré en figure III.3.

1. **Porisme** : Terme par lequel les mathématiciens des temps modernes désignent certaines propositions qui étaient en usage dans la géométrie des Grecs. Un porisme est une question dont la solution consiste à tirer une vérité géométrique de conditions assignées par l'énoncé. (*Le Littré*)

**Théorème III.5: Porisme de Steiner**

Soient  $\mathcal{C}_1$  et  $\mathcal{C}_2$  deux cercles fixés, avec  $\mathcal{C}_2$  à l'intérieur de  $\mathcal{C}_1$ . Choisissons un cercle  $\Gamma_1$  tangent à  $\mathcal{C}_1$  et  $\mathcal{C}_2$ ; cela détermine une chaîne de cercles  $\Gamma_2, \Gamma_3, \dots$  entre  $\mathcal{C}_1$  et  $\mathcal{C}_2$  avec  $\Gamma_i$  simultanément tangent à  $\Gamma_{i-1}$ ,  $\mathcal{C}_1$  et  $\mathcal{C}_2$ . Soit  $n$  le plus petit entier positif (ou infini) tel que  $\Gamma_n$  coïncide avec  $\Gamma_1$ . Alors,  $n$  est indépendant du choix de  $\Gamma_1$  et ne dépend que de  $\mathcal{C}_1$  et de  $\mathcal{C}_2$ .

*Démonstration.* Cet énoncé est trivialement vrai dans le cas où les deux cercles  $\mathcal{C}_1$  et  $\mathcal{C}_2$  sont concentriques. Comme les inversions préservent tous les termes de l'énoncé, il suffit de vérifier qu'étant donnés deux cercles  $\mathcal{C}_1$  et  $\mathcal{C}_2$  avec  $\mathcal{C}_2$  à l'intérieur de  $\mathcal{C}_1$ , il existe une inversion  $i_{\mathcal{C}}$  telle que  $i_{\mathcal{C}}(\mathcal{C}_1)$  et  $i_{\mathcal{C}}(\mathcal{C}_2)$  sont concentriques. La construction explicite est laissée en exercice.  $\square$

FIGURE III.3 – Le porisme de Steiner. Ici,  $n = 7$ .

Notons que ce type d'argument est similaire à ce que l'on a vu en géométrie projective en première partie, par exemple dans la preuve du théorème de Pappus par Poncelet. L'idée fondamentale est de transformer un énoncé donné en un énoncé plus facile au moyen d'une transformation du plan (ici : une inversion ; là : une transformation projective) qui laisse invariants tous les ingrédients de l'énoncé (ici : être deux cercles tangents ; là : être deux droites qui se coupent).

On verra de nombreuses autres applications en exercices.

**III.1.3 Projection stéréographique**

Voici une observation qui ressemble à première vue à un miracle. Posons  $\hat{\mathbb{C}} := \mathbb{C} \cup \{\infty\}$  et appelons «cercle dans  $\hat{\mathbb{C}}$ » les cercles dans  $\mathbb{C}$  ainsi que les droites dans  $\mathbb{C}$  avec le point  $\infty$ , que l'on interprète comme des cercles de rayon infini passant par  $\infty$ . Alors, toute la théorie s'uniformise de façon magique :

- l'inversion  $i_{\mathcal{C}}: \mathbb{C} \setminus \{O_{\mathcal{C}}\} \rightarrow \mathbb{C} \setminus \{O_{\mathcal{C}}\}$  s'étend de manière naturelle en une application  $i_{\mathcal{C}}: \hat{\mathbb{C}} \rightarrow \hat{\mathbb{C}}$  en posant  $i_{\mathcal{C}}(O_{\mathcal{C}}) = \infty$  (et  $i_{\mathcal{C}}(\infty) = O_{\mathcal{C}}$ ) ;
- la réflexion  $\sigma_{\ell}: \mathbb{C} \rightarrow \mathbb{C}$  s'étend en  $\sigma_{\ell}: \hat{\mathbb{C}} \rightarrow \hat{\mathbb{C}}$  en posant  $\sigma_{\ell}(\infty) = \infty$ , et s'interprète comme l'inversion de cercle  $\ell \subset \hat{\mathbb{C}}$  passant par  $\infty$ .

Au final, les propositions III.1, III.2 et III.4 peuvent se reformuler de la manière suivante.

**Théorème III.6**

Étant donné un cercle  $\mathcal{C}$  dans  $\hat{\mathbb{C}}$ , il existe une unique application  $i_{\mathcal{C}}: \hat{\mathbb{C}} \rightarrow \hat{\mathbb{C}}$  qui satisfait les propriétés suivantes :

- (i)  $i_{\mathcal{C}}(P) = P$  pour tout  $P \in \mathcal{C}$  ;
- (ii)  $i_{\mathcal{C}}$  échange les deux côtés de  $\mathcal{C}$  ;
- (iii)  $i_{\mathcal{C}}$  laisse globalement invariants les cercles orthogonaux à  $\mathcal{C}$ .

De plus, l'application  $i_{\mathcal{C}}$  est une involution, envoie les cercles sur des cercles, et renverse les angles orientés.  $\square$

Mais il n'y a pas de miracle en mathématique, ou du moins, chaque miracle a une explication. La voici. Soit  $\mathbb{S}^2 \subset \mathbb{R}^3$  une sphère de dimension 2,  $N \in \mathbb{S}^2$  son pôle nord, et  $\mathbb{C}$  son plan équatorial. La *projection stéréographique* correspondante est l'application  $p: \mathbb{S}^2 \setminus \{N\} \rightarrow \mathbb{C}$  qui envoie un point  $P$  sur l'intersection de la droite  $\overline{NP}$  avec le plan équatorial  $\mathbb{C}$ . (Voir figure III.4.) Cette projection s'étend naturellement en une application bijective  $p: \mathbb{S}^2 \rightarrow \hat{\mathbb{C}}$  en posant  $p(N) = \infty$ .

On peut vérifier relativement facilement que  $p: \mathbb{S}^2 \rightarrow \hat{\mathbb{C}}$  est une bijection conforme qui envoie les cercles ne passant pas par  $N$  sur des cercles dans  $\mathbb{C}$ , et les cercles passant par  $N$  sur des droites dans  $\mathbb{C}$ . Ainsi, via cette bijection, chaque "cercle dans  $\hat{\mathbb{C}}$ " est un véritable cercle dans  $\mathbb{S}^2$ .

De plus, l'inversion  $i_{\mathcal{C}}: \hat{\mathbb{C}} \rightarrow \hat{\mathbb{C}}$  de cercle équatorial  $\mathcal{C}$  se traduit via la bijection  $p: \mathbb{S}^2 \rightarrow \hat{\mathbb{C}}$  en l'application  $p^{-1} \circ i_{\mathcal{C}} \circ p: \mathbb{S}^2 \rightarrow \mathbb{S}^2$  :

$$\begin{array}{ccc} \mathbb{S}^2 & \longrightarrow & \mathbb{S}^2 \\ p \downarrow & & \downarrow p \\ \hat{\mathbb{C}} & \xrightarrow{i_{\mathcal{C}}} & \hat{\mathbb{C}}. \end{array}$$

Quelle est-elle ? Par le théorème III.6 et l'observation ci-dessus, nous savons que cette application est uniquement déterminée par les propriétés suivantes : elle fixe le cercle équatorial  $\mathcal{C}$  point par point, échange les deux hémisphères de  $\mathbb{S}^2$ , et préserve globalement les cercles orthogonaux à  $\mathcal{C}$ . Il s'agit donc tout simplement de la réflexion de  $\mathbb{S}^2$  par rapport à son plan équatorial.

## III.2 Transformations de Möbius et birapport

Selon Félix Klein, une géométrie est l'étude des propriétés invariantes par l'action d'un groupe  $G$  sur un ensemble  $X$ . En géométrie euclidienne, par exemple, on considère l'action du groupe  $\text{Isom}(\mathbb{R}^2)$  sur  $\mathbb{R}^2$ . Rappelons que ce groupe est engendré par les réflexions  $\sigma_{\ell}$ , et qu'une isométrie  $f$  du plan préserve l'orientation (ce qui est noté  $f \in \text{Isom}^+(\mathbb{R}^2)$ ) si et seulement si elle est composition d'un nombre pair de réflexions.

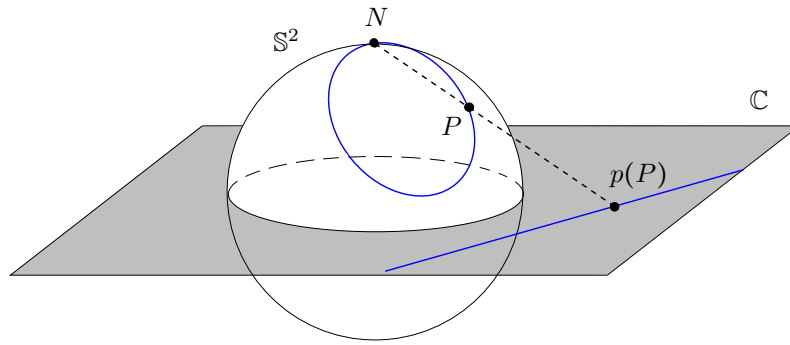


FIGURE III.4 – La projection stéréographique.

Comme on l'a vu dans la section précédente, les réflexions peuvent être interprétées comme des cas particuliers d'inversions : ce sont les inversions par des cercles passant par  $\infty$ . Ainsi, il est naturel de s'intéresser à l'action sur  $\hat{\mathbb{C}}$  du groupe engendré par toutes les inversions. C'est l'objet de la *géométrie de Möbius*.

Comme dans le cas de la géométrie projective, le groupe de transformations est élargi. Par conséquent, le nombre de notions géométriques invariantes va diminuer. (Par le théorème III.6, nous savons que les notions d'angle et de cercle en sont des exemples.) Néanmoins, il sera possible de donner des démonstrations plus faciles d'énoncés ne faisant intervenir que ces notions invariantes : c'est exactement ce que l'on vient de faire avec le porisme de Steiner.

Dans cette section, on va déterminer le groupe engendré par les inversions, puis utiliser ce résultat pour découvrir une nouvelle notion géométrique invariante : le *birapport*. Cette notion est d'une importance capitale en géométrie hyperbolique.

### III.2.1 Le groupe de Möbius

Soit  $Möb(2)$  l'ensemble de toutes les bijections  $\hat{\mathbb{C}} \rightarrow \hat{\mathbb{C}}$  qui sont compositions d'inversions. Il s'agit clairement d'un groupe, le sous-groupe du groupe symétrique  $S(\hat{\mathbb{C}})$  engendré par les inversions, et ce groupe agit naturellement sur l'ensemble  $\hat{\mathbb{C}}$ . On notera  $Möb^+(2)$  l'ensemble des éléments de  $Möb(2)$  qui peuvent s'écrire comme composition d'un nombre pair d'inversions ; c'est un sous-groupe de  $Möb(2)$ .

#### Définition III.1

Les éléments de  $Möb(2)$  sont appelées les **transformations de Möbius**, et le groupe  $Möb^+(2)$  est appelé le **groupe de Möbius**.

Notre premier but sera de déterminer le groupe de Möbius. Cela requiert la définition suivante. Étant donnée une matrice inversible  $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \in GL(2, \mathbb{C})$ ,

soit  $\varphi_A: \hat{\mathbb{C}} \rightarrow \hat{\mathbb{C}}$  l'application définie par

$$\varphi_A(z) = \begin{cases} \frac{az+b}{cz+d} & \text{si } z \neq \infty, z \neq -\frac{d}{c}, \\ \infty & \text{si } z = -\frac{d}{c}, \\ \frac{a}{c} & \text{si } z = \infty. \end{cases}$$

Notons que cette définition est naturelle, dans le sens où  $\frac{az+b}{cz+d}$  diverge lorsque  $z$  tend vers  $-\frac{d}{c}$  et converge vers  $\frac{a}{c}$  lorsque  $z$  tend vers l'infini.

Voyons quelques exemples explicites de telles transformations.

### Exemples.

1. La matrice  $A = \begin{pmatrix} 1 & b \\ 0 & 1 \end{pmatrix}$  définit l'application  $\varphi_A(z) = z + b$ ; il s'agit donc simplement de la translation par le vecteur  $b \in \mathbb{C}$ .
2. Si  $A = \begin{pmatrix} a & 0 \\ 0 & d \end{pmatrix}$ , l'application correspondante est  $\varphi_A(z) = \frac{a}{d}z$ , la multiplication par  $\frac{a}{d} \in \mathbb{C}^*$ .
3. La matrice  $A = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$  induit l'application  $\varphi_A(z) = \frac{1}{z}$ .

Voici à présent le résultat principal du paragraphe.

### Théorème III.7: Groupe de Möbius

L'application  $A \mapsto \varphi_A$  définit un homomorphisme de groupes surjectif

$$\varphi: \mathrm{SL}(2, \mathbb{C}) \longrightarrow \mathrm{Möb}^+(2)$$

dont le noyau est  $\ker \varphi = \{\pm I\}$ .

Ce résultat appelle quelques remarques.

### Remarques.

1. Par les résultats d'Algèbre I, cela implique que  $\varphi$  induit un isomorphisme entre le groupe de Möbius et le groupe quotient  $\mathrm{SL}(2, \mathbb{C})/\{\pm I\}$ . Ce groupe est habituellement noté  $\mathrm{PSL}(2, \mathbb{C})$ , et appelé le *groupe projectif spécial linéaire* de degré 2 sur  $\mathbb{C}$ . En résumé, le théorème ci-dessus nous dit que l'application  $A \mapsto \varphi_A$  définit un isomorphisme de groupes  $\mathrm{PSL}(2, \mathbb{C}) \simeq \mathrm{Möb}^+(2)$ .
2. On vérifie facilement que le groupe  $\mathrm{Möb}(2)$  est le produit semi-direct de  $\mathrm{Möb}^+(2)$  et de  $C_2 = \{\mathrm{id}, i_{\mathcal{C}}\}$ , où  $i_{\mathcal{C}}$  est une inversion quelconque. Ainsi, on obtient l'isomorphisme  $\mathrm{Möb}(2) \simeq \mathrm{PSL}(2, \mathbb{C}) \rtimes C_2$ .

La démonstration du théorème repose sur deux lemmes.

**Lemme III.8.** *Le groupe  $\mathrm{GL}(2, \mathbb{C})$  est engendré par les matrices  $\begin{pmatrix} 1 & b \\ 0 & 1 \end{pmatrix}$ ,  $\begin{pmatrix} a & 0 \\ 0 & d \end{pmatrix}$ ,  $\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$  avec  $b \in \mathbb{C}$  et  $a, d \in \mathbb{C}^*$ .*

*Démonstration.* Cet énoncé n'est en vérité qu'une traduction en théorie des groupes d'un résultat standard en algèbre linéaire, à savoir le fait que toute matrice  $2 \times 2$  à coefficients complexes est diagonalisable via des opérations élémentaires. Plus précisément, toute matrice  $A \in \mathrm{GL}(2, \mathbb{C})$  peut-être diagonalisée (i.e. transformée en une matrice de la forme  $\begin{pmatrix} a & 0 \\ 0 & d \end{pmatrix}$  avec  $a, d \in \mathbb{C}^*$ ) via addition d'un multiple d'une colonne à l'autre colonne (i.e. multiplication à droite par  $\begin{pmatrix} 1 & b \\ 0 & 1 \end{pmatrix}$  ou  $\begin{pmatrix} 1 & 0 \\ b & 1 \end{pmatrix}$ ,  $b \in \mathbb{C}$ ), addition d'un multiple d'une ligne à l'autre ligne (i.e. multiplication à gauche par  $\begin{pmatrix} 1 & b \\ 0 & 1 \end{pmatrix}$  ou  $\begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$ ,  $b \in \mathbb{C}$ ), et échange des lignes ou des colonnes (i.e. multiplication à gauche ou à droite par  $\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$ ). Comme la matrice  $\begin{pmatrix} 1 & 0 \\ b & 1 \end{pmatrix}$  peut s'écrire comme  $\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} 1 & b \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$ , on obtient donc : étant donnée  $A \in \mathrm{GL}(2, \mathbb{C})$ , il existe des matrices  $P_1, \dots, P_n$  et  $Q_1, \dots, Q_m$  de la forme  $\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$  ou  $\begin{pmatrix} 1 & b \\ 0 & 1 \end{pmatrix}$  et une matrice  $D$  de la forme  $\begin{pmatrix} a & 0 \\ 0 & d \end{pmatrix}$  telles que  $P_1 \cdots P_n A Q_1 \cdots Q_m = D$ . Suit l'identité

$$A = P_n^{-1} \cdots P_1^{-1} D Q_m^{-1} \cdots Q_1^{-1}.$$

Comme  $\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}^{-1} = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$  et  $\begin{pmatrix} 1 & b \\ 0 & 1 \end{pmatrix}^{-1} = \begin{pmatrix} 1 & -b \\ 0 & 1 \end{pmatrix}$ , cela démontre le lemme.  $\square$

**Lemme III.9.** *Pour tout cercle  $\mathcal{C}$  dans  $\hat{\mathbb{C}}$ , l'inversion  $i_{\mathcal{C}}: \hat{\mathbb{C}} \rightarrow \hat{\mathbb{C}}$  est de la forme  $z \mapsto \frac{a\bar{z}+b}{c\bar{z}+d}$  avec  $\begin{pmatrix} a & b \\ c & d \end{pmatrix} \in \mathrm{GL}(2, \mathbb{C})$ . De plus, si le cercle  $\mathcal{C}$  est orthogonal à la droite réelle  $\mathbb{R}$ , alors on peut choisir  $\begin{pmatrix} a & b \\ c & d \end{pmatrix} \in \mathrm{GL}(2, \mathbb{R})$ .*

*Démonstration.* Supposons tout d'abord que  $\mathcal{C}$  est un cercle dans  $\mathbb{C}$ , disons centré en  $a \in \mathbb{C}$  et de rayon  $r$ . Alors, on a vu au paragraphe III.1.1 que l'inversion de cercle  $\mathcal{C}$  est donnée par la formule

$$z \mapsto i_{\mathcal{C}}(z) = \frac{r^2}{\bar{z} - \bar{a}} + a = \frac{a\bar{z} + (r^2 - |a|^2)}{\bar{z} - \bar{a}}.$$

Comme  $\det \begin{pmatrix} a & r^2 - |a|^2 \\ 1 & -\bar{a} \end{pmatrix} = -r^2$  est non-nul, ce cas est vérifié. De plus,  $\mathcal{C}$  est orthogonal à  $\mathbb{R}$  si et seulement si  $a \in \mathbb{R}$ , auquel cas cette matrice est réelle. Si  $\mathcal{C}$  passe par le point  $\infty$ , i.e., si  $\mathcal{C}$  est une droite  $\ell$ , alors la réflexion  $i_{\mathcal{C}} = \sigma_{\ell}$  peut aussi s'écrire sous cette forme. La vérification est laissée en exercice.  $\square$

*Démonstration du théorème III.7.* Pour  $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$  et  $A' = \begin{pmatrix} a' & b' \\ c' & d' \end{pmatrix}$  dans  $\mathrm{GL}(2, \mathbb{C})$ , on a

$$(\varphi_A \circ \varphi_{A'})(z) = \varphi_A \left( \frac{a'z + b'}{c'z + d'} \right) = \frac{a \frac{a'z + b'}{c'z + d'} + b}{c \frac{a'z + b'}{c'z + d'} + d} = \frac{(aa' + bc')z + (ab' + bd')}{(ca' + dc')z + (cb' + dd')} = \varphi_{AA'}(z).$$

En particulier,  $\varphi_A \circ \varphi_{A^{-1}} = \varphi_I = \mathrm{id}_{\hat{\mathbb{C}}}$ , et  $\varphi_A: \hat{\mathbb{C}} \rightarrow \hat{\mathbb{C}}$  est une application bijective. Ainsi l'application  $A \mapsto \varphi_A$  définit un homomorphisme  $\varphi: \mathrm{GL}(2, \mathbb{C}) \rightarrow S(\hat{\mathbb{C}})$ .

Considérons à présent la restriction de cet homomorphisme au sous-groupe  $\mathrm{SL}(2, \mathbb{C})$ , et soit  $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$  un élément de son noyau. Cela signifie que  $\varphi_A = \mathrm{id}_{\hat{\mathbb{C}}}$ , c'est-à-dire, que  $\frac{az+b}{cz+d} = z$  pour tout  $z \in \hat{\mathbb{C}}$ . On a alors  $cz^2 + (d-a)z - b = 0$  pour tout  $z \in \mathbb{C}$ , ce qui implique  $b = c = 0$  et  $d = a$ . Mais puisque  $A$  est élément du groupe spécial linéaire,  $1 = \det(A) = a^2$ . Il suit que  $a = \pm 1$ , et donc, que  $\ker \varphi = \{\pm I\}$ .

Reste à montrer que l'image de l'homomorphisme  $\varphi$  n'est autre que le groupe de Möbius. Notons tout d'abord que par définition,  $\varphi_A$  et  $\varphi_{\lambda A}$  coïncident pour tout scalaire  $\lambda \in \mathbb{C}^*$ . En particulier, étant donnée une matrice  $A \in \mathrm{GL}(2, \mathbb{C})$ , il existe une matrice  $A' \in \mathrm{SL}(2, \mathbb{C})$  telle que  $\varphi_A = \varphi_{A'}$ . (Il suffit de choisir  $A' = \lambda A$  avec  $\lambda \in \mathbb{C}^*$  tel que  $\lambda^2 \det(A) = 1$ .) Cela implique que  $\varphi(\mathrm{SL}(2, \mathbb{C})) = \varphi(\mathrm{GL}(2, \mathbb{C}))$ , et il reste à voir que ce dernier groupe est exactement  $M\ddot{o}b^+(2)$ .

Montrons tout d'abord que  $\varphi(\mathrm{GL}(2, \mathbb{C}))$  est inclus dans le groupe de Möbius, i.e. que pour tout  $A \in \mathrm{GL}(2, \mathbb{C})$ ,  $\varphi_A$  est composition d'un nombre pair d'inversions. Par le lemme III.8, il suffit de le vérifier pour les matrices  $\begin{pmatrix} 1 & b \\ 0 & 1 \end{pmatrix}$ ,  $\begin{pmatrix} a & 0 \\ 0 & d \end{pmatrix}$ , et  $\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$ . Comme on l'a vu plus haut en exemple, ces matrices correspondent à la translation par le vecteur  $b \in \mathbb{C}$ , à l'homothétie de facteur  $\lambda = a/d \in \mathbb{C}^*$ , et à la transformation  $z \mapsto \frac{1}{z}$ , respectivement. Ces trois transformations sont compositions d'un nombre pair d'inversions. En effet, on sait bien qu'une translation est la composition de deux réflexions par des droites parallèles, i.e. de deux inversions. Quant à la multiplication par  $\lambda = r^2 e^{i\vartheta} \in \mathbb{C}^*$ , c'est la composition de la rotation d'angle  $\vartheta$  autour de l'origine (qui est la composition de deux inversions par des droites qui se coupent en l'origine) et d'une homothétie de facteur  $r^2$  (que l'on peut écrire comme la composition  $i_{\mathbb{S}^1} \circ i_{1/r\mathbb{S}^1}$ ). Finalement, l'application  $z \mapsto \frac{1}{z}$  peut s'écrire  $\sigma_{\mathbb{R}} \circ i_{\mathbb{S}^1}$  ; elle est donc elle aussi la composition d'un nombre pair d'inversions.

Il reste encore à démontrer l'inclusion  $M\ddot{o}b^+(2) \subset \varphi(\mathrm{GL}(2, \mathbb{C}))$ , i.e. que la composition de deux inversions est toujours de la forme  $\varphi_A$  pour un certain  $A \in \mathrm{GL}(2, \mathbb{C})$ . Mais par le lemme III.9, toute inversion est de la forme  $i_{\mathcal{C}}(z) = \frac{a\bar{z}+b}{c\bar{z}+d}$  avec  $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \in \mathrm{GL}(2, \mathbb{C})$ . Or, si  $i_{\mathcal{C}'}(z) = \frac{a'\bar{z}+b'}{c'\bar{z}+d'}$  avec  $A' = \begin{pmatrix} a' & b' \\ c' & d' \end{pmatrix} \in \mathrm{GL}(2, \mathbb{C})$ , une légère modification du calcul fait au début de la preuve donne

$$(i_{\mathcal{C}} \circ i_{\mathcal{C}'})(z) = \frac{(a\bar{a}' + b\bar{c}')z + (a\bar{b}' + b\bar{d}')}{(c\bar{a}' + d\bar{c}')z + (c\bar{b}' + d\bar{d}')} = \varphi_{A\bar{A}'}(z).$$

Comme  $\det(A\bar{A}') = \det(A)\det(\bar{A}') \neq 0$ , cela conclut la preuve du théorème.  $\square$

### III.2.2 Le birapport

Par les résultats de la section 3.1, on sait que les inversions préservent les cercles et les angles (non-orientés) dans  $\hat{\mathbb{C}}$ . Il en va de même pour toute transformation de Möbius, puisque ces dernières sont par définition des compositions d'inversions. Aussi, les éléments du groupe de Möbius  $M\ddot{o}b^+(2) = \{\varphi_A \mid A \in \mathrm{GL}(2, \mathbb{C})\}$  préservent les cercles et les angles orientés dans  $\hat{\mathbb{C}}$ .

Existe-t-il d'autres notions géométriques en géométrie de Möbius, i.e. d'autres objets invariants par l'action du groupe de Möbius ? Le but de ce paragraphe est d'utiliser le théorème III.7 pour apporter une réponse positive à cette question.

#### Définition III.2

Étant donnés  $z_1, z_2, z_3, z_4 \in \hat{\mathbb{C}} = \mathbb{C} \cup \{\infty\}$  avec  $z_1 \neq z_4$  et  $z_2 \neq z_3$ , leur

**birapport** est le nombre complexe

$$[z_1, z_2; z_3, z_4] = \frac{(z_1 - z_3)(z_2 - z_4)}{(z_2 - z_3)(z_1 - z_4)} \in \mathbb{C},$$

avec la convention que  $[\infty, z_2; z_3, z_4] = \frac{z_2 - z_4}{z_2 - z_3}$ , et similairement pour les autres coordonnées.

Notons que cette convention est naturelle, puisque  $\frac{z_2 - z_4}{z_2 - z_3}$  est la limite du birapport  $[z_1, z_2; z_3, z_4]$  lorsque  $z_1$  tend vers l'infini.

**Remarques.**

1. Si  $z_1 = z_2$ , on obtient  $[z, z; z_3, z_4] = \frac{(z - z_3)(z - z_4)}{(z - z_3)(z - z_4)} = 1$  indépendamment de la valeur de  $z_3$  et de  $z_4$ . Similairement,  $[z_1, z_2; z, z] = 1$  pour tous  $z_1, z_2$ .
2. Par définition, le birapport  $[z_1, z_2; z_3, z_4]$  est égal au birapport  $[z_2, z_1; z_4, z_3]$ .

### Théorème III.10: Invariance du birapport

Pour tous  $z_1, z_2, z_3, z_4 \in \hat{\mathbb{C}}$  avec  $z_1 \neq z_4$  et  $z_2 \neq z_3$  et pour tout  $\varphi \in \text{Möb}^+(2)$ ,

$$[\varphi(z_1), \varphi(z_2); \varphi(z_3), \varphi(z_4)] = [z_1, z_2; z_3, z_4].$$

De plus, pour tout  $\varphi \in \text{Möb}(2) \setminus \text{Möb}^+(2)$ ,

$$[\varphi(z_1), \varphi(z_2); \varphi(z_3), \varphi(z_4)] = \overline{[z_1, z_2; z_3, z_4]}.$$

*Démonstration.* Supposons d'abord que  $\varphi$  est un élément du groupe de Möbius  $\text{Möb}^+(2)$ . Par le théorème III.7 et le lemme III.8, il suffit de vérifier l'invariance du birapport par l'action de  $\varphi_A$  avec  $A = \begin{pmatrix} 1 & b \\ 0 & 1 \end{pmatrix}$ ,  $\begin{pmatrix} a & 0 \\ 0 & d \end{pmatrix}$ , et  $\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$ . En d'autres termes, il nous faut vérifier que le birapport est laissé invariant par la translation  $z \mapsto z + b$ , par l'application  $z \mapsto z^{-1}$  et par la multiplication  $z \mapsto \frac{a}{d}z$ . Cela se montre de manière immédiate.

Soit à présent  $\varphi \in \text{Möb}(2) \setminus \text{Möb}^+(2)$ . Alors, si l'on note par  $\sigma$  la conjugaison complexe  $\sigma(z) = \bar{z}$ , la transformation  $\sigma \circ \varphi$  est un élément du groupe de Möbius  $\text{Möb}^+(2)$ . Ainsi, par la première partie, on obtient :

$$\begin{aligned} [z_1, z_2; z_3, z_4] &= [(\sigma \circ \varphi)(z_1), (\sigma \circ \varphi)(z_2); (\sigma \circ \varphi)(z_3), (\sigma \circ \varphi)(z_4)] \\ &= [\overline{\varphi(z_1)}, \overline{\varphi(z_2)}; \overline{\varphi(z_3)}, \overline{\varphi(z_4)}] = \overline{[\varphi(z_1), \varphi(z_2); \varphi(z_3), \varphi(z_4)]}, \end{aligned}$$

ce qui conclut la démonstration. □

Concluons ce paragraphe par une jolie application de la notion de birapport.

**Proposition III.11.** *Quatre points distincts  $z_1, z_2, z_3, z_4 \in \hat{\mathbb{C}}$  sont sur un cercle dans  $\hat{\mathbb{C}}$  si et seulement si leur birapport est réel.*

*Démonstration.* Soient donc  $z_1, z_2, z_3, z_4 \in \hat{\mathbb{C}}$  quatre points distincts, et considérons la matrice

$$A = \begin{pmatrix} z_2 - z_4 & -z_3(z_2 - z_4) \\ z_2 - z_3 & -z_4(z_2 - z_3) \end{pmatrix}.$$

Notons que comme ces points sont distincts,  $\det(A) = (z_2 - z_3)(z_2 - z_4)(z_3 - z_4)$  est non-nul, si bien que  $A$  est une matrice inversible. Ainsi, on peut considérer l'élément  $\varphi = \varphi_A \in \text{Möb}^+(2)$  associé. Il est donné par la formule  $\varphi(z) = \frac{(z - z_3)(z_2 - z_4)}{(z - z_4)(z_2 - z_3)}$ , et satisfait donc  $\varphi(z_2) = 1$ ,  $\varphi(z_3) = 0$ ,  $\varphi(z_4) = \infty$  et  $\varphi(z_1) = [z_1, z_2; z_3, z_4]$ . Comme  $\varphi$  préserve les cercles dans  $\hat{\mathbb{C}}$ ,  $z_1, z_2, z_3$  et  $z_4 \in \hat{\mathbb{C}}$  sont sur un cercle si et seulement si  $\varphi(z_1), \varphi(z_2), \varphi(z_3)$  et  $\varphi(z_4)$  sont sur un cercle, i.e. si  $[z_1, z_2; z_3, z_4]$ , 1, 0 et  $\infty$  sont sur un cercle dans  $\hat{\mathbb{C}}$ . Comme les cercles passant par  $\infty$  sont les droites dans  $\mathbb{C}$ , c'est le cas si et seulement si  $[z_1, z_2; z_3, z_4]$  est sur la droite passant par 0 et 1, à savoir la droite réelle.  $\square$

### III.3 Le disque de Poincaré

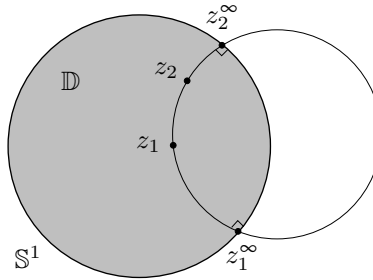
Nous allons enfin utiliser les résultats des deux sections précédentes pour définir un modèle de la géométrie hyperbolique.

#### III.3.1 La distance hyperbolique

Dans tout ce paragraphe, on notera  $\mathbb{D}$  le disque ouvert centré en l'origine de rayon 1, et  $\mathbb{S}^1$  son bord :

$$\mathbb{D} = \{z \in \mathbb{C} \mid |z| < 1\} \quad \text{et} \quad \mathbb{S}^1 = \{z \in \mathbb{C} \mid |z| = 1\}.$$

Le disque  $\mathbb{D}$  sera l'ensemble des points de notre modèle de la géométrie hyperbolique. Étant donnés deux points  $z_1 \neq z_2 \in \mathbb{D}$ , on vérifie facilement (à l'aide d'une inversion bien choisie) qu'il existe un unique cercle dans  $\hat{\mathbb{C}}$  passant par  $z_1$  et  $z_2$  orthogonal à  $\mathbb{S}^1$ . Ce cercle intersecte  $\mathbb{S}^1$  en deux points  $z_1^\infty, z_2^\infty$ , avec  $z_1^\infty$  du côté de  $z_1$  et  $z_2^\infty$  du côté de  $z_2$ , comme illustré ci-dessous.



#### Définition III.3

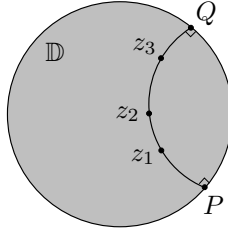
La **distance hyperbolique** entre  $z_1$  et  $z_2$  est définie par

$$d_{\mathbb{D}}(z_1, z_2) = -\log[z_1, z_2; z_1^\infty, z_2^\infty].$$

Avant de démontrer qu'il s'agit bel et bien d'une distance, quelques remarques et exemples s'imposent.

**Remarques.**

1. Si  $0$ ,  $z_1$  et  $z_2$  sont alignés, alors le cercle en question est une droite, la droite passant par  $z_1$  et  $z_2$ .
2. Par construction, les quatre points  $z_1$ ,  $z_2$ ,  $z_1^\infty$  et  $z_2^\infty$  sont sur un cercle ; par la proposition III.11, leur birapport est donc réel.
3. Si  $z_1$  et  $z_2$  coïncident, alors on ne peut pas définir  $z_1^\infty$  et  $z_2^\infty$ . Néanmoins, le birapport  $[z_1, z_2; z_3, z_4]$  est alors toujours égal à 1, comme on l'a vu ci-dessus. Ainsi, la distance hyperbolique entre deux points égaux est bien définie, et nulle comme il se doit.
4. Supposons que  $z_1, z_2$  et  $z_3$  sont trois points sur un même cercle orthogonal à  $\mathbb{S}^1$ , et soient  $P, Q$  les points d'intersection avec  $\mathbb{S}^1$  comme ci-dessous.



Alors, on a

$$[z_1, z_3; P, Q] = \frac{(z_1 - P)(z_2 - Q)(z_2 - P)(z_3 - Q)}{(z_2 - P)(z_1 - Q)(z_3 - P)(z_2 - Q)} = [z_1, z_2; P, Q][z_2, z_3; P, Q],$$

d'où  $d_{\mathbb{D}}(z_1, z_3) = d_{\mathbb{D}}(z_1, z_2) + d_{\mathbb{D}}(z_2, z_3)$ . La distance hyperbolique est ainsi additive le long des arcs de cercles orthogonaux à  $\mathbb{S}^1$ . Ces derniers jouent donc le rôle des droites dans cette géométrie ; on y reviendra.

**Exemple.** Calculons la distance hyperbolique entre l'origine et  $t \in (0, 1)$ . Dans ce cas, le cercle passant par  $z_1 = 0$  et  $z_2 = t$  n'est autre que la droite réelle, d'où  $z_1^\infty = -1$  et  $z_2^\infty = 1$ . On a donc

$$d_{\mathbb{D}}(0, t) = -\log[0, t; -1, 1] = -\log\left(\frac{(0 - (-1))(t - 1)}{(t - (-1))(0 - 1)}\right) = -\log\left(\frac{1 - t}{1 + t}\right).$$

Notons que la fonction  $f: [0, 1) \rightarrow [0, 1)$  donnée par  $f(t) = \frac{1-t}{1+t}$  est une fonction monotone décroissante qui tend vers 0 lorsque  $t$  tend vers 1. Ainsi, la fonction  $t \mapsto d_{\mathbb{D}}(0, t)$  est monotone croissante et tend vers l'infini quand  $t$  tend vers 1.

Pour démontrer que la distance hyperbolique est bien une distance, on aura besoin du lemme suivant.

**Lemme III.12.** (i) Si  $\varphi \in \text{Möb}(2)$  est composition d'inversions par des cercles orthogonaux à  $\mathbb{S}^1$ , alors  $\varphi(\mathbb{D}) \subset \mathbb{D}$  et  $d_{\mathbb{D}}(\varphi(z_1), \varphi(z_2)) = d_{\mathbb{D}}(z_1, z_2)$  pour tous  $z_1, z_2 \in \mathbb{D}$ .

(ii) Pour tous  $z_1, z_2 \in \mathbb{D}$ , il existe un cercle  $\mathcal{C}$  dans  $\hat{\mathbb{C}}$  orthogonal à  $\mathbb{S}^1$  tel que  $i_{\mathcal{C}}(z_1) = z_2$ .

*Démonstration.* Soit  $\mathcal{C}$  un cercle orthogonal à  $\mathbb{S}^1$ ; on veut vérifier que  $i_{\mathcal{C}}(\mathbb{D}) \subset \mathbb{D}$  et  $d_{\mathbb{D}}(i_{\mathcal{C}}(z_1), i_{\mathcal{C}}(z_2)) = d_{\mathbb{D}}(z_1, z_2)$ . Cela impliquera immédiatement le même résultat pour toute composition de telles inversions, i.e. le premier point du lemme.

Pour démontrer que  $i_{\mathcal{C}}$  envoie  $\mathbb{D}$  sur lui-même, on peut donner une démonstration analytique qui est laissée en exercice. Voici une autre preuve, de nature topologique. Comme  $\mathbb{S}^1$  est orthogonal au cercle  $\mathcal{C}$ , l'inversion  $i_{\mathcal{C}}$  satisfait  $i_{\mathcal{C}}(\mathbb{S}^1) = \mathbb{S}^1$ . Ainsi, la restriction de  $i_{\mathcal{C}}$  au complémentaire de  $\mathbb{S}^1$  dans  $\hat{\mathbb{C}}$  est une application continue de l'union disjointe  $\mathbb{D} \sqcup \hat{\mathbb{C}} \setminus \mathbb{D}$  dans elle-même. Par continuité de  $i_{\mathcal{C}}$ , soit  $i_{\mathcal{C}}(\mathbb{D}) \subset \mathbb{D}$ , soit  $i_{\mathcal{C}}(\mathbb{D}) \subset \hat{\mathbb{C}} \setminus \mathbb{D}$ .<sup>2</sup> Mais tout élément de  $\mathbb{D} \cap \mathcal{C}$  est fixé par  $i_{\mathcal{C}}$ ; ainsi,  $i_{\mathcal{C}}$  envoie  $\mathbb{D}$  sur lui-même.

Soient à présent  $z_1, z_2 \in \mathbb{D}$ , que l'on peut supposer distincts sans restreindre la généralité. Comme  $i_{\mathcal{C}}$  préserve les angles, préserve les cercles, et envoie  $\mathbb{S}^1$  sur lui-même, cette inversion envoie tout cercle orthogonal à  $\mathbb{S}^1$  sur un cercle orthogonal à  $\mathbb{S}^1$ . Ainsi,  $i_{\mathcal{C}}$  envoie l'unique cercle passant par  $z_1$  et  $z_2$  orthogonal à  $\mathbb{S}^1$  sur l'unique cercle passant par  $i_{\mathcal{C}}(z_1)$  et  $i_{\mathcal{C}}(z_2)$  orthogonal à  $\mathbb{S}^1$ . Il suit que  $i_{\mathcal{C}}(z_1^{\infty}) = i_{\mathcal{C}}(z_1)^{\infty}$  et  $i_{\mathcal{C}}(z_2^{\infty}) = i_{\mathcal{C}}(z_2)^{\infty}$ . En utilisant ces égalités ainsi que le théorème III.10 et la proposition III.11, on obtient

$$\begin{aligned} d_{\mathbb{D}}(i_{\mathcal{C}}(z_1), i_{\mathcal{C}}(z_2)) &= -\log[i_{\mathcal{C}}(z_1), i_{\mathcal{C}}(z_2); i_{\mathcal{C}}(z_1)^{\infty}, i_{\mathcal{C}}(z_2)^{\infty}] \\ &= -\log[i_{\mathcal{C}}(z_1), i_{\mathcal{C}}(z_2); i_{\mathcal{C}}(z_1^{\infty}), i_{\mathcal{C}}(z_2^{\infty})] \\ &= -\log[z_1, z_2; z_1^{\infty}, z_2^{\infty}] = -\log[z_1, z_2; z_1^{\infty}, z_2^{\infty}] = d_{\mathbb{D}}(z_1, z_2). \end{aligned}$$

Cela conclut la preuve du premier point du lemme. La démonstration du second est laissée en exercice.  $\square$

### Conséquences du lemme.

1. Toute rotation autour de l'origine peut être écrite comme la composition de deux réflexions par des droites passant par l'origine. En d'autres termes, une rotation autour de l'origine est la composition de deux inversions par des cercles orthogonaux à  $\mathbb{S}^1$ . Par le premier point du lemme, elle laisse donc  $d_{\mathbb{D}}$  invariant :  $d_{\mathbb{D}}(\rho_0^{\alpha}(z_1), \rho_0^{\alpha}(z_2)) = d_{\mathbb{D}}(z_1, z_2)$  pour tous  $z_1, z_2 \in \mathbb{D}$ .
2. En particulier,  $d_{\mathbb{D}}(0, z) = r$  si et seulement si  $d_{\mathbb{D}}(0, \rho_0^{\alpha}(z)) = r$  pour tous  $\alpha$ . Il suit que le cercle hyperbolique centré en l'origine de rayon  $r$ , i.e. l'ensemble  $\{z \in \mathbb{D} \mid d_{\mathbb{D}}(0, z) = r\}$ , est simplement un cercle euclidien centré en l'origine (de rayon  $t$  tel que  $r = -\log\left(\frac{1-t}{1+t}\right)$ , i.e.  $t = \frac{1-e^{-r}}{1+e^{-r}}$ ).
3. Cela implique que *tous les cercles hyperboliques sont des cercles euclidiens*. En effet, considérons un cercle hyperbolique quelconque, c'est-à-dire un ensemble de la forme  $S = \{z \in \mathbb{D} \mid d_{\mathbb{D}}(z_0, z) = r\}$  avec  $z_0 \in \mathbb{D}$  et  $r > 0$ . Par

2. Il est possible de donner une preuve de ce fait en utilisant le théorème des valeurs intermédiaires. La formalisation de cet argument se fait à l'aide de la notion topologique de *connexité*.

le second point du lemme, il existe un cercle  $\mathcal{C}$  orthogonal à  $\mathbb{S}^1$  tel que  $i_{\mathcal{C}}(z_0) = 0$ . Par le premier point du lemme,

$$i_{\mathcal{C}}(S) = \{i_{\mathcal{C}}(z) \in \mathbb{D} \mid r = d_{\mathbb{D}}(z_0, z) = d_{\mathbb{D}}(0, i_{\mathcal{C}}(z))\} = \{w \in \mathbb{D} \mid d_{\mathbb{D}}(0, w) = r\},$$

qui est un cercle euclidien centré en l'origine par la remarque précédente. Comme  $i_{\mathcal{C}}$  est une involution, il suit que  $S$  est l'image par  $i_{\mathcal{C}}$  d'un cercle euclidien, donc un cercle euclidien ou une droite euclidienne. Ce dernier cas est exclu puisque  $i_{\mathcal{C}}(\mathbb{D}) \subset \mathbb{D}$ . Notons que le cercle euclidien  $S$  ne sera pas centré en  $z_0$  en général. (En fait, il sera centré en  $z_0$  si et seulement si  $z_0 = 0$ .)

4. Pour tous  $z_1, z_2 \in \mathbb{D}$ , il existe une transformation bijective  $\varphi$  de  $\mathbb{D}$  laissant  $d_{\mathbb{D}}$  invariante, telle que  $\varphi(z_1) = 0$  et  $\varphi(z_2) = t \geq 0$ . En effet, par le second point du lemme, on peut trouver une telle transformation  $\psi$  qui envoie  $z_1$  sur l'origine ; il est ensuite possible d'envoyer  $\psi(z_2)$  sur l'axe réel positif via une rotation  $\rho_0^\alpha$  bien choisie. La composition  $\varphi = \rho_0^\alpha \circ \psi$  satisfait les propriétés requises.

Nous sommes enfin prêts à démontrer confortablement que la distance hyperbolique est bel et bien une métrique.

**Proposition III.13.** *La distance hyperbolique  $d_{\mathbb{D}}$  est une métrique sur  $\mathbb{D}$ . De plus,  $d_{\mathbb{D}}(z_1, z_3) = d_{\mathbb{D}}(z_1, z_2) + d_{\mathbb{D}}(z_2, z_3)$  si et seulement si  $z_2$  est sur l'arc de cercle entre  $z_1$  et  $z_3$  orthogonal à  $\mathbb{S}^1$ .*

*Démonstration.* Soient  $z_1, z_2 \in \mathbb{D}$ . Par la remarque 4 ci-dessus, il existe une transformation bijective  $\varphi$  de  $\mathbb{D}$  laissant  $d_{\mathbb{D}}$  invariante, telle que  $\varphi(z_1) = 0$  et  $\varphi(z_2) = t \geq 0$ . Ainsi,

$$d_{\mathbb{D}}(z_1, z_2) = d_{\mathbb{D}}(0, t) = -\log\left(\frac{1-t}{1+t}\right)$$

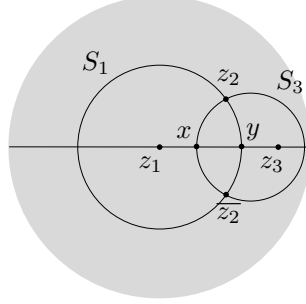
un nombre réel positif ou nul. La formule pour  $d_{\mathbb{D}}$  définit donc bien une application  $d_{\mathbb{D}}: \mathbb{D} \times \mathbb{D} \rightarrow [0, \infty)$ . Par l'égalité ci-dessus,  $d_{\mathbb{D}}(z_1, z_2) = 0$  si et seulement si  $t = 0$ , i.e.  $\varphi(z_2) = \varphi(z_1)$ . Comme  $\varphi$  est bijective, c'est le cas si et seulement si  $z_1 = z_2$ , ce qui démontre l'identité des indiscernables.

La symétrie de la distance hyperbolique découle immédiatement de la propriété de symétrie du birapport vue après sa définition :

$$d_{\mathbb{D}}(z_2, z_1) = -\log[z_2, z_1; z_2^\infty, z_1^\infty] = -\log[z_1, z_2; z_1^\infty, z_2^\infty] = d_{\mathbb{D}}(z_1, z_2).$$

Reste à vérifier l'inégalité du triangle. Soient donc  $z_1, z_2, z_3 \in \mathbb{D}$ , que l'on peut supposer tous distincts sans restreindre la généralité. Via le remarque 4 ci-dessus, on peut se ramener à la situation où  $z_1 = 0$  et  $z_3 = t > 0$ . Traçons le cercle hyperbolique  $S_1$  centré en  $z_1 = 0$  de rayon  $d_{\mathbb{D}}(z_1, z_2)$  et le cercle hyperbolique  $S_3$  centré en  $z_3 = t$  de rayon  $d_{\mathbb{D}}(z_2, z_3)$ . Par les remarques 2 et 3 ci-dessus,  $S_1$  est un cercle euclidien (centré en l'origine), tandis que  $S_3$  est un cercle euclidien (centré sur l'axe réel positif). Par définition,  $z_2$  est dans l'intersection de ces deux cercles. Comme la conjugaison complexe fixe  $z_1$  et  $z_3$ , et laisse  $d_{\mathbb{D}}$  invariante

(par le premier point du lemme III.12), le conjugué complexe  $\overline{z_2}$  est aussi dans l'intersection de  $S_1$  et de  $S_3$ . Ainsi,  $S_1$  et  $S_3$  se coupent en les deux points  $z_2, \overline{z_2}$  si  $z_2$  n'est pas sur l'axe réel, et en l'unique point  $z_2$  si  $z_2$  est réel. On notera encore  $x$  et  $y$  les points d'intersection de  $S_1$  et de  $S_3$  avec l'intervalle  $[z_1, z_3]$ , comme illustré ci-dessous.



Comme on l'a vu en remarque 4 après la définition de la distance hyperbolique, cette dernière est (en particulier) additive le long de l'axe réel. En utilisant ce fait ainsi que la croissance monotone de la fonction  $t \mapsto d_{\mathbb{D}}(0, t)$ , il suit

$$\begin{aligned} d_{\mathbb{D}}(z_1, z_3) &= d_{\mathbb{D}}(0, t) = d_{\mathbb{D}}(0, x) + d_{\mathbb{D}}(x, t) \\ &\leq d_{\mathbb{D}}(0, y) + d_{\mathbb{D}}(y, t) = d_{\mathbb{D}}(z_1, z_2) + d_{\mathbb{D}}(z_2, z_3), \end{aligned}$$

ce qui démontre l'inégalité du triangle. Finalement, il y a égalité si et seulement si  $x = y$ , i.e. si et seulement si  $z_2$  appartient à l'intervalle  $[0, t]$  qui n'est autre que l'arc de cercle entre  $z_1$  et  $z_3$  orthogonal à  $\mathbb{S}^1$ .  $\square$

Le plus dur est fait, et la définition suivante est maintenant parfaitement motivée :

#### Définition III.4

Une **droite hyperbolique** dans  $\mathbb{D}$  est un cercle orthogonal à  $\mathbb{S}^1$ .

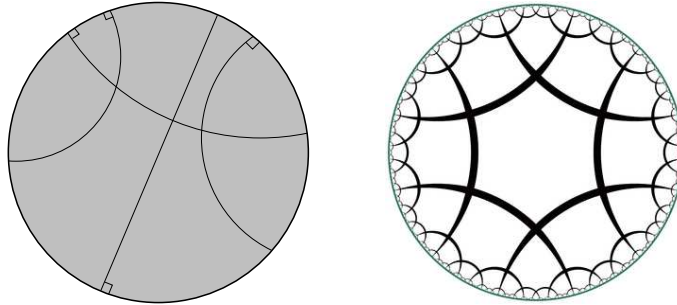


FIGURE III.5 – Quelques exemples de droites hyperboliques dans  $\mathbb{D}$ .

Quelques exemples de telles droites sont données en figure III.5.

### III.3.2 Le disque de Poincaré

Nous pouvons à présent résumer tous les résultats obtenus dans cette section sous la forme suivante.

Le **disque de Poincaré** est le modèle de la géométrie hyperbolique de dimension 2 (on parle du *plan hyperbolique*) donné comme suit.

- Les points de la géométrie sont les éléments du disque  $\mathbb{D} = \{z \in \mathbb{C} \mid |z| < 1\}$ .
- Les droites de la géométrie sont les arcs de cercles dans  $\mathbb{D}$  orthogonaux à  $\mathbb{S}^1 = \{z \in \mathbb{C} \mid |z| = 1\}$ .
- La distance est la distance hyperbolique  $d_{\mathbb{D}}: \mathbb{D} \times \mathbb{D} \rightarrow [0, \infty)$ . Notons qu'elle est additive le long des droites hyperboliques, et que les points  $z$  tendant vers le bord de  $\mathbb{D}$  sont à distance  $d_{\mathbb{D}}(z, 0)$  tendant vers l'infini. Notons encore que les cercles hyperboliques définis par cette métrique sont des cercles euclidiens.
- Les angles sont les angles euclidiens.
- Toutes les compositions d'inversions par des droites hyperboliques sont des isométries de l'espace métrique  $(\mathbb{D}, d_{\mathbb{D}})$ . Elles préservent donc la distance hyperbolique, mais aussi les droites hyperboliques, les cercles hyperboliques, et les angles.

Il est temps de vérifier que cette géométrie satisfait tous les postulats d'Euclide, à l'exception du cinquième (ou de manière équivalente, l'axiome de Playfair).

**Postulat I** Pour tous points  $A \neq B$  dans  $\mathbb{D}$ , il existe un segment de droite les joignant : on l'a vu.

**Postulat II** Un tel segment peut être prolongé indéfiniment en une droite : c'est bien clair, si l'on se rappelle que les points du bord sont à distance infinie.

**Postulat III** Pour tous points  $A \neq B$  dans  $\mathbb{D}$ , il existe un cercle centré en  $A$  et passant par  $B$  : oui, et c'est un cercle euclidien.

**Postulat IV** Tous les angles droits sont égaux, c'est-à-dire isométriques entre eux : fixons un angle droit en un point  $z \in \mathbb{D}$  formé par deux droites hyperboliques  $\ell_1, \ell_2$  ; par la remarque 4 ci-dessus, il existe une isométrie de  $(\mathbb{D}, d_{\mathbb{D}})$  qui envoie  $z$  sur l'origine et  $\ell_1$  sur la droite réelle. La droite  $\ell_2$  sera alors envoyée sur l'axe imaginaire, si bien que l'angle droit arbitraire dont nous sommes partis est isométrique à un angle droit fixé (à l'origine, entre les axes réels et imaginaires).

**Postulat V** Étant donné une droite  $\ell$  et un point  $A \notin \ell$ , il existe une *unique* parallèle à  $\ell$  passant par  $A$  (i.e. une droite  $\ell'$  disjointe de  $\ell$ ) : comme illustré en figure III.6, en géométrie hyperbolique, il y a une infinité de telles droites !

En conclusion, le postulat des parallèles ne peut pas être démontré à partir des quatre premiers, puisqu'il existe un modèle qui satisfait ces quatre postulats sans pour autant satisfaire le cinquième. Aussi, tout ce qu'Euclide démontre à partir

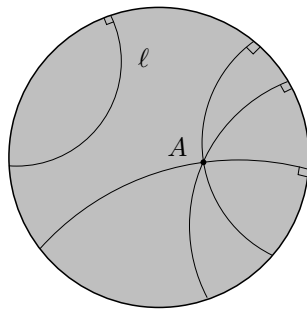


FIGURE III.6 – Le postulat V n'est pas valide en géométrie hyperbolique.

des quatre premiers postulats est valide en géométrie hyperbolique, mais *a priori* pas le reste. Par exemple, la somme des angles d'un triangle hyperbolique (i.e. d'un sous-ensemble de  $\mathbb{D}$  délimité par trois droites hyperboliques) est toujours strictement inférieure à deux angles droits. En fait, deux triangles hyperboliques sont isométriques si et seulement s'ils ont mêmes angles !

Par manque de temps, nous n'aurons malheureusement pas la possibilité d'étudier la géométrie hyperbolique comme nous avons étudié la géométrie euclidienne en première partie. Néanmoins, pour conclure ce chapitre dans l'esprit de Félix Klein, nous allons maintenant nous concentrer sur les isométries du plan hyperbolique.

### III.4 Isométries hyperboliques

Dans l'esprit du programme d'Erlangen mentionné en introduction, et dans la continuité des chapitres précédents, nous allons déterminer le groupe des isométries  $\text{Isom}(\mathbb{D})$  du plan hyperbolique. On comprendra alors la géométrie hyperbolique comme l'étude des notions invariantes par l'action de ce groupe sur  $\mathbb{D}$ . Cette analyse est facilitée par l'introduction d'un nouveau modèle de la géométrie hyperbolique, appelé le *demi-plan de Poincaré*.

#### III.4.1 Le demi-plan de Poincaré

Commençons par définir cet autre modèle du plan hyperbolique. Pas de panique, il existe un dictionnaire explicite (une isométrie) pour passer d'un modèle à l'autre.

Le **demi-plan de Poincaré** est le modèle de la géométrie hyperbolique de dimension 2 donné comme suit.

- Les points sont les éléments du demi-plan supérieur  $\mathbb{H} = \{z \in \mathbb{C} \mid \text{Im}(z) > 0\}$ .
- Les droites hyperboliques sont les cercles dans  $\mathbb{H}$  orthogonaux à  $\mathbb{R}$ . Des exemples sont illustrés en figure III.7.

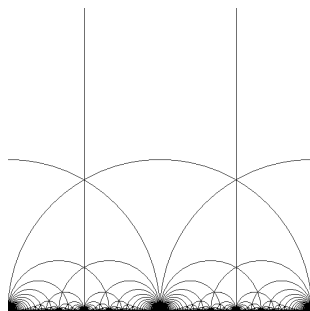
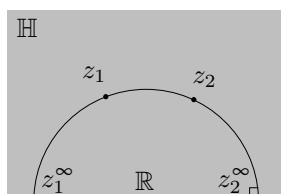


FIGURE III.7 – Des droites hyperboliques dans le demi-plan de Poincaré.

- La distance entre deux points  $z_1, z_2 \in \mathbb{H}$  est  $d_{\mathbb{H}}(z_1, z_2) = -\log[z_1, z_2; z_1^{\infty}, z_2^{\infty}]$ , où  $z_1^{\infty}$  et  $z_2^{\infty}$  sont les intersections avec  $\mathbb{R}$  de l'unique droite hyperbolique passant par  $z_1$  et  $z_2$ , comme illustré ci-dessous.



- Les angles sont les angles euclidiens.

L'identification des deux modèles se fait via ce qu'on appelle la *transformation de Cayley*. Il s'agit de la transformation de Möbius  $\varphi_A: \hat{\mathbb{C}} \rightarrow \hat{\mathbb{C}}$  induite par  $A = \begin{pmatrix} 1 & -i \\ 1 & i \end{pmatrix}$ . On vérifie facilement que  $\text{Im}(z) > 0$  si et seulement si  $|\varphi_A(z)| < 1$ , d'où une bijection  $\varphi_A: \mathbb{H} \rightarrow \mathbb{D}$  qui envoie  $\mathbb{R}$  sur  $\mathbb{S}^1$ . De plus, comme il s'agit d'un élément de  $Möb^+(2)$ ,  $\varphi_A$  préserve les cercles, les angles orientés, et les birapports. Ainsi, la fonction  $d_{\mathbb{H}}$  définie ci-dessus sur  $\mathbb{H}$  n'est autre que la distance hyperbolique sur  $\mathbb{D}$  transportée sur  $\mathbb{H}$  via  $\varphi_A$  :

$$d_{\mathbb{H}}(z_1, z_2) = d_{\mathbb{D}}(\varphi_A(z_1), \varphi_A(z_2)) \quad \text{pour } z_1, z_2 \in \mathbb{H}.$$

En particulier, il s'agit bien d'une distance, et l'on peut formuler le résultat de cette discussion de la manière suivante.

**Proposition III.14.** *La transformation de Cayley  $z \mapsto \frac{z-i}{z+i}$  définit une isométrie de l'espace métrique  $(\mathbb{H}, d_{\mathbb{H}})$  sur l'espace métrique  $(\mathbb{D}, d_{\mathbb{D}})$ , qui préserve les droites hyperboliques, les angles, et les cercles.*  $\square$

En résumé,  $\mathbb{D}$  et  $\mathbb{H}$  ne sont que deux modèles différents du même espace métrique, et le dictionnaire est donné par la transformation de Cayley. En particulier, tous les résultats vus au sujet du disque de Poincaré se traduisent en des énoncés sur le demi-plan de Poincaré via ce dictionnaire.

### III.4.2 Le groupe $\text{Isom}(\mathbb{H})$

La proposition III.14 nous dit que les espaces métriques  $\mathbb{D}$  et  $\mathbb{H}$  sont isométriques. Cela implique que les groupes d'isométries correspondants sont iso-

morphes. Mais il s'avère que  $\text{Isom}(\mathbb{H})$  est plus facile à déterminer ; c'est d'ailleurs la raison de l'introduction de ce nouveau modèle.

Le théorème que voici nous sera utile pour cette détermination. Toute ressemblance avec le théorème II.9 n'est pas du tout fortuite.

**Théorème III.15**

Tout élément de  $\text{Isom}(\mathbb{H})$  est composition d'au plus trois inversions par des droites hyperboliques.

*Démonstration.* La preuve est essentiellement la même que dans le cas euclidien (théorèmes II.9 et II.11), avec les droites hyperboliques jouant le rôle des droites et les inversions jouant le rôle des réflexions, à cela près que nous n'avons pas d'équivalent de la proposition II.6. Qu'à cela ne tienne, nous avons vu en exercices des preuves "à la main" des résultats correspondants, et ces preuves-là se traduisent parfaitement bien. Nous donnerons ici les étapes de la démonstration, sans détail.

Dans un premier temps, il faut vérifier que si  $z_1$  et  $z_2$  sont deux points fixes de  $\varphi \in \text{Isom}(\mathbb{H})$ , alors  $\varphi$  fixe toute la droite hyperbolique passant par  $z_1$  et  $z_2$ . Via une isométrie hyperbolique, on peut se ramener au cas où  $z_1$  et  $z_2$  ont même partie réelle. Alors, la droite hyperbolique en question est une droite euclidienne verticale. Comme les cercles hyperboliques sont des cercles euclidiens, l'argument de l'exercice 4 (i) de la série 6 s'applique parfaitement. Dans un second temps, il s'agit de vérifier que si  $\varphi \in \text{Isom}(\mathbb{H})$  fixe trois points non-contenus dans une droite hyperbolique, alors  $\varphi$  est l'identité. Ici encore, l'argument l'exercice 4 (ii) de la série 6 s'applique directement. Il suit que l'ensemble des points fixes d'une isométrie hyperbolique est soit vide, soit un point, soit une droite hyperbolique, soit tout le plan hyperbolique, i.e. des sous-ensembles de  $\mathbb{H}$  de dimension  $-1, 0, 1$  et  $2$ , respectivement.

On peut alors procéder par récurrence sur  $i \geq 0$  pour montrer l'affirmation suivante : si la dimension des points fixes de  $\varphi \in \text{Isom}(\mathbb{H})$  est supérieure ou égale à  $2-i$ , alors  $\varphi$  est composition d'au plus  $i$  inversions par des droites hyperboliques. La démonstration est mot pour mot la même que celle du lemme 2.13 (dans le cas  $n = 2$ ), et repose sur l'équivalent du lemme 2.12 que voici, et dont la vérification est laissée en exercice : pour tous  $z_1 \neq z_2 \in \mathbb{H}$ , le lieu des points de  $\mathbb{H}$  à égale distance hyperbolique de  $z_1$  et de  $z_2$  est une droite hyperbolique. Le cas  $i = 3$  de l'affirmation ci-dessus donne le théorème.  $\square$

En particulier, il suit de ce théorème que toute isométrie de  $\mathbb{H}$  est composition d'inversions par des droites hyperboliques. On dira qu'une telle isométrie *préserve l'orientation* si elle est composition d'un nombre pair d'inversions. L'ensemble des isométries de  $\mathbb{H}$  qui préservent l'orientation forme un sous-groupe de  $\text{Isom}(\mathbb{H})$  que l'on notera  $\text{Isom}^+(\mathbb{H})$ . Comme d'habitude, on vérifie facilement que le groupe  $\text{Isom}(\mathbb{H})$  est le produit semi-direct de  $\text{Isom}^+(\mathbb{H})$  et de  $C_2 = \{\text{id}, i_{\mathcal{C}}\}$ , où  $i_{\mathcal{C}}$  est une inversion par une droite hyperbolique quelconque.

Nous sommes maintenant prêts à déterminer les groupes d'isométries  $\text{Isom}(\mathbb{H})$  et  $\text{Isom}^+(\mathbb{H})$ . C'est l'objet du théorème suivant.

**Théorème III.16: Groupe des isométries du plan hyperbolique**

L'application  $\varphi: \text{GL}(2, \mathbb{C}) \rightarrow \text{Möb}^+(2)$  définie par  $A \mapsto \varphi_A$  induit un homomorphisme surjectif  $\text{SL}(2, \mathbb{R}) \rightarrow \text{Isom}^+(\mathbb{H})$  de noyau  $\{\pm I\}$ .

Les remarques faites après l'énoncé du théorème III.7 s'appliquent aussi ici. Tout d'abord,  $\varphi$  définit un isomorphisme entre le groupe des isométries de  $\mathbb{H}$  qui préservent l'orientation et le groupe projectif spécial linéaire de degré 2 sur  $\mathbb{R}$ , à savoir  $\text{PSL}(2, \mathbb{R}) = \text{SL}(2, \mathbb{R})/\{\pm I\}$ . En résumé, le théorème ci-dessus nous dit que l'application  $A \mapsto \varphi_A$  définit un isomorphisme de groupes  $\text{PSL}(2, \mathbb{R}) \simeq \text{Isom}^+(\mathbb{H})$ . Au final, on obtient donc les isomorphismes suivants :

$$\text{Isom}^+(\mathbb{H}) \simeq \text{PSL}(2, \mathbb{R}) \quad \text{et} \quad \text{Isom}(\mathbb{H}) \simeq \text{PSL}(2, \mathbb{R}) \rtimes C_2.$$

*Démonstration du théorème III.16.* Notons d'abord que si  $A \in \text{GL}(2, \mathbb{R})$ , alors  $\varphi_A(\mathbb{R}) \subset \mathbb{R}$ . Comme  $\varphi_A$  préserve les angles et les cercles, l'argument de la preuve du premier point du lemme III.12 montre alors que  $\varphi_A$  est une isométrie de  $\mathbb{H}$ . Par le théorème III.7,  $\varphi_A$  est composition d'un nombre pair d'inversions ; ainsi, c'est un élément de  $\text{Isom}^+(\mathbb{H})$ . Par restriction, on obtient donc un homomorphisme  $\varphi: \text{SL}(2, \mathbb{R}) \rightarrow \text{Isom}^+(\mathbb{H})$ . Le fait que son noyau est  $\{\pm I\}$  se démontre exactement comme dans le théorème III.7. Il ne reste donc plus qu'à vérifier que  $\text{Isom}^+(\mathbb{H})$  est contenu dans  $\varphi(\text{SL}(2, \mathbb{R})) = \varphi(\text{GL}(2, \mathbb{R}))$ , i.e. que toute isométrie  $h \in \text{Isom}^+(\mathbb{H})$  est de la forme  $h = \varphi_A$  pour une matrice réelle  $A \in \text{GL}(2, \mathbb{R})$ . Par le théorème III.15, une telle isométrie est composition de deux inversions par des droites hyperboliques, i.e. par des cercles de  $\hat{\mathbb{C}}$  orthogonaux à  $\mathbb{R}$ . Par le lemme III.9, une telle inversion est de la forme

$$z \mapsto \frac{a\bar{z} + b}{c\bar{z} + d} \quad \text{avec} \quad a, b, c, d \in \mathbb{R}, \quad ad - bc \neq 0.$$

La composition de deux telles transformations donne  $\varphi_A$  avec  $A \in \text{GL}(2, \mathbb{R})$ .  $\square$